

Boekbesprekingen

Redactie:

Arend D. Oosterhoorn
Van Doorne's Transmissie bv
Postbus 500
5000 AH Tilburg

telefoon 013 - 640319

V.N. Lagunov

Introduction to differential games and control theory

Helderman Verlag, Berlin, 1985, 285 pag., ISBN 3-88538-401-9, DM 88.00.

Dit boek is gegroeid uit colleges gegeven door de auteur aan de Universiteit van Novosibirsk en Kalinin. Het heeft een elementair karakter en is didactisch aantrekkelijk.

Het boek bestaat uit drie delen.

In deel I wordt een inleiding gegeven in de optimale besturingstheorie. Dynamisch programmeren en het maximumprincipe komen aan de orde.

In deel II worden spelen in normale vorm en meerstapspelen behandeld. Nash-evenwichten worden ingevoerd. Echter vooral nulsomspelen krijgen veel aandacht.

Differentiaalspelen vormen het onderwerp van deel III. Naast interessante voorbeelden worden verschillende methoden gegeven voor het oplossen van zulke spelen.

Als een introductie in het moeilijke gebied van de differentiaalspelen is dit boek zeker aan te raden.

Voor een meer geavanceerde behandeling is vooral het boek 'Dynamic Nonco-operative Game Theory' van T. Basar en G.J. Olsder (Academic Press, 1982) aan te bevelen.

S.H. Tijs

Mathematisch Instituut KUN

R. Kapadia and G. Andersson

Statistics explained, basic concepts and methods

Ellis Horwood Ltd, Chichester, 1987, 234 pag., ISBN 0-7458-03150-6, £ 12.95.

The book deals with basic statistics and is based on an idea to collect newspaper cuttings relating to statistics. It is a rewritten version in English of an original Swedish book. The publication does not pretend to be an advanced course in statistics but a serious and entertaining introduction to statistics. The book contains 12 chapters. Statistical investigations or exercises are included in each chapter and most chapters also contain computer applications.

The first two chapters deal with basic concepts (variables, discrete, continuous, validity, scales of measurement etc.) and common fallacies in statistics. A special chapter is devoted to official statistics. The next three chapters deal with tables, averages, variability and pictorial representation. Some 11 points are mentioned which should be born in mind when constructing tables. Averages (mode, median and mean), quartiles, variance and standard deviation are treated and illustrated. Several diagrams are shown: bar charts, pie charts, histograms, stem-and-leaf diagrams, frequency and cumulative frequency curves, trend lines, etc.

The chapters 7 and 8 deal with chance, polls and questionnaires, while chapter 11 considers inferences. Notions of chance (classical, frequentist and objective) are explained. Furthermore probability rules, conditional probability, expected value, etc. are treated. Not mentioned are probability distributions, not even the binomial distribution. To solve problems random numbers and simulation are used. Common sampling and non-probability sampling is mentioned and some attention is paid to non-response. The use of confidence intervals for the population proportion p is extensively illustrated. Only some words are said about rejecting and accepting hypothesis and confidence intervals of averages are treated in only three quarters of a page. Nothing is said about the situation where the standard deviation of the population is not known.

Chapter 9 relate to measuring inflation. General indices, Laspeyres (not Laspeyere as is denoted in the book) and Paasche indices are explained. The Laspeyres price index is elucidated on behalf of an example of food prices, using both the direct and the indirect method. The relation between the two

methods however is not established. The principles of correlation and regression are outlined in chapter 10. Also spurious correlation and the variance of the residuals are discussed.

The treatment of the subject-matter of the book is concise. In general in simple terms the heart of the matter is touched. Except for some misprints I mention the following critical notes. With regard to the scales of measurement qualitative data are not only nominal as is told in the book but also may be ordinal. Numbers in parenthesis in a table should be avoided. Indicating class intervals in frequency tables of continuous variables as $90 < 100$, $100 < 110$, etc. is more clear than $90-100$, $100-110$, etc., that is used in the book. In a skew distribution not only a mean is to be used as an average, but sometimes a mode is preferred (income distribution). Nothing is said about frequency density (in histograms). A log scale is mentioned, but can only be understood if one already has some knowledge of it. It is a pity that graphs and tables are not connected with the scales of measurement.

A tree diagram is used to illustrate chances, but is not related with the probability rules. The formulae of the confidence intervals are printed (with a finite population correction) without any explanation or derivation. It seems not well-balanced applying confidence intervals to the proportion of a population or even to differences between the proportions of two populations (with a finite population correction!) without any treatment of probability distributions. Using in consequence random numbers and simulation for problem solving is rather much to ask.

Two alternative formulae for the correlation coefficient are presented, however with no theoretical justification or explanation. Moreover it becomes not quite clear why there are two alternative formulae for the regression coefficient and not for the variance. Most chapters contain computer applications using the statistical package Minitab, which however is very poor for statistical diagrams.

Most chapters conclude with statistical investigations. However, in chapter 3 one is asked to represent some data in a pie chart, graph, line of best fit, bar diagram and to use Spearman's rank correlation coefficient. All of them are treated in later chapters or are treated not at all (Spearman's rank correlation coefficient). The same is the case using the normal distribution and the standard deviation in chapter 4, estimating the probability in chapter 5 and calculating the product moment correlation coefficient in chapter 10.

The book intends to be entertaining on behalf of newspaper cuttings. Indeed many statistical subjects are demonstrated by several pictures taken from newspapers. However these are mainly restricted to subjects that are already suitable for diagrams and pictures like common fallacies, official statistics, tables, pictorial representation and inflation in the chapters 2,3,4,6 and 9. In other chapters there are only few cuttings. As for me it would be better to incorporate even more cuttings or distribute them better amongst the chapters.

I think the book is a useful background to elementary statistics courses, though it is not a book for beginners in statistics. The book can neither be a substitute for an introductory statistics textbook. It is an interesting treatise on statistics to be used for background reading, provided one already has some knowledge of statistics.

H.W.J. Donkers
CBS, Voorburg

M.E. Johnson, ed.

Simulated Annealing & Optimization: Modern Algorithms with VLSI, Optimal Design, & Missile Defense Applications

American Science Press, Inc., New York, 1989, 246 pag., ISSN 0196-6324, \$ 89.75.

Simulated annealing is een enkele jaren geleden geïntroduceerde en sindsdien sterk in de belangstelling staande techniek ter oplossing van combinatorische optimaliseringsproblemen. Een combinatorisch optimaliseringsprobleem wordt gekarakteriseerd door een (meestal) eindige verzameling oplossingen en een kostenfunctie die aan iedere oplossing een reëel getal toekent. Het probleem bestaat dan uit het vinden van een oplossing waar de kostenfunctie haar minimale (of maximale) waarde aanneemt. Een reeds lang bekende techniek om dit soort problemen bij benadering op te lossen staat bekend onder de naam iteratieve verbetering (of: lokale zoekprocedures). Naast een verzameling oplossingen en een kostenfunctie is er dan tevens voor iedere oplossing een buurruimte gedefinieerd als de (deel)verzameling oplossingen die vanuit de gegeven oplossing door middel van een simpele manipulatie (bv. verwisseling) kunnen worden verkregen. Een iteratief verbeteringsalgoritme laat zich dan als volgt omschrijven. Het algoritme bestaat uit een aantal stappen: aan het begin van iedere stap is een oplossing gegeven en genereert het algoritme een nieuwe oplossing, die behoort tot de buurruimte van de gegeven oplossing. Het algoritme kiest dan uit de twee oplossingen die met de laagste kosten. De gekozen oplossing vormt vervolgens het uitgangspunt voor de volgende stap van het algoritme. Het algoritme termineert per definitie in een oplossing, waarvan de buurruimte uitsluitend oplossingen bevat met hogere of gelijke kosten (locaal minimum). Daarmee belanden we direct bij het grootste nadeel van een dergelijke aanpak: de kosten van een lokaal minimum kunnen ver boven die van een globaal minimum liggen. Daar staat tegenover dat de aanpak vrij algemeen toepasbaar is omdat het meestal niet moeilijk is een geschikte buurruimtestructuur te definiëren.

Een van de mogelijkheden het genoemde bezwaar te ondervangen is een van de uitgangspunten te veranderen door de keuze, die tijdens iedere stap moet worden gemaakt, niet langer deterministisch te laten zijn door bv. met een bepaalde kans de oplossing met de hoogste kosten te kiezen. Op deze wijze is het algoritme in staat een lokaal minimum te verlaten.

De zojuist beschreven wijziging leidt tot het simulated annealing algoritme: als de nieuwe oplossing hogere kosten heeft, kiest het algoritme deze oplossing met kans $\exp(-\Delta C/c)$. Hierbij is ΔC het (positieve) verschil in kosten tussen de nieuwe en de oude oplossing en c een controle parameter die tijdens de executie van het algoritme langzaam in waarde wordt verlaagd. Uit de uitdrukking blijkt dat de kans op het kiezen van een slechtere oplossing kleiner is naarmate de verslechtering groter of de waarde van c kleiner is. Dit komt ook tot uitdrukking in de karakteristieken van het algoritme: bij hoge waarden van c genereert het algoritme min of meer een random walk (ook verslechtingen worden met kans ≈ 1 gekozen), bij lage waarden van c verschilt het algoritme nauwelijks met iteratieve verbetering.

De convergentie kan worden geanalyseerd met behulp van de convergentie eigenschappen van Markov keten (de oplossingen die achtereenvolgens door het algoritme worden bezocht vormen een inhomogene Markov keten). Het is onder andere mogelijk te bewijzen dat het algoritme (asymptotisch) naar een globaal minimum convergeert mits de verlaging van de controle parameter 'voldoende langzaam' is. Ten aanzien van het gedrag in eindige tijd zijn helaas weinig analytische of probabilistische resultaten bekend.

Sinds de introductie van simulated annealing in 1982 zijn er enkele honderden artikelen over theorie en vooral over (succesvolle) toepassingen verschenen. Een overzicht van deze artikelen beslaat de eerste helft van het onderhavige boek. Nigel Collins, Richard Eglese en Bruce Golden geven eerst een korte beschrijving van het algoritme en daarna een classificatie van meer dan 300 artikelen over simulated annealing door middel van een veertigtal categorieën. Daarna volgt een geannoteerde bibliografie van deze artikelen, waarin de inhoud van ieder artikel in enkele zinnen wordt beschreven. De bibliografie is redelijk compleet en geeft een goed overzicht van de zeer uiteenlopende gebieden waarin simulated annealing wordt toegepast.

De tweede helft van het boek bestaat uit drie artikelen, waarin toepassingen worden behandeld, en drie artikelen die aan computationele aspecten van deze aanpak zijn gewijd. Tot de eerste categorie behoren een artikel van Diegert over plaatsingsproblemen bij het ontwerp van IC's (het toepassingsgebied waarin simulated annealing de meeste success stories heeft opgeleverd), van Meyer en Nachtsheim over het construeren van optimale ontwerpen en van Bohachevsky, Johnson en Stein over bepaalde militaire toepassingen.

In de tweede categorie vallen de artikelen van Tovey, over bepaalde hulpmiddelen om de executie van het algoritme te versnellen, van Goldstein en Waterman over de invloed van de grootte van de buurruimte op de kwaliteit van de door het algoritme gevonden oplossingen en van Brooks en Verdini, over continue optimaliseringsproblemen.

Alhoewel bruikbaar als naslagwerk heb ik toch mijn twijfels over het nut van dit boekje: het raadplegen van een computerbestand levert al snel een recentere en dus completere bibliografie op dan die van dit boek en de artikelen over toepassingen vormen niet meer dan een aanvulling op de uitgebreide bestaande literatuur. De lezer kan de recensent een plezier doen door voor dezelfde prijs een standaardwerk over simulated annealing (1) aan te schaffen, waarin het algoritme veel uitvoeriger wordt geanalyseerd en veel meer toepassingen worden besproken.

(1) P.J.M. van Laarhoven en E.H.L. Aarts, *Simulated Annealing: Theory and Applications*, Kluwer, Dordrecht, 1987.

P.J.M. van Laarhoven
Nederlandse Philips Bedrijven bv, Eindhoven

O.J. Dunn, V.A. Clark

Analysis of variance and regression (Second Edition)

John Wiley & Sons, New York, 1987, vii + 445 pag., ISBN 0-47-81269-2, £ 33.90.

De auteurs hebben zich met deze uitgave een driedelig doel gesteld:

- a. Volledig tekstboek voor een 1-jarige cursus regressie- en variantie-analyse;
- b. Naslagwerk;
- c. Hulp voor onderzoekers die na een jaartje statistiekonderwijs geavanceerde statistische technieken met behulp van de computer toe gaan passen.

In tegenstelling tot wat de titel aangeeft behandelt dit boek 4 onderwerpen:

1. Inleiding in de statistiek
2. Variantie analyse
3. Proefopzetten
4. Regressie analyse.

Onderwerp 3. wordt behandeld verstrengd met onderwerp 2.

Als voorkennis kan volstaan worden met elementaire algebra, het boek maakt geen gebruik van matrices. De statistiek wordt vanaf het nulnivo opgebouwd, zo worden in de eerste hoofdstukken o.a. het Σ -teken, het rekenkundig gemiddelde en histogrammen uitgelegd. In de loop van het boek neemt de complexiteit toe tot het nivo waarop zaken als kovariantie-analyses en omgaan met gemengde modellen worden behandeld. Het boek legt sterk de nadruk op het netjes opschrijven van het model en geeft veel ellenlange algebraïsche afleidingen om vanuit de modelvergelijking de opsplitsing in kwadraatsommen te vinden. Vervolgens worden de resultaten omschreven naar een voor handrekenmachientjes makkelijker vorm en tenslotte aan de hand van een getallenvoorbeeld ingevuld. Proefopzetten worden verscholen in de hoofdstukken over variantie-analyse behandeld. Bij de hoofdstukken over regressie wordt het verband met proefopzetten grotendeels losgelaten, er wordt nog slechts onderscheid gemaakt tussen data verkregen uit opgezette proeven en uit survey's. Als slottekst van elk hoofdstuk is een paragraaf opgenomen getiteld: 'with the computer', waarvan de inhoud samengevat kan worden als: 'De in dit hoofdstuk beschreven technieken zitten in SAS PROC ... en in programma BMDP...'. Geannoteerde output ontbreekt. Elk hoofdstuk wordt afgesloten met een aantal opgaven (geen uitwerkingen).

Als positief heb ik ervaren: de nadruk op het berekenen van betrouwbaarheidsintervallen en het opnemen van een groot aantal tabellen met verwachtingen van kwadraatsommen, en formules voor de berekening van betrouwbaarheidsintervallen en toetsen. Als negatief: De veelheid van afleidingen die het moeilijk maken de grote lijn vast te houden, de geringe toegevoegde waarde van de computer-paragrafen, en het tussen de regels door behandelen van proefopzetten.

Al met al wekt het boek de indruk dat het oorspronkelijk geschreven is vanuit de optiek dat berekeningen 'met de hand' uitgevoerd worden, en dat deze wat achterhaalde visie aangepast wordt door in ieder hoofdstuk afzonderlijk op te merken dat het ook met de computer kan. Refererend aan de gestelde doelen: a. en c. worden niet gehaald; Als leerboek niet geschikt; te algebraïsch voor de lezer met weinig wiskundige voorkennis, te weinig gebruik makend van matrixrekening voor de wiskundige. b. wordt redelijk gehaald.

Konklusie: geen aanrader

Jan van Gellecum

Volvo Car bv, Helmond

B.Jones & M.G.Kenward

Design and Analysis of Cross-Over Trials

Chapman and Hall, London, 1989, xii + 340 pag., ISBN 0-412-30000-1, £ 27.50

In een cross-over experiment ondergaan groepen individuen meerdere behandelingen in een volgorde die verschillend is per groep. De eenvoudigste proefopzet vergelijkt twee behandelingen, A en B, in twee groepen. In één groep is de volgorde AB, in de andere BA. Omdat de variabiliteit binnen een individu doorgaans kleiner is dan tussen de individuen, verwacht men van een cross-over proefopzet een hogere efficiëntie dan van de proefopzet met maar één behandeling per individu.

De winst van een cross-over proefopzet is gemakkelijk te incasseren als er geen problemen optreden in de vorm van differentiële residuele (of carry-over) effecten; dat wil zeggen als het effect van een behandeling niet afhangt van de behandeling toegepast bij dezelfde persoon in de voorgaande periode(n) van het experiment. Het niet kunnen uitsluiten van zulke effecten leidt tot complicaties, de proefopzet heeft dan veel van de bekende 'can of worms' (once you open a can of worms, the only way to recan it is to use a larger can). Dit verklaart waarom de op het eerste gezicht zo eenvoudige proefopzet relatief vaak in de literatuur opduikt (vijftien bladzijden aan referenties in dit boek). Nu is er voor het eerst een boek geheel aan dit onderwerp gewijd.

Het eerste hoofdstuk (15 blz.) geeft historische achtergronden en introduceert relevante terminologie, modellen en notatie. Het tweede hoofdstuk (73 blz.) gaat over de meest voorkomende vorm van de cross-over proefopzet: de bovengenoemde AB/BA proefopzet (2x2 cross-over design) en wel in het geval van continue responsies. Aan bod komt onder andere grafische presentatie, gebruik van baseline-gegevens, verdelingsvrije analyse, Bayesiaanse benadering. De theorie wordt in detail geadstrueerd met voorbeelden uit de praktijk. Het eerste deel van dit hoofdstuk is ook geschikt voor iemand met weinig ervaring. De auteurs gaan wel een beetje slordig om met schatbaarheidsrestricties op modelparameters. Een meer systematische en expliciete bespreking, bv. in Hoofdstuk 1 waar het probleem van aliasing genoemd wordt, zou op zijn plaats zijn. Dit gebrek aan precisie is wellicht de oorzaak dat van de gepresenteerde toets voor periode effect ten onrechte beweerd wordt dat deze geldig is bij aanwezigheid van een (niet-differentiële) carry-over effect.

Hoofdstuk 3 (51 blz.) veronderstelt dat de responsies binaire of nominale variabelen zijn. Na een beschrijving van gangbare toetsen volgt een modelmatige benadering op basis van log-lineaire modellen voor afhankelijke waarnemingen. De 2x2 proefopzet met binaire responsies krijgt de meeste aandacht. Aan bod komen ook meer gecompliceerde situaties. Zo is er ook een voorbeeld met tri-nomiale responsies, maar de bespreking van de analyse is erg summier (en de analyse bevat een foutief aantal vrijheidsgraden). De interpretatie van parameters van log-lineaire modellen is een vrij delicate zaak. Op dit gebied verschaffen de auteurs niet volledige informatie. Hun omschrijving van het werk van de eerste ondertekenaar van deze bespreking is niet correct; een verwijzing naar gepubliceerde onderlinge correspondentie over dit onderwerp ontbreekt.

Hoofdstuk 4 (49 blz.) gaat over 'Higher order designs for two treatments' en Hoofdstuk 5 (53 blz.) over 'Designs and analyses for three or more treatments'. Behalve een summiere voortzetting van het boven genoemde voorbeeld met trinomiaale responsies werken beide hoofdstukken met uitsluitend continue responsies. Hoofdstuk 4 bespreekt op een systematische wijze, compleet met uitgewerkte voorbeelden, analyse van optimale proefopzetten met 2 tot en met 5 perioden. Hoofdstuk 5 geeft vooral een overzicht en vergelijking van de grote hoeveelheid van proefopzetten in de literatuur. Aan bod komen o.a. variance balanced, partially balanced en control balanced designs. Beide hoofdstukken besteden veel aandacht aan diverse theoretische en praktische aspecten van optimaliteit van een proefopzet.

De analyses in hoofdstukken 5 en 6 zijn gebaseerd op ANOVA-modellen met vaste individu-effecten. Zulke modellen veronderstellen dat waarnemingen binnen een individu onafhankelijk zijn, wat tamelijk onrealistisch is. Gelukkig is het zo dat de verkregen methoden - voorzover gebaseerd op binnen-individu contrasten - correct zijn ook onder de iets meer realistische veronderstelling van gelijkmatige afhankelijkheidsstructuur. In de praktijk zal de correlatie tussen twee metingen vaak groter worden als de tijdsafstand afneemt, zodat er plaats is voor complexere afhankelijkheidsmodellen. Deze problematiek komt aan bod in Hoofdstuk 6 (25 blz.), The analysis of repeated measurements within periods, en in Hoofdstuk 7 (45 blz.), Cross-over data as repeated measurements. Hoofdstuk 6 werkt met vaste individu-effecten, maar veronderstelt dat er binnen elke periode herhaalde metingen beschikbaar zijn die mogelijk gecorreleerd zijn. Twee benaderingen van deze situatie worden geadstrueerd, split-plot-in-time ANOVA en een analyse gebaseerd op 'ante-dependence' modellen voor afhankelijkheid. In Hoofdstuk 7 wordt een algemene covariantiestructuur ook voor de waarnemingen uit verschillende perioden toegestaan en worden de theoretische consequenties hiervan besproken.

Het boek bevat bijna 300 referenties, enkele daarvan echter (nog) niet bibliografisch toegankelijk. Er is een register van auteurs en van onderwerpen; er zijn geen opgaven.

Het verschijnen van dit boek is zeker welkom. Voor de praktisch georiënteerde lezer zullen vooral de hoofdstukken 1,2,4 en 5 direct bruikbare informatie verschaffen. In een nieuwe editie zal wel iets te verbeteren zijn, maar ondertussen kunnen we u dit boek zeker aanraden.

Vaclav Fidler en Hans Burgerhof
Medische Statistiek RUG

Alan F. Karr

Point processes and their statistical inference

Marcel Dekker, Inc., New York, 1989, 504 pag., ISBN 0-824-77513-9, \$ 107.50.

This is a beautiful book-smooth and enjoyable reading, a mine of information, stimulating to the imagination. Statistical inference for stochastic processes and in particular point processes is a vast and rapidly growing area, covering such disparate areas as martingale based inference for intensity based models for counting processes on the real line with its now classical application in survival analysis; empirical process ideas (Vapnik-Cervonenkis classes, Hungarian embeddings) applied to the abstract problem of estimating the distribution of a point process given i.i.d. realizations of it and to special problems concerning Poisson processes; nonparametric estimation of Palm distributions and Markov renewal processes; thinning and Cox processes; state estimation; point process sampling of another (continuous) process; and so on.

The book presents a wide panorama and what unifies it is not so much the view itself as the viewpoint from which it is taken. In this case the view-point is from classical probability theory. The probabilistic structure of various kinds of point processes is carefully and thoroughly developed; the various limit theorems, inequalities and so on are discussed at length and often completely proved. The statistical problems with which the book is nominally concerned appear rather as illustrations of probabilistic limit theorems than as serious practical problems and some of these problems appear to be (in this reviewer's opinion) definitely 'academic' in nature. The book does not systematically address practical data analytic

problems of model choice, complications arising from various types of incomplete observation, computational issues, and so on.

These last remarks are not meant to imply criticism, but just to make it quite clear what the book aims to be and what it does not aim to be. As the author writes in his preface, 'the goal is to present a unified description of point processes before the field becomes so broad that such a goal is hopeless. By depicting the probabilistic and statistical heart of the subject, as well as several of its boundaries, I have endeavored to convey simultaneously the unity and diversity whose interplay and tension I find so alluring'.

A unifying strand to all these different problems is the systematic use of the modern language of point processes. A point process - a random collection of points in some space - is just a certain kind of random measure, counting the random number of points in any given region of space. A marked point process, where each point is accompanied by a random variable taking values in some other space of quantitative or qualitative characteristics, is just a point process on the Cartesian product of the original point space and the mark space. Once we have given the space of measures (possible realizations of our point process) a suitable, natural topology and the corresponding Borel σ -algebra, the distribution of a point process becomes a probability measure on this measurable space, and we can talk about convergence in distribution of point processes using the usual terminology of weak convergence.

So the book is not for the faint hearted. But once the reader has taken the plunge, this abstraction does pay off, unifying and making precise notions which otherwise stay vague. And the author is skilled at gently accustoming the reader to this world.

Despite the fact that this book is now more than three years old (apologies to the reviewer!) it is startlingly up to date, many subjects still receiving here their first accessible book treatment. What I especially like about the book is its completeness (in a technical sense); for instance, not only are Palm process used and even estimated here, but their definition and properties are nicely and fully introduced; similarly, the famous Lenglart inequality and Rebolledo's martingale central limit theorem, which play such an important role in intensity-based models for counting processes, are not just used and explained, but also proved in just a few paragraphs (albeit a sketch), which vividly bring to life the ideas behind them and the meaning of the conditions.

Obviously, not every possible family of point processes is studied in this book - the most striking omission for me being that of Markov point processes. However this book is most valuable, both as a reference work and as a starting point, for anyone interested in the statistical analyses of point processes.

Richard D. Gill
RUG

J.J. Hox en G. de Zeeuw (Redactie)

De microcomputer in sociaal- wetenschappelijk onderzoek

Swets & Zeitlinger, Lisse, 1988, viii + 248 pag. ISBN 90-265-0965-0, f 37.50.

Het boek vormt de neerslag van een wintercursus van de faculteit der Pedagogische, Andragogische en Onderwijskundige Wetenschappen van de Universiteit van Amsterdam. In het boek willen de auteurs laten zien in welke vormen het gebruik van de microcomputer zich manifesteert in sociaal-wetenschappelijk onderzoek. Het richt zich met namen op studenten en afgestudeerden in de sociale wetenschappen. In 13 hoofdstukken van verschillende auteurs wordt op de diverse aspecten nader ingegaan. Deze hoofdstukken kunnen los van elkaar worden gelezen.

In hoofdstuk 1 plaatst G. de Zeeuw het gebruik van microcomputers in het perspectief van de sociale wetenschappen. Naar analogie met allerlei andere uitvindingen uit de geschiedenis wordt uitgelegd dat de microcomputer niet alleen leidt tot verbetering van bestaande concepten, maar ook tot nieuwe ideeën. De computer is bij uitstek het instrument om ongestructureerde data om te zetten in bruikbare informa-

tie, door het kiezen van een geschikte representatie. Aldus draagt de microcomputer bij tot kennisvergroting.

Na dit wat abstracte verhaal gaat G.F. Heyting in hoofdstuk 2 in op het proces van wetenschappelijke verslaglegging. De rol van de micro is hierbij meer dan alleen die van een geavanceerde schrijfmachine. De micro biedt door het manipuleren met blokken tekst, de mogelijkheid van een modulaire aanpak. Informatie kan worden opgebouwd door het ordenen of herordenen van stukken tekst. Met name in de eerste fase van het onderzoek, bij het verzamelen van ideeën, vondsten en literatuurverwijzingen kan dit een krachtig hulpmiddel zijn. Gewone tekstverwerkers kunnen hierbij al nuttig zijn, maar ook heel handig is een outline-programma voor het aanbrengen van structuur in tekstuele informatie. Helaas wordt niet concreet ingegaan op beschikbare programma's hiervoor.

In hoofdstuk 3 beschrijft P. Debets een aantal programma's voor statistische analyse op de microcomputer. Hij schetst de steeds groter wordende rol van de microcomputer in alle fasen van het proces van statistische informatieverwerking. Genoemd worden ondermeer computer assisted interviewing, data entry en analyse. Na een stukje geschiedenis worden enkele programma's voor statistische analyse wat nader bekeken. Dit overzicht is zeer summier en bevat zeker niet alle gangbare en bekende pakketten. Ook is het jammer dat niet nader wordt ingegaan op het verschil tussen het gebruik voor exploratief en inductief onderzoek. Ook hadden de grafische aspecten wat meer nadruk mogen hebben. Al met al helpt de beschrijving niet echt bij het maken van een keuze voor een geschikt pakket, maar dat was misschien ook niet de bedoeling.

Hoofdstuk 4 is gewijd aan het verzamelen van gegevens met de computer. E.D. de Leeuw gaat met name in op computergestuurd telefonisch interviewen (CATI), terwijl kleine uitstapjes worden gemaakt naar computergestuurd persoonlijk interviewen (CAPI) en computergestuurd zelf invullen van vragenlijsten (CASAQ). Belangrijke voordelen van computergestuurde dataverzameling worden genoemd (controle/correctie-mogelijkheid en computergestuurde route). Vergeten wordt dat het verwerkingsproces van de gegevens aanzienlijk wordt versneld, omdat de gegevens tijdens het interview worden ingebracht in de computer. Ook maakt deze vorm van interviewen zaken als random routing en random response order mogelijk, wat tot kwalitatief betere resultaten kan leiden. Er wordt niet ingegaan op de overeenkomsten (volgens mij zeer veel) en verschillen (zeer weinig) in de programmatuur voor CAPI of CATI. Er wordt een lans gebroken voor CATI, maar het zou mij redelijk hebben geleken ook wat problemen te noemen die te maken hebben met telefonisch interviewen (het steekproefkader, identificatie van de interviewer, afhandeling van lange ingewikkelde vragenlijsten).

De micro bij kwalitatief onderzoek is het onderwerp van hoofdstuk 5. G. van den Wittenboer verstaat hierbij onder 'kwalitatief' het onderzoek van niet-numerieke gegevens die niet in een datamatrix kunnen worden ondergebracht ten behoeve van statistische analyse. Ook wordt geen generalisatie naar een populatie nagestreefd. Wat dan wel met dit onderzoeksmateriaal kan worden gedaan, wordt uitgesplitst in twee categorieën. In de eerste plaats is dat de gefundeerde theoriebenadering. Hierbij kan de micro goede diensten bewijzen als tekstverwerker, maar meer valt te bereiken met programma's waarbij trefwoorden aan tekstfragmenten kunnen worden toegekend, en met relationele databanken. Ten onrechte wordt hierbij de indruk gewekt dat dBase II zo'n relationele database zou zijn. In de tweede plaats kan kwalitatief onderzoek worden uitgevoerd via interactieve gegevensverzameling. Hierbij worden bewerkte gegevens gebruikt om nieuwe informatie aan belanghebbenden of deskundigen te ontlocken. Ook hierbij speelt database programmatuur een belangrijke rol.

Ook hoofdstuk 6 gaat over kwalitatief onderzoek, al omschrijft J.J. Hox het hier als onderzoek van gegevens die op ongestructureerde wijze zijn verzameld. In dit hoofdstuk wordt de analyse beschreven van de antwoorden op een open vraag in een enquête naar wat ouders met de opvoeding van hun kinderen beogen. De nadruk ligt meer op de gevolgde analysemethode dan op de wijze waarop de computer is toegepast. Dat kan ook moeilijk want volgens de auteur is er nauwelijks geschikte programmatuur hiervoor beschikbaar. Er bestaat de behoefte aan een database-programma waarin ongestructureerde teksten kunnen worden opgeslagen. Via frequentietellingen van woorden komt de onderzoeker tot het opstellen van een lijst van trefwoorden. Zijn de trefwoorden eenmaal vastgesteld, dan kunnen deze een uitgangspunt vormen voor een meer gestructureerde kwantitatieve analyse.

In hoofdstuk 7 behandelen P. de Greef, J. Breuker en B. Wielinga kennisverwerving voor het bouwen van expertsystemen. Ze beschrijven de KADS-methodologie, die elementen bevat uit de methoden die

worden gebruikt in de conventionele software-ontwikkeling (analyse, technisch ontwerp, implementatie, testen en onderhoud). KADS richt zich met name op de analyse van de problematiek en onderscheid daarbij dataverzameling (bij experts), modelvorming door exploratie van de verzamelde gegevens, en toetsing van het opgestelde model aan de data. Het is een vrij abstract verhaal dat niet wordt geïllustreerd aan de hand van een voorbeeld uit de sociale wetenschappen, maar aan de hand van een analyse van wijnbereiding. De computerspecten komen niet zo erg aan de orde. Alleen aan het einde van het verhaal worden de KADS power tools even genoemd als hulpmiddel bij het analyseren van grote hoeveelheden verbale data.

M. Masuch en R. Verhorst beschrijven in hoofdstuk 8 computersimulatie van sociale systemen. Na een bespreking van de verschillende functies van simulatie (educatie, kennisvergroting en toetsing van modellen), concentreren de auteurs zich op een specifiek voorbeeld, een model voor besluitvormingsprocessen in een organisatie. Terecht wijzen zij erop dat de waarde van dergelijke simulatiemodellen valt of staat met de validiteit van het geïmplementeerde model. De rol die de computer speelt in simulatie is overduidelijk, maar toch wordt hier maar weinig over verteld. Wel besluiten de auteurs met het ventileren van de opvatting dat microcomputers veelal niet krachtig genoeg zijn voor dit soort applicaties. Dit is een opvatting die op zijn minst twijfelachtig aan het worden is met de groeiende kracht en mogelijkheden van de micro's.

In hoofdstuk 9 bespreekt H. Koppelaar hoe kennissystemen kunnen worden toegepast bij sociaal-wetenschappelijke theorievorming. Hierbij gaat het om computersystemen waarin menselijke kennis wordt opgeslagen, en met behulp van inferentiemethoden kan deze kennis gericht worden gebruikt om problemen op te lossen. Het hoofdstuk gaat met name in op het gebruik van heuristische representatie en een logische redeneermethode voor de verheldering en verbetering van theorieën in de sociale wetenschappen. De auteur illustreert zijn verhaal aan de hand van een tweetal concrete voorbeelden waarbij een theorie wordt doorgerekend met het programma QUASIL, dat gebruik maakt van fuzzy algebra technieken.

In hoofdstuk 10 geeft P.T. van de Hoef een elementaire beschrijving van de microcomputer en zijn werking. J.J. Hox gaat in hoofdstuk 11 in op programmatuur en behandelt daarin met name tekstverwerking, spreadsheets, bestandsbeheer, en geïntegreerde pakketten. Bij het onderwerp programmeren wordt een lans gebroken voor gestructureerd programmeren. Verder wordt ingegaan op hulpprogramma's en besturingssystemen. En in hoofdstuk 12 geeft M.D. Lewis een uiteenzetting over de voor- en nadelen van het gebruik van netwerken van microcomputers.

In het laatste hoofdstuk bespreken J.J. Hox en G. de Leeuw de toekomstige ontwikkelingen. De kracht en de mogelijkheden van de microcomputers zullen toenemen, en daardoor zullen de toepassingsmogelijkheden in de sociale wetenschappen worden vergroot. En een verantwoord gebruik van deze nieuwe mogelijkheden zal kunnen leiden tot een verbetering van de kwaliteit en uitkomsten van dit soort onderzoek.

Het boek maakt zijn titel niet in alle hoofdstukken even waar. Soms ligt de nadruk meer op de inhoudelijke aspecten van het onderzoek dan op het specifieke gebruik van de microcomputer daarbij. Ook valt de grote mate van aandacht op die wordt gegeven aan kwalitatief onderzoek; het kwantitatieve onderzoek komt er eigenlijk maar bekaaid af. Niettemin geeft het boek toch een breed spectrum van de toepassingsmogelijkheden, maar het laat ook zien dat er op dit gebied nog veel werk valt te verrichten. Uiteraard moet worden opgemerkt dat het boek slechts een momentopname geeft. De ontwikkelingen snellen voort, zodat op het moment van publicatie een aantal zaken eigenlijk alweer verouderd zijn.

J.G. Bethlehem
Centraal Bureau voor de Statistiek

F.H.C. Marriott

A dictionary of statistical terms, fifth edition

Longman Group, Essex, 1990, vi + 223 pag., ISBN 0-582-01905-2, f 25.00.

In 1951 begonnen Kendall en Buckland op uitnodiging van de International Statistical Institute aan het samenstellen van dit 'woordenboek van statistische termen'. De eerste editie werd gepubliceerd in 1957 en bevatte ongeveer 1700 termen. Vrij snel (1960) verscheen een tweede uitgave en de derde uitgave van 1971 was gegroeid tot ongeveer 2500 begrippen. Na de vierde editie van 1982 (ruim 3000 uitdrukkingen) is nu de vijfde uitgave verschenen met een uitbreiding van meer dan 400 termen ten opzichte van de vorige. De meeste omschrijvingen zijn gelijk aan die uit de vorige uitgave.

Dit is de eerste editie onder de auspiciën van Marriott. Deze editie, zoals ook al de eerdere twee uitgaven, bevat de termen alleen in de Engelse taal opgenomen. De eerste twee uitgaven bevatten ook nog een samenvatting van de equivalenten termen in het Frans, Duits, Spaans en Italiaans.

De omschrijvingen van de begrippen zijn over het algemeen kort, zo rond de vijftig woorden. Termen welke elders in het boek nader worden beschreven staan in de beschrijvingen vet gedrukt. Ter illustratie de volgende begrippen (font-wijziging en vetgedrukte woorden als in de tekst)

average article run length The **average run length** is the average number of *samples* taken before a signal showing the need for action is given. The average article run length is the average number of *items* sampled before action is taken.

average run length (ARL) The average run length of a sampling inspection scheme at a given level of quality is the average number of samples of *n* items taken in the period between the time when the process commences to run at the stated level and that at which the scheme indicates a change from acceptable to rejectable quality level is likely to have occurred. This concept must not be confused with **average article run length**.

average sample run length An alternative name sometimes used for **average run length** in order to distinguish it from the concept of **average article run length**.

Overigens gaat dat in minstens één geval fout, want het begrip 'multidimensional scaling', dat vet gedrukt staat in de beschrijving van 'horseshoe effect' is niet meer in het boek terug te vinden.

In veel omschrijvingen staan literatuurverwijzingen, meestal ook met vermelding van bedenker van het begrip. Dat maakt het leuk om gewoon eens door het boek te bladeren om een indruk te krijgen van wie wat wanneer heeft voorgesteld. Het boek bevat geen lijst met referenties. Dat is begrijpelijk als je bedenkt dat wel opnemen van de aangehaalde publicaties het boek in omvang zou kunnen verdubbelen en dat het niet onder de doelstelling valt, maar het zou toch een luxe zijn als dat wel was opgenomen. Bovendien iedere pagina staan de eerste drie letters van het eerste en het laatste begrip welke op de betreffende pagina worden behandeld en dat maakt het boek goed toegankelijk.

Zomaar twee pagina's leveren de volgende begrippen op

Laplace's theorem, law of large numbers, Laspeyres' index, Laspeyres-Konyus index, latent root, latent structure, latent variabele, Latin cube, Latin rectangle, Latin square,

en

probability moment, probability paper, probability ratio test, probability surface, probable error, probit, probit analysis, probit regression line, procedural bias, process average fraction defective, process with independent increments, processing error, Procrustes methods, produces's risk, product moment, product moment correlation

Het boek is (natuurlijk) niet compleet. Het staat ook nergens dat het de intentie van de auteur is om compleet te zijn. Zo ben ik niet tegengekomen 'partial least squares', 'exponentially weighted moving average chart' en ook geen HOMALS, PRIMALS en andere ALS-en. Maar toch is het leuk om te gebruiken en leerzaam om doorheen te bladeren. De prijs maakt het ook aantrekkelijk.

Arend Oosterhoorn
Van Doorne's Transmissie bv

Gregory M. Constantine

Combinatorial Theory and Statistical Design

John Wiley & Sons Ltd., New York, 1987, xiv + 472 pag., ISBN 0-471-84097-1, £ 48.60.

Dit boek gaat over discrete wiskunde, zoals die kan worden toegepast in statistische proefopzetten. Het boek is een weinig voorkomend voorbeeld van de combinatie van twee vakgebieden, waarbij aan het eerstgenoemde vakgebied en aan de interactie van beide ruime aandacht wordt besteed. Bij de lezer wordt echter een goede kennis van het tweede vakgebied, de statistische proefopzet, aanwezig geacht. Het is zeker geen gemakkelijk boek, hetgeen onder andere moge blijken uit een flinke hoeveelheid gebruikte groepen-theorie. De uitspraak van de auteur dat 'Part of the material will be usefull to experimenters (engineers and managers) who are concerned with the actual planning of statistical studies' lijkt mij, zeker voor de manager, aan de optimistische kant.

Het boek is een zeer duidelijke illustratie van het feit dat eenvoudig ogende problemen (die in het gehele boek steeds terugkomen) diepe wiskundige resultaten vereisen wanneer men ze wil oplossen.

Het boek omvat 9 hoofdstukken, een literatuur-overzicht bij elk hoofdstuk, 3 appendices en een onderwerpen-index. Het boek bevat de bij dit onderwerp gebruikelijke (goed verzorgde) illustraties. Aanwijzingen en antwoorden bij de opgaven staan achterin het boek vermeld.

De hoofdstukken zijn: Ways to Choose / Generating Functions / Classical Inversion / Graphs / Flows in Networks / Counting in the Presence of a Group / Block Designs / Statistical Designs / Mobius Inversion.

De appendices gaan over: Finite Groups / Finite Fields, Vector Spaces, and Finite Geometries / The Four Squares Theorem and Witt's Cancellation Law.

Het boek is geschreven in een 'engaging conversational style' en begint eenvoudig met het tellen van mogelijkheden, maar die eenvoud verdwijnt zeer snel. Het boek is sterk formule-georiënteerd en geeft een overdaad aan resultaten. Naast de formele afleiding van formules besteedt de auteur echter ook steeds aandacht aan de intuïtieve en didactische aspecten van de afleiding.

Het boek is door zijn structuur zeker niet geschikt als naslagwerk. De proefopzetter vindt er geen pasklare receptuur voor conventionele of onconventionele experimentele situaties. Hij zou zich door een veelheid van definities, notaties, specialistische terminologie en afspraken moeten heenworstelen. Naar mijn smaak is het boek eigenlijk alleen toegankelijk voor de goed geschoolde wiskundige, die ook geïnteresseerd is in statistische proefopzetten. Voor die (kleine?) groep liefhebbers heeft het dan wel heel veel te bieden.

J.J. Dik
Volvo Car B.V.

Computer Intensive Methods for Testing Hypotheses, An Introduction

John Wiley & Sons, New York, 1989, ix + 229 pag. ISBN 0-471-61136-0, £ 19.40.

Dit boek beschrijft een aantal technieken om met de computer overschrijdingskansen uit te rekenen van toetsingsgrootheden. De verdeling van de toetsingsgrootheden wordt niet op theoretische gronden afgeleid, maar door voor een groot aantal deelverzamelingen van de waarnemingen de toetsingsgrootheden uit te rekenen. De auteur voert twee belangrijke voordelen aan voor het gebruik van deze computer-intensieve methoden boven de klassieke toetsingstheorie. In de eerste plaats vereisen veel klassieke toetsen veronderstellingen omtrent de vorm van de verdeling, waarbij uiteraard de normale verdeling de meest populaire is. En in de tweede plaats gaat de klassieke theorie er altijd vanuit dat de steekproef uit onafhankelijke waarnemingen bestaat, en dat hoeft bij de computer-intensieve methoden lang niet altijd.

In het boek wordt een drietal technieken onderscheiden: randomizatietechnieken, Monte Carlotechnieken en Bootstraptechnieken. Na een inleidend hoofdstuk wordt in hoofdstuk 2 begonnen met de randomizatietechnieken.

Deze technieken worden toegepast voor toetsen op samenhang tussen (groepen van) variabelen. Van bijna elke relevant toetsingsgrootte kan hiermee de overschrijdingskans worden berekend. Er worden geen veronderstellingen gemaakt omtrent de aard van de verdeling, en de steekproef hoeft ook geen aselechte te zijn. De verdeling van de toetsingsgrootte wordt afgeleid door voor een (groot) aantal gevallen de volgorde van de waarnemingen van de ene variabele te permuteren ten opzichte van de andere variabele, en vervolgens de toetsingsgrootte uit te rekenen. Uit voorbeelden blijkt dat in situaties waarin ook klassieke toetsen kunnen worden toegepast (de t-toets voor het verschil tussen twee gemiddelden, en de toets op correlatie tussen twee variabelen), de randomizatietoetsen daarvoor niet of nauwelijks onderscheiden.

In hoofdstuk 3 worden Monte Carlotechnieken behandeld. Hierbij gaat het om het toetsen of een steekproef afkomstig is uit een nader gespecificeerde verdeling. De verdeling van de toetsingsgrootte wordt afgeleid door voor een groot aantal trekkingen uit de hypothetische verdeling de waarde van de toetsingsgrootheden te berekenen. Een voordeel van Monte Carlotoetsen boven de klassieke toetsen is dat de Monte Carlotoets ook kan worden toegepast in situaties waarin de verdeling van de toetsingsgrootte onder de nulhypothese niet, of niet volledig, is gespecificeerd. Nadeel van zowel de klassieke als de computer-intensieve benadering blijft dat bij verwerking van de nulhypothese het niet duidelijk is of dat komt omdat de nulhypothese inderdaad niet juist is, of omdat niet aan de modelveronderstelling is voldaan. In voorbeelden toont de auteur aan dat het onderscheidingsvermogen van Monte Carlotoetsen vrijwel gelijk is aan dat van (optimale) klassieke toetsen, zelfs al is voldaan aan de modelveronderstellingen van de klassieke toets.

Hoofdstuk 4 is gewijd aan bootstraptechnieken. Deze kunnen worden gebruikt voor het toetsen van hypothesen omtrent de verdeling waaruit de steekproef is getrokken. Met name wordt ingegaan op technieken om te toetsen of de verwachting gelijk is aan een gegeven waarde. De verdeling van de toetsingsgrootte wordt bepaald door het uittrekken van de toetsingsgrootte voor een groot aantal trekkingen van een steekproef met een teruglegging uit de oorspronkelijke steekproef. Ook hier kan weer als nadeel worden aangevoerd (van zowel klassieke als computer-intensieve technieken) dat de nulhypothese ook verworpen kan worden doordat niet aan de gemaakte modelveronderstellingen is voldaan. Daarom hebben bootstraptoetsen eigenlijk alleen maar voordelen in situaties waarin met klassieke methoden de verdeling van de toetsingsgrootte niet kan worden bepaald.

In hoofdstuk 5 worden alle technieken nog eens samengevat en vergeleken met de klassieke technieken. Er vallen twee situaties aan te wijzen waarin de klassieke toetsen het winnen van de computer-intensieve technieken. In de eerste plaats, als zekerheid bestaat dat aan de modelveronderstelling van de klassieke toets is voldaan, dan zijn ze in het algemeen de beste. En in de tweede plaats kan voor hele grote steekproeven het rekenwerk voor de computer-intensieve methoden wel eens te omvangrijk en dus te

landurig worden. Bovendien spelen bij grote steekproeven de onderliggende modelveronderstellingen een veel minder kritische rol.

De appendix van het boek bevat een flinke verzameling computerprogramma's, zowel in Basic, Fortran als (Turbo) Pascal. Het zijn overwegend eenvoudige programma's, die kunnen dienen als matrijzen voor verdere ontwikkeling van programmatuur op dit gebied. Cruciaal bij al deze methoden is de beschikbaarheid van een goede random generator in de gebruikte programmeertaal. De bij de compiler geleverde random generatoren zijn niet altijd even goed. De auteur geeft zelf ook een procedure voor een random generator, maar het nadeel van in een in hogere taal geprogrammeerd exemplaar is uiteraard dat hij langzamer is.

Al met al is het een aardig boekje, waar het vast leuk mee spelen is op de (micro)computer. En met het toenemen van de kracht van de computer zal er steeds meer empoloi zijn voor deze reken-intensieve methoden.

J.G. Bethlehem
Centraal Bureau voor de Statistiek

Sudhakar Dharmadhikary en Kumar Joag-dev.

Unimodality, Convexity, and Applications

Academic Press Inc., New York, 1988, xiii + 278 pag., ISBN 0-12-214690-5, \$ 64.50.

Dit is een degelijk en goed leesbaar boek over de onderwerpen in de titel. Het is geschreven met het oog op toepassingen in de kansrekening en statistiek, maar het is vooral ook een wiskunde-boek: de behandeling is volledig en precies. Merkwaardigerwijs bevat het boek geen definitie van het begrip convexe functie; het begrip komt ook niet in de index voor. Men zoekt ook tevergeefs naar de ongelijkheid van Jensen (wel in het voorwoord). Het boek gaat in feite over verdelingsfuncties.

Hoofdstuk 1 behandelt ééndimensionale verdelingen met onder andere de drie- σ -regel en 'mean-median-mode'. In hoofdstukken 2 en 3 worden generalisaties van het begrip unimodaal besproken; hier werkt het aantal begrippen en de relaties ertussen soms verwarrend. Hoofdstuk 4 behandelt diskrete verdelingen en geeft recente resultaten over het niet-noodzakelijk-unimodaal-zijn van hoge orde convoluties. In hoofdstuk 5 wordt onder meer de door fouten gekweld geschiedenis verteld over de unimodaliteit van zelf-ontbindbare verdelingen. De korte hoofdstukken 6 en 7 geven verbanden aan tussen unimodaliteit, vormen van onafhankelijkheid en order relaties tussen kansverdelingen. De hoofdstukken 8 en 9 geven toepassingen van unimodaliteit en convexiteit in de statistiek (likelihood-functie, power-functie) en in de 'reliability' (IFR- en IFRA-verdelingen).

Er zijn 4 korte appendices over: convexe verzamelingen, convergentie van kansmaten, convexe verzamelingen kansverdelingen, en de ongelijkheid van Brunn-Minkowski.

Het boek besluit met een literatuurlijst van ca. 160 titels (vele van na 1980), een auteursregister en een onderwerpenregister.

Het is een degelijk, verzorgd uitgegeven boek, dat een schat van, dikwijls recente, informatie bevat. Natuurlijk mist men sommige zeer speciale onderwerpen, maar over het algemeen lijkt het boek tamelijk volledig. Het bevat onder meer de oplossing van een recente opgave in Statistica Neerlandica! Van harte aanbevolen voor alle kansrekenaars en mathematisch statistici.

F.W. Sleutel
TU Eindhoven

Experimental design in biotechnology

Marcel Dekker, Inc., New York, 1989, xv + 259, ISBN 0-2847-7881-2, \$ 107.50.

Dit boek heeft een beetje een verkeerde titel. Het zou beter zijn de titel zodanig om te vormen dat er tot uitdrukking kwam dat het boek gaat over probleem oplossen middels het gebruik van statistische proefschemata's, met voorbeelden uit de biotechnology. Het boek is geschreven voor onderzoekers die niet goed thuis zijn in de statistische proefopzetten, maar onderzoek willen doen met gebruik van kleine, efficiënte proefschemata's. In het boek wordt niet uitgebreid ingegaan op analyse methoden als de variantie analyse. Het begrip kwadraatsom komt in het boek niet voor. De opzet is de lezer in te voeren in de principes van netjes experimenteren met gebruik van, meestal verzadigde, schemata's voor het verkrijgen van informatie over veel factoren. De basis gedachte daarbij is dat het doen van onderzoek start met een situatie dat er veel kandidaat-invloedrijke factoren zijn maar dat er weinig kennis is omtrent welke precies. Veel factoren beproeven in zo weinig mogelijk proeven is dan een efficiënte mogelijkheid om te komen tot meer kennis over de werkelijk belangrijke factoren, vaak minder in aantal. Met deze kennis wordt dan naar een optimum gezocht en wordt een verificatie experiment uitgevoerd.

De bovenstaande filosofie wordt in dit boek uitgedragen. De aandacht van het verhaal wordt niet afgeleid door de mathematische beslommingen van de analyse van de experimenten, daarvoor wordt verwezen naar de bestaande literatuur. De resultaten worden allemaal gepresenteerd op grafische wijze. Daaronder valt dan het op normaal waarschijnlijkheidspapier uitzetten van de geschatte effecten, staafdiagrammen van de geschatte effecten om de grootste eruit te halen (Pareto diagram), waarbij wordt verwezen naar de regel van Juran over de 'vital few' en de 'trivial many', en het grafisch weergeven van de gemiddelden op de hoekpunten van de kubus, waarmee de experimentele punten worden aangeduid. Tevens worden interactie plaatjes gemaakt, waarbij het effect van een factor wordt uitgezet bij verschillende niveaus van een andere factor.

In hoofdstuk 1 worden enige begrippen en principes uitgelegd, zoals experimentele ruis en het scheiden van de effecten, het signaal, van de ruis. Daarbij introduceert hij het begrip 'clear signal design', waarmee factoriële proefschemata's met twee niveaus worden bedoeld. In een plaatje laat hij zien dat deze schemata's effectief zijn, maar misleid daarbij de lezer enigszins door een kubus met experimenteerpunten van een 5^3 schema (125 punten) te vergelijken met een Central Composite design voor drie dimensies (15 punten).

Hoofdstuk 2 bevat een voorbeeld studie naar de optimalisatie van in vivo productie van monoclonale antilichamen. Het gehele onderzoeksproces wordt omschreven en het gebruik van een 2^{2^2} schema en de analyse van de resultaten komen aan bod. Voor de optimalisatie van het proces wordt daarna een Central Composite schema behandeld. Hier komt ook, voor het eerst, een regressiemodel om de hoek kijken, maar verder dan de representatie van het resulterende kwadratische model komt het niet. Wel wordt dit model gebruikt voor het maken van een contourplaatje.

In hoofdstuk 3 wordt het visgraatdiagram geïntroduceerd als eenvoudig hulpmiddel bij het identificeren van mogelijk belangrijke factoren. Tevens worden de factoriële schemata's met twee niveaus behandeld, met begrippen als interactie, fractioneren, verstrengeling van effecten en de resolutie van een schema. Uitvoeren van de proefschemata's in blokken komt hier niet aan de orde. Dat wordt omzeild door te wijzen op de noodzaak de experimenten uit te voeren onder homogene omstandigheden. Het volgende hoofdstuk is gewijd aan de grafische weergave van de resultaten, waarvan hiervoor al melding is gemaakt. Daarbij worden ook 'active contrast plots' behandeld, een soort van Bayesiaanse aanpak. Dit onderwerp komt zomaar uit de lucht vallen en verdwijnt daar na één pagina ook weer in. Er wordt veel te weinig over geschreven om gebruikt te kunnen worden. Tevens staat in dit hoofdstuk de eerste variantie analyse tabel (van in totaal twee), wat overigens een combinatie is van variantie analyse tabel en een regressietabel met overzicht van geschatte effecten en t-waarden. Als belangrijke grootheden worden de t-waarden, de overeenkomstige p-waarden en de 'R-squared' behandeld. De rol van de p-waarde wordt enigszins, en terecht, gerelativeerd.

In hoofdstuk 5 wordt aangegeven hoe je moet komen tot de keuze van een proefschemata voor een bepaald probleem. In het boek zijn 45 schemata's opgenomen, aangeduid als de 'Design Digest'. Onder deze

schema's fractionele factoriële schema's op twee niveaus, Plackett-Burman schema's, meestal met centrepunten. Tevens staan vermeld Central Composite schema's en een Box-Behnken schema. De schema's bevatten maximaal 11 factoren uit te voeren in maximaal 32 experimenten. Hoe men tot een keuze uit deze schema's moet komen staat hier vermeld.

Tevens worden in dit hoofdstuk de drie peilers van de statistische proefopzetten behandeld, randomisatie, replicatie en blokvorming. Met deze laatste wordt echter niet actief iets gedaan en dat is een gemis.

Transformeren van de waarnemingen is het onderwerp van hoofdstuk 6 en het is geheel gebaseerd op de Box-Cox familie van transformaties. Tevens komen hier lambda plaatjes aan de orde. Problemen waarbij het toepassen van proefschema's met een factor op drie niveaus van belang is worden aangekaart in het zevende hoofdstuk van het boek. De behandeling is, als in het voorgaande, grafisch van aard.

Een strategie voor het doen van onderzoek wordt in hoofdstuk 8 aangedragen. Definieer eerst het probleem, identificeer de mogelijke invloedrijke factoren en bepaal de respons variabelen. Selecteer het geschikte schema en voer het experiment uit. Analyseer de data, pak de invloedrijke eruit en stel het proces optimaal in. Verifieer de uitkomsten met een vervolg experiment. Behoed je voor het maken van een fout van type III, het juist beantwoorden van een niet op zijn plaats zijnde vraag. Dit relaas wordt uitgebreid uitgewerkt.

Het laatste hoofdstuk bevat weer een uitgewerkt voorbeeld uit de biotechnologische praktijk.

Ik vind het boek waardevol. De boodschap van het netjes experimenteren komt goed over, maar waarom nu precies een bepaald schema moet worden gekozen is misschien niet geheel duidelijk als een niet-ingewijde het boek leest. Het boek is duidelijk geschreven door iemand die op verschillende kortdurende cursussen de stof behandelt, veel willen vertellen maar niet de diepere achtergrond aanbieden. Dat kan goed werken als er vervolg is op die stof, en die wordt aangeboden in de vorm van verwijzingen. De aanpak van de grafische weergave van de resultaten zonder statistisch-technische zaken als variantie analyse tabellen en de daarbij horende grootheden werk goed, vooral in een omgeving waarbij de resultaten gepresenteerd moeten worden aan leken op statistisch gebied. De 'Design Digest', waar ook de alias structuur staat vermeld, is buitengewoon handig.

Arend Oosterhoorn
Van Doorne's Transmissie bv

***** CWI publicaties *****

J.K. Lenstra, H. Tijms en T. Volgenant (eds)

Twenty-five years of operations research in the Netherlands: Papers dedicated to Gijs de Leve

CWI, Amsterdam, tract 70, 1990, 181 pag., ISBN 90-6196-385-0, f 48.00.

Deze bundel bevat bijdragen van verschillende auteurs ter gelegenheid van het 25-jarig jubileum van de Leve als hoogleraar in de operations research. Naast een overzicht van de publicaties door of mede door de Leve zijn de volgende bijdragen opgenomen

Decision support: challenges, opportunities and limitations of operations research

J. Chr. van Dalen

Assignment and shortest path problems

B. Dorhout

Heuristics for the hierarchical network design problem

C. Duin, A. Volgenant

Some approximation approaches in large-scale Markov decision problems: applications to one warehouse multiple retailer systems

A. Federgruen

On the transition probabilities of a Markov process with interventions

A. Hordijk

Teaching linear assignment by Mack's algorithm

R. Jonker, A. Volgenant

Parallel local search for the time-constrained traveling salesman problem

G.A.P. Kindervater, J.K. Lenstra, M.W.P. Savelsbergh

Constructing valid inequalities for combinatorial optimization problems

A.W.J. Kolen

On the differentiability of the set of efficient (μ, σ_2) combinations in the Markowitz portfolio selection method

J. Kriens

Probabilistic analysis of algorithms

A.H.G. Rinnooy Kan, L. Stougie

On the strong connectivity problem

A. Schrijver

Good probabilistic ideas are often simple

H.C. Tijms

Partial contraction in Christofides' lower bound for the traveling salesman problem

A. Volgenant

Stochastic games and the total reward criterion

O.J. Vrieze

Comparing fragmentary and integral production and inventory control concepts

P.J. Weeda

Capacity analysis of automatic transport systems in an assembly factory

W.H.M. Zijm

