

Bockbesprekingen

Redactie

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

telefoon 013 - 640319
 fax 013 - 636590

Besproken boeken in Kwantitatieve Methoden nr. 44

V.P. Godambe (editor)	155
<i>Estimating Functions</i>	
S.T. Thompson	158
<i>Sampling</i>	
B.S. Everitt	160
<i>The Analysis of Contingency Tables Second edition</i>	
Richard M. Heiberger	161
<i>Computation for the analysis of designed experiments</i>	
Rachev, S.T.	162
<i>Probability metrics and the stability of stochastic models</i>	
Lessler, J.T. & W.D. Kalsbeek	163
<i>Nonsampling error in surveys</i>	
John Mandel	165
<i>Evaluation and control of measurements</i>	
Jost Reinecke & Gaby Krekeler (Hg.)	167
<i>Methodische Grundlagen und Anwendungen von Strukturgleichungsmodellen</i>	
Nanda Piersma	168
<i>Combinatorial optimizatoin and empirical processes</i>	

V.P. Godambe (editor)

Estimating Functions

Clarendon Press, Oxford, 1991, xii + 344 pag., ISBN 0-19-852228-2, \$ 35.00

For a statistical problem the statistician models the phenomenon under study such that the sequence $Y=(Y_1, \dots, Y_n)$ of potential observations has a distribution belonging to some specified family, and desires to reach some conclusion concerning the value θ of some functional over this family. Denote the set of its possible values by Θ . An *estimator* of θ is any function of Y he may use as a conjecture (to use J. Bernoulli's terminology) of the value of θ . But clearly not just any function is useful for the statistical problem under consideration. Statistical theories of estimation deals with properties of different functions considered for this purpose and the choice among them.

Estimators in wide use are frequently solutions of an *estimating function* of the type $f(\theta; Y)$ set equal to 0, in which f arises in a natural way. Thus the following methods of estimation all involve minimizing or maximizing a function of θ (and of the observations) and thus can, under suitable conditions, be rewritten in the form $f(\theta; Y) = 0$: Gauss-Markov, least squares, minimum chi-square, and modifications of these; maximum likelihood and modifications thereof. In general θ will be a vector, (in which case f will be vector-valued), but in the following we shall usually confine ourselves to real θ . Denote the unique solution of $f(\theta; Y) = 0$ by $\hat{\theta} = \hat{\theta}(Y)$. There is of course, a one-to-one correspondence between the properties of $\hat{\theta}$ and the properties of f . Nonetheless, as we shall see in some instances below, it often turns out to be useful to consider a class of functions f rather than the corresponding class of solutions of $f(\theta; Y) = 0$.

Often it is reasonable to restrict oneself to unbiased estimating functions, so that not only $f(\hat{\theta}; Y) = 0$ but also $E\{f(\theta; Y)|\theta\}=0$ throughout Θ . (This is certainly much more reasonable than a restriction to unbiased estimators of θ itself; in the present volume, Chapter 6 is devoted to a discussion of this requirement of f .) One may then try to find among such functions one which has the smallest variance. However, if $\hat{\theta}$ solves $f(\theta; Y)=0$, so does the solution of a nonvanishing multiple $k(\theta)f(\theta; Y)=0$; and the mean of the left hand side is still zero, while its variance is $k^2(\theta)$ times the variance of $f(\theta; Y)$. Consequently, if one wishes to minimize the variance of the estimating function, one needs to standardize the latter by division by a nonrandom nonvanishing quantity. For smooth estimating functions of θ , it is reasonable to choose this quantity to be the mean of $\partial f(\theta; Y)/\partial \theta$ when it does not vanish: intuitively, the larger its value (taken positively), the easier it should be to discriminate between a particular value of θ and neighboring ones. (Incidentally, the argument in the book that it does not vanish is mistaken.) One can generalize the foregoing in a number of ways to vector-valued θ ; Chandrasekar and Kale (1984) proved the equivalence of three such ways that are often used.

Consider the following example: the Y_i are independently distributed and have mean $\alpha_i(\theta)$ and variance $\gamma b_i^2(\theta)$, where the α_i and $b_i > 0$ are specified differentiable functions and γ is an unspecified positive number. The (weighted) least squares approach, which requires minimization of $\sum \{Y_i - \alpha_i(\theta)\}^2 / b_i^2(\theta)$ leads to the estimating equations $f_{LS}(\theta; Y) \equiv f(\theta; Y) + B(\theta; Y) = 0$, where

$$f(\theta; Y) = \sum \{Y_i - \alpha_i(\theta)\} \alpha_i'(\theta) / b_i^2(\theta), \quad B(\theta; Y) = \sum \{Y_i - \alpha_i(\theta)\}^2 b_i'(\theta) / b_i^3(\theta),$$

with means 0 and $\gamma \sum b_i'(\theta) / b_i(\theta)$, respectively. The latter does not usually vanish, so that, while under some regularity conditions the solution of $f(\theta; Y) = 0$ will converge in probability to θ , the solution of $f_{LS}(\theta; Y)$ could easily converge to some quantity far from θ . Note that the function f is unbiased, and that in the class of unbiased estimating functions $\{g_d\}$ with $g_d(\theta; Y) \equiv g_{d(\theta)}(\theta, Y) = \sum d_i(\theta) \{Y_i - \alpha_i(\theta)\}$, f actually minimizes the variance of the standardized estimating functions.

Consider now the case in which the family of distributions can be specified in terms of a family of density functions $p(\bullet; \theta)$ with respect to a sigma-finite measure on the observation space; and in which one is concerned with estimating θ , or $\eta = h(\theta)$ for given h . Fisher considered $f(\theta; Y) = \partial \log p(Y; \theta) / \partial \theta$ and called it the score function. Under certain regularity conditions it is unbiased and has variance equal to the negative of the mean of $\partial f(\theta; Y) / \partial \theta$. He called the latter the information on θ in the sequence Y of observations; let us denote it by $I_Y(\theta)$. When it is positive and finite, the variance of the standardized score function is $I_Y(\theta) I_Y^{-2}(\theta) = I_Y^{-1}(\theta)$, which is the lower bound of the Cramer-Rao inequality for standardized estimating functions. Another optimum property of score function is that under weak conditions it is minimal sufficient, whereas a real-valued sufficient statistic in very many cases of interest does not exist. The score function is sometimes modified to lessen the *sensitivity* of the implied estimator to gross errors the data; one then speaks of an M-estimator (and could speak of an M estimating function). Much more complicated considerations arise (both for the theory of estimators and the theory of estimating functions) when θ is a vector, but one is interested in estimating only some of its components. Many Chapters deal with such considerations.

Estimating functions $f^*(\theta; Y)$ whose distribution do not depend on θ were called pivotal quantities by Fisher. Exact and approximate pivotal quantities often lead in an obvious way to *confidence intervals*. Thus, since 1945 the U.S. Bureau of the Census has been constructing confidence intervals for q -iles of a finite population (in particular for the median, which corresponds to $q = 1/2$ using considerations which we shall now describe for the case in which one uses simple stratified sampling: Let a population consist of the couples (j, θ_j) for $j = 1, \dots, N$, where the first N_1 couples constitute stratum 1, ..., the last N_H constitute stratum H . If $g(\theta_j, u)$ is defined as 1 when $\theta_j \leq u$ and 0 otherwise, then $\eta = h(\theta)$ is that u for which $G(u; \theta) = \sum \{g(\theta_j, u) - q\} = 0$, that is, $\eta = G^{-1}(0; \theta)$ with some conventional definition for $G^{-1}(0; \theta)$

when there is a nondegenerate interval of u 's (usually a narrow one) satisfying $G(u;\theta)=0$. One draws a simple random sample S_h without replacement of size $n_h(1 < n_h < N_h)$ from stratum h , independently for each h , and defines

$$f(\eta; S_1, \dots, S_h) = \sum \left(\frac{N_h}{n_h} \right) \sum \{g(\theta_j; \eta) - q\}$$

where the inner sum is over $j \in S_h$ and the outer sum over h . Its mean is 0 and its variance for given η may be estimated in the usual way by

$$\hat{V}\{f(\eta; S)\} = \sum N_h (N_h - n_h) \bar{g}_h(\eta) \{1 - \bar{g}_h(\eta)\} / (n_h - 1),$$

with $\bar{g}_h(\eta)$ the average of $g(\theta_j; \eta)$ over $j \in S_h$. Since η can be (roughly) estimated from $f(\eta; S)=0$ as $\hat{\eta}$, the ratio $\hat{f}(\eta; S) = f(\eta; S) / \hat{V}\{f(\eta; S)\}^{1/2} |_{\eta=\hat{\eta}}$ would, for large n_h and $N_h - n_h$, or for large H , usually have an approximate standard normal distribution. However, if the stratifying variable is highly correlated with θ , many of the $\bar{g}_h(\hat{\eta})$ may vanish and the approximation may become poor. This will also be the case when q is small. (A possible partial remedy one may adopt is to replace vanishing $\bar{g}_h(\hat{\eta})$ by cn_h^{-1} with c equal to $1/2$, say). The resulting confidence interval, of confidence coefficient $1-\alpha$, is then $[\hat{f}^{-1}(z_{\alpha/2}), \hat{f}^{-1}(-z_{\alpha/2})]$, where z_k is the k -ile of the standard normal distribution. (Clearly, a similar approach works when one samples the population in more complicated ways, as long as, for the function of θ_j for which (j, θ_j) is included in the sample, one uses a weight proportional to the probability of including this unit in the sample.) Note that it is not at all clear how an approximate confidence interval could be obtained from the distribution of $\hat{\eta}$ itself, since $\hat{\eta}$ is a highly irregular function of S , unless use is made of resampling. Godambe and Thompson (1986) noted that, if one uses instead of $g(\theta_j; \eta) - q$ some other function $h_j(\theta_j; \eta)$, one can again obtain confidence intervals by this method. For example, for $h_j(\theta_j; \eta) = (\theta_j - \eta a_j) a_j$ (where the a_j are the known numbers), $G^{-1}(0; \theta)$ is the population regression coefficient and a suitable estimator of the variance of the corresponding $f(\eta; S)$ is easily written down. They proved an optimality property of the resulting confidence intervals under the assumption that $\theta_1, \dots, \theta_N$ may themselves be modeled in a certain way. For the case of the q -ile, the model is simply that $\theta_1, \dots, \theta_N$ are independently distributed, each with the same (superpopulation) q -ile, and that the estimating functions have means of the form $\sum k_j(\theta_j; \eta)$, with each summand having 0 mean in the superpopulation.

Possibly the first author to consider explicitly classes of estimating functions was Wilks (1938), who wished to define a class among which the score function divided by $I_Y(\theta)^{1/2}$ gives confidence intervals which are asymptotically of shortest length. (The book refers several times to Wilks (1962) where the proof is clearly incorrect, rather than to Wilks (1938) and to Wilks and Daly (1939).)

The book has some general papers on estimating functions (Chapters 18-22, 24), papers which consider stochastic processes (Chapters 2 and 8-13), and papers involving mixed models, capture-recapture studies, and quadratic exponential models; of course, it is not possible to discuss these in a brief review, but the above remarks will make evident that the book should be useful to workers in a variety of fields.

References

- Chandrasekar, B., and Kale, B.K. (1984). Unbiased statistical estimation functions in the presence of nuisance parameter. *J. Statist. Plann. Inference* 9, 45-54.
- Godambe, V.P., and Thompson, U.E. (1986). Parameters of superpopulation and survey population: their relationship and estimation. *Internat. Statist. Rev.* 54, 127-138.
- Wilks, S.S. (1938). Shortest average confidence intervals from large samples. *Ann. Math. Statist.* 9, 166-175.
- Wilks, S.S. (1962). *Mathematical Statistics*. Wiley, New York.
- Wilks, S.S., and Daley, J.F. (1939) A optimum property of confidence regions associated with the likelihood function. *Ann. Math. Statist.* 10, 225-235.

H.S.Konijn

Tel Aviv University -- Statistics Department

S.T. Thompson

Sampling

John Wiley & Sons, New York, 1992, 343 pag., ISBN 0-471-54045-5, £ 41.95

Enige maanden geleden verscheen het (zoveelste) inleidende boek over steekproeftheorie, simpelweg getiteld 'Sampling'. Auteur is Stephen Thompson, verbonden aan de Pennsylvania State University. Het boek verzorgt in de eerste 150 bladzijden een algemene inleiding en legt in het resterende deel het accent bij het sampelen van 'moeilijk grijpbare' populaties en populaties die gekenmerkt worden door een kleine trefkans.

opzet van het boek

Het boek is onderverdeeld in 6 delen en totaal 26 hoofdstukken. Deel I steekt van wal met een overzicht van de basisprincipes van steekproeven en schatters. Deel II gaat in op het gebruik van secundaire informatie bij schattingstechnieken. Ratio- en regressieschatters passeren hier onder meer de revue. Interessant is dat Thompson een apart hoofdstuk wijdt aan het onderscheid tussen design-based sampling en model-based sampling. Deel III gaat in op iets geavanceerdere technieken, zoals gestratificeerde steekproeven, clustersteekproeven, getrapte steekproeven, double sampling en network-/multiplicity sampling.

In deel IV komt Thompson op onderwerpen die niet tot de standaardmonstering van een boek over steekproeven behoren. Het gaat hier om steekproeven bij moeilijk grijpbare populaties, zoals zoogdieren en vogels. Deel V is gewijd aan spatiële steekproeven. Deze vinden hun toepassingsgebied met name in de geologie en ecologie (bijv. meting van luchtverontreiniging). Deel VI gaat in op adaptive sampling methoden, wat inhoudt dat het sampling design wordt aangepast gedurende het trekkingsproces: afhankelijk van de gevonden waarden op reeds gesampled eenheden wordt het steekproefdesign aangepast. Adaptive sampling vindt zijn toepassing bij populaties waarbij de target elementen schaars verdeeld zijn en geclusterd voorkomen. Dit kan zich bijvoorbeeld voordoen bij fossiele brandstoffen en bij schaarse ziekten.

evaluatie

Thompson begint met de correcte observatie dat steekproeftheorie geconsumeerd wordt door personen uit velerlei disciplines. Een inleidend boek dat zich richt op de volledige markt van geïnteresseerden moet dus van goede huize komen: het vereist dat de auteur rekening houdt met een sterk divers lezerspubliek.

De flaptekst vermeldt dat het boek bedoeld is als naslagwerk voor wetenschappelijk onderzoekers en daarnaast als inleiding voor studenten op een redelijk gevorderd niveau. Het boek lijkt goed aan deze doelgroep tegemoet te komen. De betoogtrant van Thompson is helder, de opbouw van het boek is begrijpelijk en daarmee is het uitstekend geschikt als algemene inleiding. Ook als naslagwerk heeft het voldoende te bieden.

Voor degenen die specifiek geïnteresseerd zijn in Thompson's expertise, namelijk het sampelen van de genoemde moeilijke populaties is het boek ongetwijfeld extra aantrekkelijk.

Verschijsning in de 'Wiley Series in Probability and Mathematical Statistics' betekent dat de tekst door een uitvoerige review-cyclus is gegaan en dat we ons weinig zorgen hoeven te maken over het niveau en de inhoudelijke correctheid van het werk. Dat geldt uiteraard ook voor technische zaken als verwijzingen, inhoudsopgave en register.

W. van Vreden

Instituut voor Telefonische Marktonderzoek, Amsterdam

B.S. Everitt

The Analysis of Contingency Tables Second edition

Chapman & Hall, London, 1992, ix + 164 pag., ISBN 0-412-39850-8, £ 22.50

Dit boekje is een herziening van de gewaardeerde eerste druk, die in 1977 het licht zag (en bracht) en reeds enkele malen een volgende oplage beleefde. Het bevat: chi-kwadraattoetsen (algemeen), de analyse van 2×2 - en $r \times c$ -tabellen, meerdimensionale tabellen, log-lineaire en lineair-logistische modellen en speciale kruistabellen (onder meer: geordende categorieën en vierkante tabellen).

Het boek is in een prettige stijl geschreven en de notaties zijn overzichtelijk. Voorbeelden maken de lezer duidelijk waar de beschreven toetsen kunnen worden toegepast. Er is een omvangrijke literatuurlijst, die ten opzichte van de vorige druk met meer dan 50% is uitgebreid. Het register had iets uitgebreider gekund. Er zijn trefwoorden (of combinaties daarvan) die er zeker bij hadden gekund, en bij sommige trefwoorden bleken te weinig verwijzingen te staan vermeld. Dat is jammer, want zo loopt men bij gebruik van het boek als naslagwerk, de kans het boekje terzijde te leggen, terwijl het gezochte onderwerp er wel in staat.

Ten opzichte van de vorige druk zijn er enkele wijzigingen. Hier en daar is een voorbeeld veranderd en de mogelijkheden die computers tegenwoordig bieden, zijn niet onvermeld gelaten. Het berekenen van exacte overschrijdingskansen en het maken van sommige grafieken zijn enkele van de inmiddels ontgonnen terreinen.

Voor methoden bij kruistabellen met geordende categorieën is er nu een apart hoofdstuk. Er is meer aandacht besteed aan associatiematen voor deze gevallen dan in de vorige druk. Grotendeels nieuw is het hoofdstuk over lineair-logistische modellen.

De layout is hier en daar licht verbeterd, maar ook de vorige druk bood een overzichtelijke aanblik en was goed te lezen. Betreurenswaardig is het feit dat de letter χ , die toch een voorname rol vervult in dit boek, is vervangen door een ordinaire hoofdletter X (behalve in de kop van de tabel). Ook het toevoegen van leestekens (punten en komma's) aan losstaande formules, geniet mijn afkeuring, temeer daar punten in de formules soms ècht een betekenis hebben (bijvoorbeeld n_i). Leestekens zijn er om zinnen beter leesbaar en niet om formules slechter leesbaar te maken. Daar houdt de kritiek wel mee op.

Grote delen van het boekje zijn begrijpelijk voor mensen met een matige voorkennis van de statistiek. Men dient wel de beginselen van de toetsings- en schattingstheorie te kennen. Everitt grijpt echter alle begrippen uit de kast die er bij zijn onderwerpen te pas komen, dus een partieel geïnteresseerde zal hier en daar een fragmentje moeten overslaan. Dat kan, zonder bij het bestuderen van de rest van de tekst gedupeerd te zijn. Aanbevolen, maar niet goedkoop.

J.M. Buhrman

Statistical, Diemen

Richard M. Heiberger

Computation for the analysis of designed experiments

John Wiley & Sons, New York, 1989, 683 pagina's, ISBN 0-471-82735-5

Doelgroep

Dit boek is niet in de eerste plaats bestemd voor gebruikers van programmatuur voor lineaire modellen, maar voor ontwikkelaars hiervan. De uitgever meent volgens de omslag dat hieronder de volgende categorieën vallen: statistici, informatici, ontwerpers van programmatuur, consultants, systeem analisten en studenten in de statistiek, informatica, werktuigbouwkunde, experimentele psychologie en economie. Mij lijkt deze lijst wat te ruim. De meeste mensen uit deze groepen zullen goed kunnen werken met manuals van bestaande pakketten, eventueel uitgebreid met wat standaardwerken over variantie-analyse en proefopzetten. Toch is het wel denkbaar dat ook anderen dan statistische programmeurs iets aan dit boek hebben. Er wordt namelijk op pakket onafhankelijke manier behandeld hoe de algoritmen voor de diverse deeltaken van variantie-analyse in elkaar zitten. Wellicht kan dit bijdragen tot begrip voor onverwachte uitkomsten bij zalen als multicollineariteit, verstrengeling en niet-orthogonaliteit.

Geen Proefopzetten

De titel spreekt terecht van Designed Experiments. Er wordt namelijk vanuit gegaan dat het experiment reeds is ontworpen en de data reeds zijn vergaard. Dit is een belangrijke beperking; er worden geen algoritmen behandeld voor het genereren van bijvoorbeeld fractionele factoriële designs. Maar alles wat men in de analyse achteraf zou kunnen doen, komt ruimschoots aan bod. En dit gaat duidelijk verder dan een verzameling formules voor kwadratsommen, vrijheidsgraden en verdelingsfuncties. Uitgebreid wordt bijvoorbeeld ingegaan op de constructie en interpretatie van uitdrukkingen voor het specificeren van modellen. Dit geldt ook voor hypothesen, alsmede voor de context waarbinnen men deze wenst te toetsen.

Floppy Disk

Bij het boek wordt een floppy disk geleverd voor IBM compatible personal computers. Alle gehandelde programmatuur staat hierop in Fortran, Basic, Apl en C. De voorbeelden zijn uitgevoerd met P-STAT, GENSTAT en SAS. De hierbij behorende programma's en datasets staan ook op de floppy. Het nut van deze voorbeelden ontgaat me een beetje, maar de files in de vier programmeertalen lijken me zeer nuttig voor ontwikkelaars van software.

Notatie

Op sommige plaatsen valt duidelijk te herkennen dat het boek eerder voor de numerieke programmeur dan voor de statisticus is geschreven. Merkwaardig genoeg blijkt de gekozen notatie ook aantrekkelijke eigenschappen voor de interpretatie te hebben. Enkele methoden

van Cochran en Cox (1957) worden bijvoorbeeld op een begrijpelijker wijze gepresenteerd dan in de oorspronkelijke publikatie.

De Auteur

Richard M. Heiberger heeft het ANOVA programma voor P-STAT geschreven en heeft een zomer doorgebracht met de ontwerpers van GENSTAT in het Rothampstead Experimental Station. De auteur weet duidelijk waar hij het over heeft en plaatst de behandelde methoden steeds in een relevante context. Waar bijkomende onderwerpen niet behandeld worden (zoals het generen van de proefopzet) treft men verwijzingen naar andere bronnen aan.

Aanbeveling

De doelgroep is (zoals in de eerste alinea werd opgemerkt) wat beperkt, maar zij die tot deze doelgroep behoren zullen zeker wat aan dit boek hebben. Voor zover ik weet is er geen concurrerend alternatief. Voor regressie bestaat het boek van Maindonald (1984) maar in situaties met uitsluitend of overwegend nominale predictoren is dit minder geschikt.

Jan B. Dijkstra
TUE

Rachev, S.T.

Probability metrics and the stability of stochastic models

John Wiley & Sons, New York, 1991, xiv + 494 pag., ISBN 0-471-92877-1, £ 55.00

Het te bespreken boek representeert het werk van de Russische school in de kansrekening rond Zolotarev met illustere voorgangers als Prokhorov en Kolmogorov. De auteur is Bulgaar, maar gepokt en gemazeld in Moskou. Al verscheidene jaren trekt de auteur ook de westerse wereld door via gastposities aan instituten waar op zijn gebied gewerkt wordt. Door deze achtergrond lijkt de auteur de ideale persoon om een samenvattende monografie te schrijven over afstandsmaten voor kansverdelingen en het gebruik daarvan voor het behandelen van stabiliteitsproblemen bij kansmodellen. Onder stabiliteitsproblemen verstaan we de vraag naar de gevoeligheid van het gedrag van een kansmodel voor kleine perturbaties in de aannamen. Bijvoorbeeld: blijft in een wachtrijmodel de kansverdeling van de doorlooptijd in de buurt van de oude als de verdeling van de bedieningstijd een beetje verandert? Als we hier uitspraken over proberen te doen, worden uiteraard de veranderingen in de bedieningstijdverdeling en de doorlooptijdverdeling beschreven met behulp van de afstandsmaten voor kansverdelingen. Kenmerkend voor de aanpak van Zolotarev voor dit soort stabiliteitsvragen is dat eerst gezocht wordt naar een soort ideale afstandsmaat voor het betreffende type model. Het stabiliteitsprobleem wordt dan opgelost ten opzichte

van deze ideale afstandsmaat. Daarna wordt het resultaat eventueel vertaald naar de gewenste afstandsmaten.

Het boek bestaat uit vier delen. Het eerste deel introduceert verschillende typen afstandsmaten voor kansverdelingen en behandelt wat basiseigenschappen. Deel twee is gewijd aan relaties tussen verschillende typen afstandsmaten. In deel drie wordt ingegaan op het gebruik van de zogeheten minimale afstandsmaten. Een van de onderwerpen die hier aan de orde komen, is de stabiliteit van wachtrijsystemen. Daarnaast worden nogal wat limietstellingen besproken. In deel vier tenslotte komt het gebruik van ideale afstandsmaten voor verscheidene problemen aan de orde. Voorbeelden zijn hier de robuustheid van de chikwadraat toets op exponentialiteit, stabiliteit in risictheorie, stabiliteit in de stelling van De Finetti.

De behandeling in het boek is zeer degelijk en volledig, met veel en goede verwijzingen naar de literatuur op dit terrein. Ook de resultaten van Europese en Amerikaanse onderzoekers zijn goed verwerkt. De behandeling is in het algemeen nogal abstract en technisch. Het boek is daardoor meer voor de potentiële gebruiker. Het boek is zeker een aanwinst omdat het heel veel kennis en ervaring op overzichtelijke wijze bundelt.

J. Wessels

TUE

Lessler, J.T. & W.D. Kalsbeek

Nonsampling error in surveys

John Wiley & Sons, New York, 1992, xiii + 412 pag., ISBN 0-471-86908-2, £ 56.00

Het doel van het boek is het vastleggen en samenvatten van onderzoek dat is verricht naar niet-steekproeffouten. Het is een algemene inleiding, bedoeld voor studenten en beoefenaars van survey-onderzoek. Het geeft inzicht in factoren die de resultaten van surveys kunnen beïnvloeden en biedt, daar waar mogelijk, handvaten om fouten te bestuderen, te voorkomen en/of voor fouten te corrigeren. Volgens de auteurs kan het boek worden gezien als een eerste stap in de ontwikkeling van een survey design informatie systeem, zoals Daniel G. Horvitz in 1978 heeft voorgesteld.

Niet-steekproeffouten worden in het boek geïnclassificeerd volgens een zeer gangbare indeling: 1. kaderfouten; fouten ten gevolge van het steekproefkader, 2. non respons fouten; fouten ten gevolge van het niet meewerken van respondenten aan onderzoek, en 3. meetfouten; fouten ten gevolge van het verkeerd verstrekken of registreren van antwoorden.

Het boek begint in hoofdstuk 1, met een algemene en historische discussie over fouten in survey-onderzoek (zowel steekproef- als niet-steekproeffouten), die wordt opgevolgd door een beschouwing wáár in het onderzoeksproces problemen kunnen ontstaan, in hoofdstuk 2. De overige hoofdstukken, met uitzondering van hoofdstuk 12, waar twee algemene modellen voor fouten worden gepresenteerd, volgen de classificatie zoals hierboven beschreven: kaderfouten (hoofdstuk 3 tot en met hoofdstuk 5), non respons fouten (hoofdstuk 6 tot en met hoofdstuk 8) en meetfouten (hoofdstuk 9 tot en met hoofdstuk 11). Tenslotte is een uitgebreid definitie-overzicht opgenomen van alle gehanteerde begrippen.

De presentatie in de hoofdstukken 3 tot en met hoofdstuk 11 is gelijk. Per foutenbron worden eerst de belangrijkste begrippen gedefinieerd. Vervolgens wordt aangegeven welke effecten bekend zijn uit de literatuur en in welke mate deze effecten optreden, en worden methoden aangedragen om de invloed van deze effecten te meten. Met name de invloed op variantie en bias, ofwel de bijdrage aan het zogenaamde mean-square-error model, wordt behandeld. Tenslotte wordt ingegaan op de wijze waarop voor deze effecten mogelijkerwijs, vooraf of achteraf, kan worden gecontroleerd. De aandacht die uitgaat naar definities, effecten en methoden varieert per foutenbron, afhankelijk van hetgeen tot nu toe bekend is in de literatuur.

Ten aanzien van het steekproefkader worden de volgende fouten behandeld: onder- en overdekking van elementen, het voorkomen van elementen in meer dan één kader, en onjuiste of onvolledige informatie ten aanzien van elementen. Met betrekking tot non respons wordt uitgebreid ingegaan op unit non respons en item non respons. Zowel voor fouten met betrekking tot het kader als voor non respons worden relatief veel methoden aangegeven, waarmee voor specifieke fouten kan worden gecorrigeerd en/of gecontroleerd.

Ten aanzien van meetfouten is, volgens de auteurs, minder bekend in de literatuur. Er wordt meer aandacht besteed aan de statistische modellen die zijn voorgesteld om inzicht te krijgen in de omvang van de fouten, dan er oplossingen worden geboden om voor specifieke fouten te corrigeren en/of controleren.

Het aardige van het boek is dat een uitgebreid en duidelijk overzicht wordt geboden van diverse niet-steekproeffouten in survey-onderzoek, zowel op een praktisch als formeel niveau. Gezien de hoeveelheid literatuur die hieromtrent is verschenen, zal duidelijk zijn dat veel onderwerpen slechts in algemene zin worden behandeld.

Nelly Kalfs

NIMMO, Universiteit van Amsterdam

John Mandel

Evaluation and control of measurements

Marcel Dekker Inc., New York, 1993, xi + 249 pag., ISBN 0-8247-8531-2, £ ?

Direct de eerste paragraaf in het eerste hoofdstuk van het boek wordt al een belangrijke uitgangspunt geformuleerd: *'manufactured products are the end product of a process called a manufacturing process. Similarly, measurements are the end product of a process that we call a measuring process'*. Door het meten te zien als een proces met input, throughput en output, worden heel veel methoden en technieken uit de 'gebruikelijke' omgeving van productie transformeerbaar naar het meetproces.

In de inleiding van het boek heeft de auteur het standpunt ingenomen dat het niet de bedoeling is data uit dergelijke meetprocessen te analyseren tegen de achtergrond van een statistisch model. Volgens hem worden de veronderstellingen welke ten grondslag liggen aan variantie analyse, covariantie analyse en andere technieken niet vervuld door praktijkgegevens. Toch begint het boek met een aantal hoofdstukken over statistische principes en wordt ook verder in het boek rijkelijk gebruik gemaakt van statistische methoden en technieken.

Na een inleidend hoofdstuk worden in hoofdstuk 2 verdelingen en momenten behandeld, compleet met integralen. In hoofdstuk 3 worden zuiverheid en nauwkeurigheid behandeld. Tevens gaat hij in op de centrale limietstelling, de normale verdeling, gewogen gemiddelden en poolen van varianties.

In 'Sources of variability' (hoofdstuk 4) wordt met behulp van de éénweg variantie analyse besproken hoe de binnen groep variabiliteit en tussen groep variabiliteit van elkaar zijn te onderscheiden. De variantie analyse wordt toegepast op variatie binnen en tussen laboratoria.

Hoofdstuk 5 behandelt de lineaire regressie, de gewogen lineaire regressie en de situatie waarbij (meet-)fouten optreden in zowel de x- als de y-variabele. De regressie wordt in hoofdstuk 6 uitgebreid tot meervoudige regressie, overigens zonder gebruik van matrices. Wel komen orthogonale polynomen en collineariteit aan bod.

In hoofdstuk 7 komt de tweeweg variantie analyse aan bod, welke wordt toegepast in hoofdstuk 8. Het voorbeeld is echter niet ontleend aan een toepassing voor het onderzoeken van een meetproces, het is een voorbeeld van een variantie analyse met twee fysische factoren (druk en golflengte). In hoofdstuk 9 draagt de auteur een methode aan die geschikt is in situaties dat niet aan de voorwaarden van lineariteit van het model uit hoofdstuk 7 wordt voldaan. Hij gebruikt daarbij eigenwaarden en eigenvectoren en de singuliere waarden decompositie.

In hoofdstuk 10 keert de auteur weer terug naar het probleem van de tussen-laboratoria studies. Daarbij wordt nu het idee uitgewerkt dan p laboratoria q testprodukten n maal meten. Daarbij komen zuiverheid, herhaalbaarheid en reproduceerbaarheid aan bod. Tevens wordt er stil gestaan bij uitbijter testen. Eén voor het geval dat er een enkele waarneming onder de n waarnemingen een mogelijke uitbijter is (k -statistic) en één voor het geval dat de metingen van een enkel testmateriaal bij een enkel laboratorium mogelijk afwijkend wordt gemeten (h -statistic).

Voor het beheersen van de meetmiddelen in de tijd wordt het gebruik van statistische proces analysekaarten behandeld. Op deze wijze, welke in de kwaliteitszorg reeds ingeburgerd zijn voor het volgen van productieprocessen, kan ook de stabiliteit van de meetprocessen zichtbaar worden gemaakt.

In hoofdstuk 12 wordt gesproken over de vergelijking van meetmethoden als deze methoden niet in dezelfde schaal of eenheid wordt uitgedrukt. Tevens wordt gekeken naar de vergelijking van benodigd aantal waarnemingen per methode om dezelfde precisie te behalen, een argument met economische motieven.

Een klein stuk over het verleden, het heden en de toekomst van de data analyse sluit het boek af. Daarin overpeinst de auteur de rol van de modelvorming. Ook geeft hij daarin duidelijk aan dat het analyseren van een rij getallen zinloos is als je niet weet wat het proces is dat deze rij getallen heeft voortgebracht. Data analyse gaat dus verder dan alleen het analyseren van data.

Het is inderdaad een onorthodox boek als je kijkt naar de behandeling van de onderwerpen. Naar mijn mening had deze behandeling netter gekund als er meer moeite was gestoken in een behandeling waarbij een betere aansluiting was gemaakt met de bestaande statistische principes. Veelal wordt er toch gebruik van gemaakt of een beroep op gedaan. In de inleiding stelt de auteur ook (misschien voorzichtigheidshalve) dat het boek niet geschreven is voor statistici.

Het boek is wel aan te bevelen aan mensen die het meetproces vanuit de statistische kant willen bekijken. Wellicht is het voor onderwijsgevendenden een goede tip om het boek eens te bekijken om de basisgedacht van het meten als een proces over te dragen. Het voordeel is dat methoden en technieken nu worden gepresenteerd alleen in het kader van het meten en dat zorgt voor methodische aandacht voor het meetproces.

Arend Oosterhoorn

Jost Reinecke & Gaby Krekeler (Hg.)

Methodische Grundlagen und Anwendungen von Strukturgleichungsmodellen

FRG, Mannheim, 1993, 179 pag., ISBN 3-924725-05-5, DM 28.00.

Dit is het verslag van de 13e bijeenkomst van de studieclub over het genoemde onderwerp op 2 en 3 april 1992. Bijdragen aan deze bundel zijn

Bedeutung individueller Kennwerte für Optimierung von Strukturgleichungsmodellen

Iris Pieper & Christian Tarnai

Konsumentenzufriedenheit als Determinante der Marken- und Händlerloyalität am Beispiel der Automobilindustrie

Christoph Burmann

Analyse des umweltorientierten Unternehmensverhaltens auf der Grundlage kausalanalytischer Modelle

Manfred Kirchgeorg

Eine Bemerkung zur Effizienz der Schätzung von Probitmodellen oder: Warum ordinale Einflußgrößen stets als gemeinsam abhängig betrachtet werden sollten

Gerd Ronning

Effects of Response Categories on the Reliability and Validity of Life Satisfaction Measurements

Dagmar Krebs & Peter Schmidt

Güte der Schätzer bei Strukturgleichungsmodellen mit mehrstufig ordinalen Variablen

Roland Brandmaier & Harald Mathes

Mode effects in structural modelling; A LISREL multi-group comparison of mail, telephone and face to face survey data

Edith de Leeuw & Joop J. Hox

Erlebte Belastung in der Partnerschaft: Eine multivariate Panelanalyse unter Verwendung eines Multi-Methoden-Ansatzes

Thomas Blank & Jost Reinecke

Arend Oosterhoorn

Nanda Piersma

Combinatorial optimizatoin and empirical processes

Thesis Publishers, Amsterdam, 1993, 138 pag., ISBN 90-5170-211-6, f 37.50

Samenvatting van het proefschrift:

Voor het onderzoek van dit proefschrift worden een aantal problemen uit de praktijk van de besliskunde bestudeerd. De onderzochte problemen hebben als kenmerk dat er een keuze wordt gemaakt van een "beste" of "optimale" oplossing uit een eindige, maar wellicht zeer grote verzameling van alternatieven. Zulke problemen worden aangeduid als combinatorische optimaliseringsproblemen.

Een voorbeeld van een combinatorisch optimaliseringsprobleem is het *Set Covering Problem*. dat wordt bestudeerd in hoofdstuk 3. Een toepassing van dit probleem is de bepaling van de locaties van een aantal brandweerkazernes in een stad. Gegeven zijn een eindig aantal mogelijk locaties, het doel is een zodanige keuze te maken dat de kosten voor de bouw minimaal zijn en dat elke gebouwde kazerne een aangewezen omgeving binnen een vaste, gegeven tijd kan bedienen. De totale bouwkosten vormen de *doelstellingswaarde* van dit probleem.

Voor het Set Covering Probleem wordt in hoofdstuk 3 een schatting afgeleid voor de doelstellingswaarde. In de toepassing zijn dat dus de bouwkosten van de kazernes. Daarvoor hoeft niet daadwerkelijk bekend te zijn welke lokaties worden aangewezen voor de bouw van een kazerne, en zelfs het aantal te bouwen kazernes hoeft niet te worden bepaald. Ook voor de andere combinatorische optimaliseringsproblemen die onderwerp zijn van dit proefschrift wordt de waarde van een doelstellingsfunctie bestudeerd zonder dat de optimale oplossing voor het probleem wordt berekend.

Naast het afleiden van een schatting voor de doelstellingswaarde wordt ook de kwaliteit van deze schatting onderzocht. Een gebruikelijke aanpak is om het asymptotische gedrag van de doelstellingswaarde te bestuderen. Daartoe wordt een reeks doelstellingswaarden beschouwd die correspondeert met een oplopend aantal waarnemingen en dus met problemen van oplopende grootte.

Een eerste kwaliteitscriterium is dat de schatting met overweldigende kans de doelstellingswaarde juist zal weergeven als het aantal waarnemingen oneindig groot wordt. Men spreekt dan van convergentie met kans 1 van de doelstellingswaarde naar de schatting. In dit proefschrift wordt bovendien de snelheid van deze convergentie bestudeerd.

Een tweede kwaliteitscriterium is een bovengrens op de kans dat het verschil tussen de schatting en de werkelijke doelstellingswaarde groter is dan een gegeven constante. In dit proefschrift worden uitspraken gedaan als "de kans dat de werkelijke bouwkosten meer dan f 30.000 afwijken van de geschatte bouwkosten is kleiner dan 5%".

Er worden zogenaamde exponentiële bovengrenzen afgeleid op deze kansen. Deze kansen worden al snel erg klein als het aantal waarnemingen waarop de schatting is gebaseerd stijgt. Opvallend is dat deze kansuitspraken kunnen worden gedaan voor elke probleemgrootte en niet alleen als het aantal waarnemingen zeer groot wordt.

In dit proefschrift wordt een speciale techniek onderzocht om combinatorische optimaliseringsproblemen te bestuderen met behulp van resultaten uit de empirische proces theorie. Deze techniek blijkt uitstekend te werken voor een aantal combinatorische optimaliseringsproblemen. In het algemeen zijn er nogal wat technische details die toepassing moeilijk maken, maar het is aannemelijk dat de methode veel toepassingsmogelijkheden heeft.

> > > > > > > > > **Binnengekomen boeken** < < < < < < < < <

Om de nieuws waarde van de boekbesprekingsrubriek te verhogen wordt een lijst gepresenteerd van boeken die in de afgelopen periode bij de redactie zijn binnengekomen. Omdat er soms nogal wat tijd kan zitten tussen de uitgave van het boek en de publicatie van de recensie, hoopt de redactie hiermee de dienstverlening aan de lezers te vergroten.

Boeken waarvoor een  symbool staat zijn nog niet ondergebracht bij een recensent. Mocht er belangstelling bestaan voor het bespreken van een boek, dan kan contact opgenomen worden met de boekbesprekingsredacteur. Met nadruk wordt gesteld dat niet zeker is of het boek nog 'in de aanbieding' is.

Pukelsheim, F.

Optimal design of experiment

John Wiley & Sons, New York, 1993, xxiii + 454 pag., ISBN 0-471-6, 1971-X, £ 58.00

Barnett, V. & K.F. Turkman (eds)

Statistics for the environment

John Wiley & Sons, New York, 1993, xix + 427 pag., ISBN 0-471-93467-4, £ 49.95

Venugopal, N. (editor)

Contributions to stochastics

John Wiley & Sons, New Dehli, 1992, xvi + 216 pag., ISBN 0-470-22050-3, £ 32.95

Elworthy, K.D. & N. Ikeda ☞

Asymptotic problems in probability theory: stochastic models and diffusion on fractals

Longman Group UK, Harlow, 1993, v + 288 pag., ISBN 0-582-08656-6, £ 36.00

Elworthy, K.D. & N. Ikeda ☞

Asymptotic problems in probability theory: Wiener functionals and asymptotics

Longman Group UK, Harlow, 1993, v + 241 pag., ISBN 0-582-22683-X, £ 26.00

Andersen. P.K, O. Borgan, R.D. Gill & N. Keiding ☞

Statistical models based on counting processes

Springer Verlag, Berlijn, 1993, xi + 767 pag., ISBN 0-387-97872-0, DM 128.0

Johnson, N.L., S. Kotz & A.W. Kemp ☞

Univariate discrete distributions, second edition

John Wiley & Sons, New York, 1993, xx + 565 pag., ISBN 0-471-54897-9, £ 58.00

Tavecchio, L.W.C. & J. Hoogeweij ☞

Oefenboek statistiek, methoden en technieken

Wolters-Noordhoff, Groningen, 1993, 307 pag., ISBN 90-01-83820-2, f 49.50

Houwelingen, P.C. van, Th. Stijnen & R. van Strik ☞

Inleiding tot de medische statistiek

Wetenschappelijke uitgeverij Bunge, Utrecht, 1993, ix + 325 pag., ISBN 90-6348-127-6, f 79.50

Fisher, L.D. & G. van Belle ☞

Biostatistics, a methodology for the health sciences

John Wiley & Sons, New York, 1993, xxii + 991 pag., ISBN 0-471-58465-7, £ 62.00

Anastassiou, G. A. ☞

Moments in probability and approximation theory

Longman Group UK, Harlow, 1993, ix + 411 pag., ISBN 0-582-22770-4, £ 39.00

Campbell, M.J. & D. Machin ☞

Medical statistics, a commonsense approach (2nd edition)

John Wiley & Sons, New York, 1993, xi + 189 pag., ISBN 0-471-93764-9, £ 14.95

Rao, T. Subba (editor) ☞

Developments in time series analysis, in honour of Maurice B. Priestley

Chapman & Hall, Londen, 1993, xxvii + 433, ISBN 0-412-49260-1, £ 49.95

Oproep voor boekbespekers

Als redacteur van de boekbesprekingsrubriek van de VVS, welke gepubliceerd wordt in Kwantitatieve Methoden, ben ik steeds op zoek naar mensen die bereid zijn een pas verschenen boek te willen lezen met het doel daarover een oordeel te geven. Dat oordeel wordt op prijs gesteld door de leden van de VVS. Het verschaft inzicht in de kwaliteiten van het boek. Voor de schrijver(s) kan een oordeel van het publiek waarvoor het boek geschreven is een mogelijkheid bieden tot verdere verbetering. Ook de uitgever heeft er uiteraard baat bij. Als tegenprestatie voor het beoordelen mogen de recensenten de besproken boeken behouden.

Behalve de algemene boeken over statistiek, waar veel mensen zinnige dingen over kunnen zeggen, betreft het vaak ook boeken met een specialistisch onderwerp. Als ik, soms op de gok, iemand een dergelijk boek toestuur, dan hoor ik een enkele keer als reactie dat 'het boek precies aansluit bij het onderzoek' dat de recensent verricht. Dat is prettig, omdat veel mensen daarbij voordeel hebben. De recensent, omdat hij/zij het nieuw verschenen boek op het onderzoeksgebied ter beschikking heeft, en de auteur(s) omdat er een terzake kundig oordeel wordt gegeven.

Om nu de succeskans van het toesturen van een boek te vergroten, wil ik graag in contact komen met veel collega statistici, kansrekenaars en OR-beoefenaren, al zijn de boeken op dit laatste gebied schaars. Stuur het onderstaande antwoordstrookje naar het adres van de boekbesprekingsrubriek als u interesse heeft om in de toekomst een boek te bespreken. Als er dan boeken binnenkomen heb ik een grotere keuze van mogelijke recensenten en u wellicht een gloednieuw boek op uw onderzoeksgebied.

Ja, ik ben graag bereid een boek op mijn interessegebied te bespreken voor de leden van de VVS

naam

adres

interessegebied (graag enigszins nauw-
keurig omschrijven)
