

BOEKBESPREKINGEN

Redactie: Arend D. Oosterhoorn
 Correspondentieadres: Zie bladzijde 1.

M.L. Tiku, W.Y. Tan and N. Balakrishnan

Robust inference

Marcel Dekker Inc., New York and Basel, 1986, 321 pag.,
 ISBN 0-8274-7532-5, f 191,52

In dit boek worden statistische procedures voorgesteld voor schatting en toetsing die door schrijvers "Robuust" worden genoemd. Wat men hieronder verstaat wordt in het voorwoord uitgelegd:

"a robust (to deviations from normality) procedure is defined to be one that is nearly as efficient as the classical procedure (based on the sample mean $\bar{x}$ and variance s^2) for a normal distribution but is considerably more efficient overall for nonnormal distributions, particularly distributions that constitute plausible alternatives to normality."

Allereerst moet worden gezegd dat de "plausible alternatives to normality" in het boek zijn beperkt tot enkele geparametrizeerde modellen die de normale verdeling als element bevatten. De "nonnormal distributions" zijn zodoende zeer specifieke alternatieven die, afhankelijk van de model parameter, schever zijn dan de normale verdeling of dikkere staarten bezitten dan die van de normale verdeling.

Verder blijven alle welbekende robuustheids concepten, zoals het "breakdown point" (zijnde het % uitschieters dat de schatting oneindig groot kan maken) of de "influence function" IF (die de mate van invloed weergeeft van een uitschieter op de schatter), ongenoemd.

Centraal in het boek staan de "Modified Maximum Likelihood" (MML) procedures. Deze aangepaste maximum likelihood methode, die gebaseerd is op "censored samples" $x_{(r_1+1)} \leq \dots \leq x_{(n-r_2)}$,

waar de r_1 kleinste en r_2 grootste waarnemingen uit de oorspronkelijke steekproef zijn weggelaten, wordt besproken in hoofdstuk 2.

De daarin ontwikkelde MML methode wordt gebruikt voor

1. het schatten van (een-dimensionale) parameters van verschillende verdelingen, zoals lokatie μ en variantie σ^2 van de normale verdeling.
2. Het toetsen van hypothesen zoals $H_0: \mu = 0$;
 $H_0: \mu_1 = \mu_2$ of
 $H_0: \sigma_1 = \sigma_2$.
3. het schatten van een (meer-dimensionale) regressie parameter θ

Hoofdstuk 1 bespreekt de klassieke t toets en de F toets. Beide toetsen heten asymptotisch robuust, waarmee wordt bedoeld dat in beide gevallen de verdeling van de toetsings-grootheden blijft convergeren naar een normale verdeling, mits de onderliggende verdeling van de steekproef niet te veel afwijkt van de normale verdeling. Een nogal beperkte opvatting van robuustheid aangezien de lengte van de betrouwbaarheidsintervallen wel gevoelig is voor afwijkingen (uitschieters) van de normale verdeling, hetgeen zijn weerslag heeft op het al dan niet verwerpen van H_0 .

In de hoofdstukken 2 en 3 worden enkele MML schatters afgeleid waarvan wordt aangetoond dat ze robuust zijn in de zin zoals boven is geciteerd. Ter vergelijking dienen de M-schatters van Huber, die een iets kleinere efficiëntie blijken te hebben. Echter bezitten deze wel een 50 % "breakdown point" en is hun "influence function" begrensd, maar dit wordt in het boek genegeerd.

Enkele MML-toetsen worden besproken in hoofdstuk 4. Ze zijn robuust in die zin dat het onderscheidingsvermogen niet al te veel verandert wanneer de onderliggende verdeling niet al te gek afwijkt van de normale verdeling (zoals eerder gezegd alleen binnen een geparametriseerd model). Echter zijn ook deze toetsen niet robuust in die zin dat de lengte van de betrouwbaarheidsintervallen wel gevoelig is voor afwijkingen van de normale verdeling.

Hoofdstuk 5 is gewijd aan het multiple lineaire model. Opnieuw dienen de Huber M-schatters als vergelijkingsmateriaal en kunnen dezelfde kanttekeningen worden geplaatst als bij de MML schatters uit de hoofdstukken 2 en 3. Bovendien worden uitschieters in de onafhankelijke veranderlijken, die in de praktijk voor veel grotere problemen zorgen, niet beschouwd.

Tenslotte levert het boek nog een hoofdstuk over "Goodness-of-fit" toetsen en een hoofdstuk over toetsen op uitschieters.

Alle hoofdstukken gaan gepaard met zeer veel numerieke voorbeelden en het boek staat dan ook bol van de tabellen. Verder is aan elk hoofdstuk een aantal opgaven toegevoegd.

Het boek maakt een rommelige indruk en is nogal onoverzichtelijk (iemand met een IBM-typemachine met verwisselbare koppen had hetzelfde resultaat kunnen bereiken). Ondanks de appendices bij de hoofdstukken 1 en 2 zijn de bewijzen over het algemeen weggelaten en wordt men verwezen naar de artikelen van de betreffende auteurs. Het voornaamste nadeel van dit boek is dat het geen enkele aandacht besteedt aan de recente ontwikkelingen in de robuustheidstheorie, vanaf de werken van Huber in de zestiger jaren tot heden.

R. Lopuhaä, TU Delft

Warren Gilchrist

Statistical Modelling

John Wiley & Sons, New York, 1984, 339 pag.,

ISBN-nummer 0 471 90391 4, £ 9.50

Het doel van dit inleidende boek is, zoals de auteur het zelf omschrijft, het op een gestructureerde manier samenbrengen van die statistische en verwante ideeën, die nodig zijn voor de constructie van kwantitatieve modellen.

Gilchrist is daarin goed geslaagd.

De schrijver heeft zich bij de vormgeving van het boek laten leiden door de belangrijkste processen, die bij de constructie van een model aan de orde komen en door de rol die statistische methoden daarbij spelen.

Zo is de volgende opbouw in zeven delen tot stand gekomen:

Part I	The Modelling process	24 blz.
Part II	Common Statistical models	59 blz.
Part III	Identification	69 blz.
Part IV	Estimation	51 blz.
Part V	Validation	30 blz.
Part VI	Application	21 blz.
Part VII	Iteration	

In delen I en II wordt de statistische basis van het boek gelegd, bestaan de uit o.a. een behandeling van het begrip "likelihood", de bekende reeks kansverdelingen (Bernouilli tot Gauss), een inleiding op de statistische besluitvorming, e.d.

De beschrijving van al deze onderwerpen is betrekkelijk beknopt en zal voor statistisch niet-geschoolde lezers onvoldoende zijn. Wat uitvoeriger wordt ingegaan op regressie en zgn. "progressie" modellen (o.a. "autoregressive model" en "moving average model").

Deel III behandelt algemene principes m.b.t. de vraag op welke wijze de beschikbare voorkennis bij een probleem kan worden gebruikt bij de constructie van een model, hoe conceptuele informatie verwerkt kan worden (conceptual identification). Het volgende hoofdstuk (empirical identification) bespreekt op beknopte wijze een beperkt aantal mogelijkheden (w.o. een aantal grafische presentaties van een- en twee dimensionale gegevens) om getallenmateriaal te leren kennen. De laatste twee hoofdstukken van dat deel proberen een combinatie van empirische en conceptuele informatie te verheffen tot een zgn. "eclectische" identificatie. De schrijver behandelt in dit verband o.a. het gebruik van probability-plots, analyse van residuen, stapsgewijze regressie-analyse en het probleem van de identificeerbaarheid van modellen.

Deel IV behandelt de wijze waarop een verondersteld model kan worden "aangepast" aan de data.

In dit verband komen o.a. aan de orde de kleinste kwadraten methode en de maximum likelihood methode, en vervolgens een aantal eigenschappen van schatters.

Ook de rol die Bayesiaanse methoden kunnen spelen bij het gebruik van conceptuele en empirische informatie in dit verband wordt behandeld.

Deel V (Validation) behandelt verschillende mogelijkheden om de geldigheid van een model m.b.t. het doel waarvoor het model ontwikkeld is na te gaan. De schrijver maakt hierbij o.m. onderscheid tussen conceptuele validatie (vergelijking van de eigenschappen van het model met de ideeën van de onderzoeker hierover), empirische validatie (vergelijking met de - eventueel voor dat doel verzamelde - data) en "eclectische" validatie.

Deel VI (Application) behandelt de vraag op welke wijze rekening gehouden kan worden met de aard van de toepassing waarvoor het betreffende model wordt ontwikkeld.

Deel VII ten slotte gaat in op het feit dat in de praktijk van modelconstructie de verschillende behandelende fases elkaar niet netjes opvolgen, maar meer op iteratieve wijze worden gehanteerd. De schrijver behandelt daartoe diverse concrete voorbeelden.

Het boek bevat verder nog 18 probleemgeëïntende opgaven met gedeeltelijke uitwerkingen, en per deel een lijst van boeken en tijdschriftartikelen.

De vele voorbeelden uit diverse takken van wetenschap en de gestructureerde opzet maken het boek prettig leesbaar en zeer relevant voor iedereen die zich wil oriënteren op de mogelijkheden van kwantitatieve modellen en statistische aspecten daarvan.

L. Stijbosch, KUB

D.H. Kaye & Mikel Aickin, ed.,

Statistical Methods in Discrimination Litigation

Marcel Dekker Inc., New York and Basel, 1986, 218 bladzijden, ISBN 0-8247-7514-7, \$ 59,50.

Het onderwerp van dit boek is de rol die de statistische besluitvorming speelt bij discriminatie-rechtzaken in de Verenigde Staten. Hoewel de bijdragen van de onderscheiden auteurs ongelijksoortig zijn voor wat betreft detaillering en diepgang, geeft dit boek een zeer goed overzicht van de specifieke problemen die het toepassen van de statistiek in de rechtszaal met zich meebrengt.

Neem het geval van een in Georgia ter dood veroordeelde zwarte Amerikaan die via een beroep op de federale rechter het doodvonnis ongedaan wilde doen maken en aanvoerde dat bij het ter dood veroordelen in Georgia op huidskleur gediscrimineerd wordt. De beschikbare gegevens werden uitputtend geanalyseerd m.b.v. regressie methoden waarbij in eerste instantie 230 factoren werden onderscheiden, die mogelijk van invloed waren op het toekennen van de doodstraf aan meer dan 1000 wegens moord veroordeelde personen in de periode 1973 tot 1978.

Het uiteindelijke verklarende model bevatte 39 onafhankelijke veranderlijken.

Hoewel sprake was van een uitputtende analyse van de beschikbare gegevens, bleek het hof door de overdaad aan statistische analyse

in zijn oordeel geschaad te zijn. Het oordeelde dat een model met minder dan de oorspronkelijke 230 variabelen wellicht niet de hele waarheid weerspiegelt en in een adem dat de aanwezigheid van multicollineariteit de bewijswaarde van het waarnemings materiaal ernstig aantast. Het concludeerde dat er een bewijs was van discriminatie bij het toekennen van de doodstraf.

Uitgebreid komt aan de orde het probleem van de relevante populatie. Als een werkgever aangeklaagd wordt wegens discriminatie bij het werven van personeel, dan kan de reputatie van zijn bedrijf zo zijn, dat vanuit gediscrimineerde groepen überhaupt niet naar een betrekking gesolliciteerd wordt. De klager zal dan naar voren brengen dat de fractie werknemers uit deze groep(en) moet worden vergeleken met de samenstelling van de werkende bevolking als geheel. De tactiek van de beklaagde kan zijn de relevante populatie te verkleinen d.m.v. een groot aantal zgn. relevante karakteristieken. Immers: hoe kleiner de populatie, hoe kleiner het onderscheidingsvermogen van statistische toetsen. Dit is een door statistici wel, maar door rechters en jury's soms niet onderkend fenomeen.

Een zeer interessant hoofdstuk gaat over de problemen die optreden bij regressie als de verklarende variabele met fout wordt waargenomen. Indien dit niet onderkend wordt, kan er onderschatting of overschatting van discriminatie optreden afhankelijk van welke variabele als verklarende wordt genomen.

Het boek bevat ook een pleidooi tegen het gebruik van Bayesiaanse statistiek in de rechtszaal. Hoewel de Bayesiaanse terminologie beter geschikt lijkt voor de rechtszaal dan de frequentistische, moet toch voorkomen worden dat de eeuwige discussie over de status van priorverdelingen de aan de orde zijnde juridische kwesties vertroebelt.

"Selectie-effecten" worden besproken in het laatste hoofdstuk. De statisticus kan het waarnemingsmateriaal reduceren en hij kan een test kiezen uit een aantal mogelijke. Ook is onbekend waarom het ene discriminatie-geval voor de rechter komt en het andere niet. Dit alles heeft onbekende gevolgen voor de waarde van de significantie van het waarnemingsmateriaal.

Dit boek is een absolute aanrader voor juristen en statistici, die als aanklager, verdediger, rechter of getuige-deskundige met statistisch bewijsmateriaal te maken hebben. Enige basiskennis van statistiek is vereist.

H.N. Linssen, TUE