

SA

**stichting  
mathematisch  
centrum**



AFDELING MATHEMATISCHE STATISTIEK

SD 108/74

DECEMBER

HENK ELFFERS

EEN KOVARIANTIE-ANALYSE BIJ EEN ONDERZOEK NAAR HET  
VEROUDERINGS-PROCES

**2e boerhaavestraat 49 amsterdam**

BIBLIOTHEEK MATHEMATISCH CENTRUM  
AMSTERDAM

*Printed at the Mathematical Centre, 49, 2e Boerhaavestraat, Amsterdam.*

*The Mathematical Centre, founded the 11-th of February 1946, is a non-profit institution aiming at the promotion of pure mathematics and its applications. It is sponsored by the Netherlands Government through the Netherlands Organization for the Advancement of Pure Research (Z.W.O), by the Municipality of Amsterdam, by the University of Amsterdam, by the Free University at Amsterdam, and by industries.*

Een kovariantie-analyse bij een onderzoek naar het verouderings-proces

door

Henk Elffers

#### SAMENVATTING

Bij een onderzoek bij honderd mannen naar de invloed van sommige variabelen op het verloop van andere variabelen met de leeftijd werd een kovariantie-analyse toegepast.

KEY WORDS & PHRASES: *Kovariantie-analyse*.



## 1. PROBLEEMSTELLING EN TERMINOLOGIE

De in dit rapport beschreven methode is toegepast bij een onderzoek naar het cardiovasculaire verouderingsproces met non-invasieve meetmethoden van U. Zuiderveld, waarvan de verslaggeving is geschied in ZUIDERVELD(1974). Nadere gegevens omtrent de metingen, de betrokken variabelen, alsmede de uitslagen van de verrichte analyses staan daar vermeld. In dit rapport wordt daar bij tijd en wijle naar verwezen.

Aan een aantal mensen zijn gemeten een aantal variabelen zoals activiteit, rookgewoonten, overgewicht, verder te noemen variabelen van de 1e categorie (of splitsende variabelen), en een aantal zoals bloeddruk en zuurstofgebruik tijdens inspanning, te noemen variabelen van de 2e categorie (of te splitsen variabelen).

Het centrale probleem werd gesteld als:

*Welke variabelen van de 1e categorie hebben invloed op het verloop met de leeftijd van variabelen van de 2e categorie, en hoe is die invloed?*

Bij invloed van variabelen van de 1e categorie is slechts gedacht in termen van "veel" en "weinig". Een nadere bepaling van deze begrippen wordt in §3.4 uitgewerkt. Zie ook §4.5.

Een verdere inperking van het centrale probleem vloeit voort uit de beslissing ons slechts te bekommeren om het lineaire verloop met de leeftijd, enerzijds om redenen van eenvoud, anderzijds wegens ontstentenis van positieve aanwijzingen in andere richting.

De enige uitzondering vormt het in §5.1 beschreven procedé, dat trouwens een gedeeltelijke rechtvaardiging voor deze beperking opleverde.

Het onderzoek had een exploratief karakter, hetgeen in opzet en werkwijze tot uiting kwam. Het analyseplan werd eerst na de verzameling van de gegevens opgesteld, zodat de analyse hier en daar mogelijk door de opgetreden waarnemingen is beïnvloed. Bovendien zijn verscheidene hypothesen bestudeerd aan de hand van hetzelfde materiaal, zonder expliciet rekening te houden met afhankelijkheid van de verrichte toetsen (§4.2).

De uitslag van het onderzoek moet dan ook veeleer gezien worden als richtinggevend voor nader onderzoek, dan als formeel statistisch bewijs van het gestelde.

## 2. HERKOMST VAN HET MATERIAAL

Voor een onderzoek naar verloop met de leeftijd komt in eerste aanleg een longitudinale opzet, waarbij een aantal proefpersonen van de wieg tot het graf worden gevolgd, in aanmerking. Alleen al om redenen van de duur van een dergelijk onderzoek, is dat vrijwel nooit doenlijk, ook in dit geval niet. Daarom had de onderzoeker besloten tot het trekken van een naar leeftijd gelaagde steekproef, d.w.z. dat elke leeftijdskategorie gelijkelijk vertegenwoordigd is. Dit geeft, onder enig voorbehoud (§4.3) een redelijke basis voor onderzoek naar verloop met de leeftijd.

Een steekproef van omvang 100 werd getrokken uit het "uitgedunde" patientenbestand van de onderzoeker, groot ongeveer 2000 à 2500 personen. "Uitgedund" wil zeggen, dat tevoren waren weggelaten al degenen die geregistreerd stonden als lijdende aan een hart- of vaatziekte, alsmede alle mensen jonger dan 20 of ouder dan 70 jaar. De steekproef werd zo getrokken dat in elke leeftijdsklasse van vijf jaar (20-25, 25-30, ..., 65-70) precies 10 personen door het lot werden aangewezen. (De zo optredende leeftijden zijn dus eigenlijk realisaties van stochastische grootheden, maar zij zijn in het gehele onderzoek behandeld als vaste getallen. Dit is geen bezwaar, gezien de klassebreedte van vijf jaar, die klein is t.o.v. de verwachte verschillen met de leeftijd). De omvang van de leeftijdsgroepen in het uitgedunde patientenbestand is niet bekend. Terzijde zij opgemerkt dat de steekproefopzet niet tevoren was afgestemd op het in de volgende paragraaf beschreven model, daar eerst na de materiaalverzameling contact tussen onderzoeker en statistikus tot stand kwam.

Statistisch verantwoorde generalisatie op grond van het in §3 te behandelen model heeft, onder het voorbehoud van §4, derhalve betrekking op het uitgedunde bestand, waarmee niets gezegd wil zijn ten nadele van conceptueel verantwoorde generalisatie naar andere populaties, die buiten de kompetentie van de statistikus valt.

Aan elke proefpersoon werden, naast zijn leeftijd, gemeten 7 variabelen van de 1e categorie, en 19 variabelen van de 2e categorie.

Bij een aantal proefpersonen boven de 60 jaar, was het niet mogelijk bepaalde variabelen van de 2e categorie te meten, omdat de metingen een te zware belasting voor hen zouden betekenen. Voor deze variabelen is de leef-

tijdsbalans dus zoek. (Naast deze 26 variabelen zijn nog andere variabelen gemeten, die evenwel niet aan de hier beschreven analyse werden onderworpen).

### 3. EEN MODEL

We concentreren ons op de bestudering van het verloop van één bepaalde variabele van de 2e categorie, zeg  $y$ , met de leeftijd, onder invloed van één bepaalde variabele van de 1e categorie, zeg  $x$ .

We stellen ons voor dat de populatie,  $B$ , d.w.z. het uitgedunde patiëntenbestand verdeeld is in twee delen,  $B_1$  en  $B_2$ , op een wijze, die met de variabele  $x$  samenhangt. Gedacht is aan:  $B_1$  bevat de mensen met verhoudingsgewijs hoge  $x$ -skore,  $B_2$  hen die verhoudingsgewijs lage  $x$ -skore hebben. In §3.4 wordt hierop nader ingegaan.

Voor het gemak nummeren we de personen in  $B$  als  $b_1, b_2, \dots, b_N$ .

In de volgende paragrafen zullen verscheidene aspecten van het gebruikte model, een zogeheten kovariantie-analyse-model, worden belicht.

#### 3.1. Model voor het verloop van $y$ met de leeftijd

Voor elke persoon  $b$  in  $B$  nemen we aan dat meting van de  $y$ -variabelen geen vaste uitkomst  $y(b)$  oplevert. We zullen zo'n uitkomst daarom als een stochastische variabele, genoteerd  $\underline{y}(b)$ , opvatten. (Onderstreping duidt hier en elders op stochastische variabelen). Laat  $\underline{y}(b)$  gemiddeld over alle mogelijke metingen de waarde  $z(b)$  hebben, en noem het verschil tussen  $\underline{y}(b)$  en  $z(b)$   $\underline{\varepsilon}(b)$ . Men kan  $\underline{\varepsilon}$  zien als veroorzaakt door storende invloeden. Dan geldt:

$$(3.1.) \quad \text{voor alle } b \text{ in } B: \underline{y}(b) \equiv z(b) + \underline{\varepsilon}(b), \text{ met } E \underline{\varepsilon}(b) = 0 \quad *)$$

De centrale veronderstelling is nu dat  $z(b)$  een relatie met de leeftijd van  $b$ , genoteerd  $l(b)$ , heeft van de volgende vorm (voor beide populatiedelen verschillend):

\*) het teken  $E$  (mathematische verwachting), worde geïnterpreteerd als "het gemiddelde over alle mogelijke metingen van"

$$(3.2) \quad \left\{ \begin{array}{ll} z(b) = \alpha_1 + \beta_1 l(b) + r_1(b) & \text{als } b \text{ in } B_1 \text{ (fig.3.1)} \\ z(b) = \alpha_2 + \beta_2 l(b) + r_2(b) & \text{als } b \text{ in } B_2 \\ \text{met: } \sum_{b \in B_1} r_1(b) = 0 \text{ en } \sum_{b \in B_2} r_2(b) = 0 & \text{voor elke leeftijd } a \\ & \text{tussen 20 en 70 jaar} \\ l(b)=a & l(b)=a \end{array} \right.$$

Dat wil zeggen: er bestaat een lineair verband tussen  $z$  en  $l$ , behoudens afwijkingen  $r_1$ , resp.  $r_2$  waarvoor geldt:

het gemiddelde van de individuele afwijkingen  $r_1$  (resp.  $r_2$ ) van alle personen van gelijke leeftijd in populatiedeel  $B_1$  (resp.  $B_2$ ) is nul (fig.3.2)

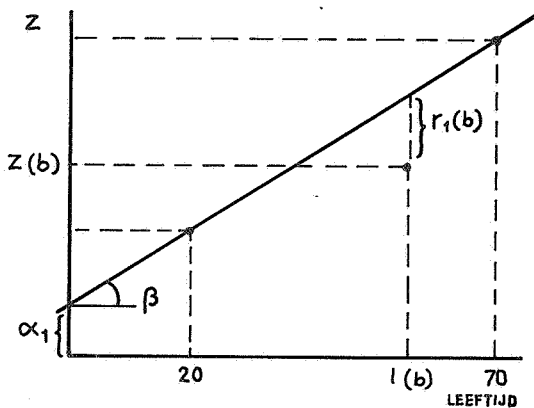


fig. 3.1

de regressielijn  $z = \alpha_1 + \beta_1 l$ ,  
met  $t_g \beta = \beta_1$

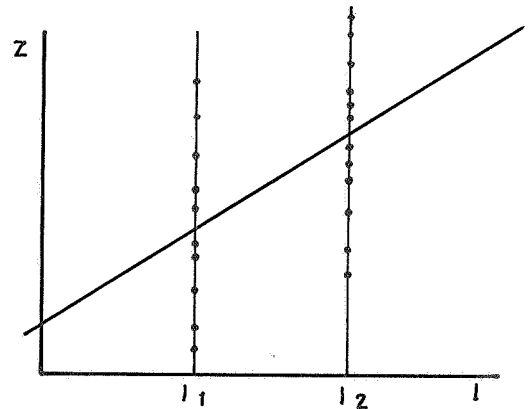


fig. 3.2

de afwijkingen  $r(b)$  voor vaste leeftijd

Daar het onmogelijk is alle personen te meten, loten we er een aantal zoals beschreven in §2.

Laten de gelote personen  $p_1, p_2, \dots, p_{100}$  zijn (d.w.z.  $p_1 = b_{128}$  betekent: de eerste loting wees de 128<sup>e</sup> persoon in het bestand aan, enz.)

Voor de  $z$ -waarden van de gelote personen geldt dan:

$$(3.3) \quad z(p_j) = \begin{cases} \alpha_1 + \beta_1 l(p_j) + r_1(p_j) & \text{als } p_j \text{ in } B_1 \\ \alpha_2 + \beta_2 l(p_j) + r_2(p_j) & \text{als } p_j \text{ in } B_2 \end{cases} \quad j = 1, 2, \dots, 100;$$

Zoals reeds in §2 vermeld is, werd de lotingsprocedure zo ingericht, dat wij  $l(p_j)$  als een vaste (d.i. niet stochastische) uitkomst, kunnen beschouwen, verder te noteren als  $l_j$ .

In werkelijkheid kunnen we  $z(p_j)$  niet observeren, maar alleen  $y(p_j)$ . Daarvoor geldt evenwel

$$(3.4^a) \quad y(p_j) = \begin{cases} \alpha_1 + \beta_1 l_j + r_1(p_j) + \varepsilon(p_j) & \text{als } p_j \text{ in } B_1 \\ \alpha_2 + \beta_2 l_j + r_2(p_j) + \varepsilon(p_j) & \text{als } p_j \text{ in } B_2 \end{cases} \quad j = 1, 2, \dots, 100.$$

In verband met het in §3.4 te vermelden, merken wij op dat de geobserveerde skores op tweeërlei wijze een stochastisch karakter tonen: enerzijds wegens de door de experimentator uitgevoerde lotingsprocedure, anderzijds door de, niet door de experimentator gemanipuleerde storende invloeden die  $y$  van  $z$  doen afwijken.

Door de assumpties in (3.2) over de som van de  $r_1$ 's en  $r_2$ 's en de definitie van  $\varepsilon(b)$  geldt nu:

$$(3.4^b) \quad \begin{cases} E(r_1(p_j) + \varepsilon(p_j)) = 0 & \text{voor alle } j \text{ met } p_j \text{ in } B_1 \\ E(r_2(p_j) + \varepsilon(p_j)) = 0 & \text{voor alle } j \text{ met } p_j \text{ in } B_2 \end{cases}$$

Voor de verdere analyse is de precieze samenstelling van de restterm niet van belang, en we zullen schrijven

$$(3.5) \quad y(p_j) = \begin{cases} \alpha_1 + \beta_1 l_j + \eta_{1j} & \text{als } p_j \text{ in } B_1 \\ \alpha_2 + \beta_2 l_j + \eta_{2j} & \text{als } p_j \text{ in } B_2 \end{cases} \quad j = 1, 2, \dots, 100$$

met  $E\eta_{1j} = 0$  en  $E\eta_{2j} = 0$ , alle  $\eta$ 's stochastisch onafhankelijk.

De onafhankelijkheid in (3.5.) geldt strikt genomen slechts als de lotingsprocedure met teruglegging wordt uitgevoerd. De verhouding steekproefgroot-

te/populatieomvang laat echter wel toe te negeren, dat we in feite zonder teruglegging hebben gewerkt.

(3.5) heet de regressievergelijking van  $y$  op  $1$ ;  $\alpha_1, \alpha_2, \beta_1, \beta_2$  heten de regressiecoëfficiënten of kortweg de parameters.

Onze opgave is nu uit de waarnemingen uitspraken over de regressiecoëfficiënten af te leiden.

### 3.2. Schatten van regressiecoëfficiënten

Wanneer we aannemen dat:

- (i) het model (3.5) korrekt is (vgl. §5)
- (ii) bij elke  $j$  bekend is of  $p_j$  in  $B_1$  dan wel in  $B_2$  terechtkomt (vgl. §3.4), dan geldt het volgende

- (1) De regressiecoëfficiënten kunnen zuiver geschat worden, d.w.z. zodanig geschat worden dat, wanneer dezelfde proef talloze malen herhaald zou worden, de resulterende schattingen gemiddeld de werkelijke waarden aannemen. Hiervoor gebruiken we kleinste kwadratenschatters (zie bijv. DE JONGE (1960)).
- (2) De 20-70 waarden van de regressielijnen kunnen zuiver geschat worden. Deze waarden zijn de ordinaten van de punten op de lijnen waarvoor  $l=20$  respectievelijk  $l=70$  (fig.3.1). Zij kunnen evenzeer dienen om de lijnen te karakteriseren als de  $\alpha$ 's en  $\beta$ 's. Voor presentatiedoeleinden zijn schattingen van deze waarden vermeld (ZUIDERVELD(1974)).

### 3.3. Toetsen van regressiecoëfficiënten

Zoals in §1 uiteengezet richt het onderzoek zich op de vraag: is het verloop van  $y$  met de leeftijd in  $B_1$  (hoge  $x$ ) en  $B_2$  (lage  $x$ ) verschillend?

In termen van het model (3.5) kunnen we deze vraag formuleren als:

geldt  $\alpha_1 = \alpha_2$  en  $\beta_1 = \beta_2$  of niet?

Merk op dat, zelfs indien inderdaad  $\alpha_1 = \alpha_2$  en  $\beta_1 = \beta_2$ , de schattingen van deze parameters niettemin meestal zullen verschillen, vanwege de aanwezigheid van de stochastische termen  $\underline{n}_1$  en  $\underline{n}_2$  in het model. Men kan zich dan afvragen: kunnen de gevonden verschillen toevalligerwijs optreden, terwijl  $\alpha_1 = \alpha_2$  en  $\beta_1 = \beta_2$ , of is dat wel heel onwaarschijnlijk?

Wanneer we bereid zijn een aanvullende assumptie over de stochastische termen  $\eta_{1j}$  en  $\eta_{2j}$  in (3.5) te maken, en (i) en (ii) uit §3.2 weer gelden, is het mogelijk een toetsingsprocedure uit te voeren. We stellen daarbij als (nul)hypothese:

$$H_0: \alpha_1 = \alpha_2 \text{ en } \beta_1 = \beta_2 \text{ (geen verschil)}$$

en berekenen uit de waarnemingen een statistische grootte, de toetsingsgrootte. Deze heeft de eigenschap, dat hij, wanneer  $H_0$  juist is, zeer grote waarden slechts met zeer kleine kans aanneemt. Komt er toch zo'n grote waarde voor, dan verwerpen we  $H_0$ , en besluiten tot het alternatief

$$K: \alpha_1 \neq \alpha_2 \text{ of } \beta_1 \neq \beta_2$$

De grens van wat "zeer groot" genoemd wordt is daarbij zo gekozen, dat de kans dat men de nulhypothese verworpt alhoewel hij in feite juist is, kleiner is dan een tevoren gekozen kleine waarde  $\alpha$ , de zogenaamde onbetrouwbaarheidsdrempel. In dit onderzoek is voor  $\alpha$  de waarde 0.05 gekozen. [Men zegt ook wel, wanneer de nulhypothese wordt verworpen, dat de gevonden verschillen significant zijn op nivo  $\alpha$ ]. Naast deze mogelijkheid tot het nemen van een foute beslissing (fout van de eerste soort), kan men ook een fout van de tweede soort maken, namelijk ten onrechte de nulhypothese niet verwerpen. Over de kans op een fout van de tweede soort is in het algemeen minder te zeggen, maar hij kan aanzienlijk groter dan  $\alpha$  uitvallen. De kans op een fout van de tweede soort wordt voornamelijk bepaald door de in werkelijkheid bestaande verschillen, en door het aantal waarnemingen. Grotere feitelijke verschillen en grotere aantallen waarnemingen doen de kans op een fout van de tweede soort afnemen.

Wanneer de nulhypothese niet wordt verworpen kan men konkluderen: er is geen doorslaggevende reden de nulhypothese te verwerpen. Dit houdt niet in dat de nulhypothese bevestigd wordt. Ook andere nulhypothesen zouden wellicht niet verworpen zijn, bijv.  $H_1: \alpha_1 = \alpha_2 + 0,01$ ,  $\beta_1 = \beta_2 * 1.1$ , om maar eens wat te noemen.

De voor deze werkwijze benodigde assumptie is

(3.6)  $\eta_{1j}$  en  $\eta_{2j}$  zijn stochastisch onafhankelijke trekkingen uit een  $N(0, \sigma^2)$  verdeling (met onbekende  $\sigma^2$ )

Hierin betekent  $N(0, \sigma^2)$  verdeling: normale verdeling met verwachting 0 en variantie  $\sigma^2$ .

Onder deze assumptie kan  $H_0$  met een variantie-ratiotoets (F-toets) worden getoetst (zie bijv. DE JONGE (1960)).

T.o.v. de aannames in (3.5) voegt (3.6) toe:

(1) normaliteit van de resttermen

Door de eindige omvang van de populatie kan de normaliteit nooit exakt vervuld zijn. Bekend is evenwel dat de F-toets ook zeer goed voldoet als de betrokken verdeling enigszins van de normale verschilt.

(2) Konstante variantie van de  $\eta$ 's voor elke leeftijdsgroep en populatiedeel.

Is dit niet het geval dan kan de kans op een fout van de eerste soort wel eens veel groter dan  $\alpha$  uitvallen. Aan deze assumptie is een onderzoekje gewijd in §5.2.

### 3.4. De verdeling van de populatie in twee delen.

Voor de verdeling van de populatie in twee delen op grond van de x-waarden van de personen is gekozen voor het criterium:

b behoort tot  $B_1$  als  $x(b)$  groot is voor de leeftijd  $l(b)$  van b, en tot  $B_2$  als  $x(b)$  klein is voor de leeftijd  $l(b)$

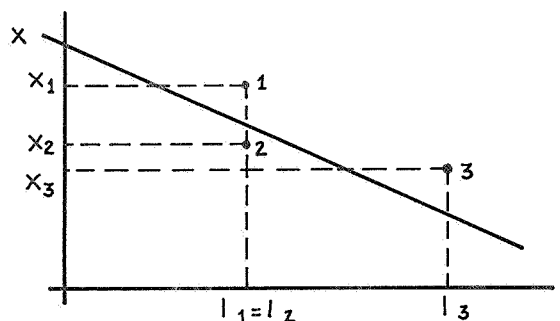


fig. 3.3 de verdeling in  $B_1$  en  $B_2$

Daarbij is "groot/klein voor zijn leeftijd" aan de hand van de regressie van  $x$  op  $l$  gedefinieerd. In bovenstaande figuur hebben de punten 1 en 3 aan grote  $x$  voor hun leeftijd, 2 een kleine, (niettemin is  $x_3$  in absolute zin kleiner dan  $x_2$ ).

Daarbij veronderstellen we dan

$$(3.7) \quad \begin{cases} \text{voor alle } b: x(b) = \gamma + \delta l(b) + q(b) \\ \text{met } \sum_{l(b)=a} q(b) = 0 \text{ voor elke } a \text{ (vgl. (3.2)).} \end{cases}$$

Met deze notatie is het criterium:

$$(3.8) \quad \begin{cases} b \text{ behoort tot } B_1 \text{ als } q(b) > 0 \\ b \text{ behoort tot } B_2 \text{ als } q(b) \leq 0 \end{cases}$$

Door de loting van  $p_1, \dots, p_{100}$  ontstaat het model

$$(3.9) \quad \begin{cases} x(p_j) = \gamma + \delta l_j + q(p_j), \quad j = 1, \dots, 100 \\ \text{met } E q(p_j) = 0, \quad q(p_j) \text{ en } q(p_k) \text{ onafhankelijk als } j \neq k \end{cases}$$

Merk op dat we hier, in tegenstelling tot in (3.4) veronderstellen dat de enige afwijking van de regressielijn hier door het steekproeftrekken (d.i. het loten van personen) optreedt. Er wordt derhalve aangenomen dat herhaalde  $x$ -metingen aan een bepaalde persoon konstante uitkomsten opleveren.

Bij de metingen is er uitdrukkelijk zorg gedragen niet-persoonskarakteristieke invloeden uit te sluiten. Hierdoor is er bij de splitsing in  $B_1$  en  $B_2$  één storingsbron geëlimineerd, hetgeen de kans op het maken van een fout van de tweede soort vermindert.

Ook nu kunnen we  $\gamma$  en  $\delta$  in (3.9) slechts schatten, met onvermijdelijke fouten. Dit brengt met zich mee dat we sommige personen wellicht misklassificeren als tot  $B_1$  of  $B_2$  behorend. In figuur 3.4 wordt punt 1 ten onrechte tot  $B_2$  gerekend, punt 2 en 3 worden korrekt geklassificeerd. Voor het ana-

lyseren van model (3.5) zullen we evenwel afzien van het stochastische karakter van de verkregen klassificatie.

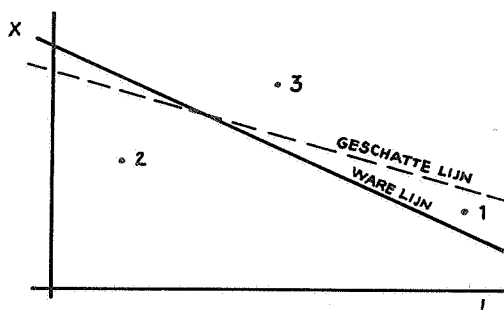


fig. 3.4 fouten in de indeling in  $B_1$  en  $B_2$

Voor een eventueel gewenste toetsing omtrent  $\gamma$  en  $\delta$  in (3.9) is weer een normaliteits-veronderstelling nodig.

Indeling volgens criterium (3.8) heeft als voordeel boven het criterium:

(3.10)  $b$  behoort tot  $B_1$  als  $x(b)$  groter is dan de mediaan van de  $x(b_i)$  ( $i=1, \dots, N$ ), anders tot  $B_2$

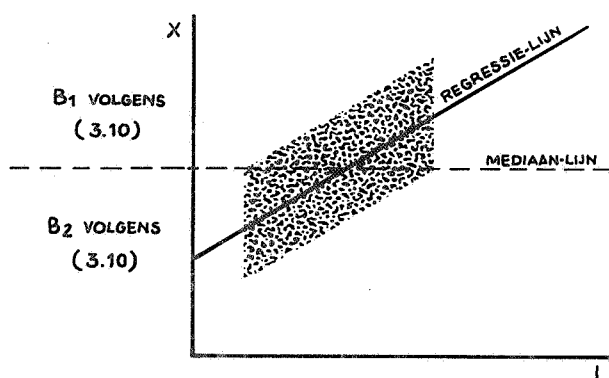


fig. 3.5 de criteria (3.8) en (3.10)  
(elk punt stelt een persoon in  $B$  voor)

dat in beide klassen oude en jonge mensen gelijkelijk vertegenwoordigd zijn, hetgeen bij indeling naar (3.10) niet het geval behoeft te zijn.

Als de situatie ligt als in fig. 3.5 heeft indeling naar (3.10) tot gevolg dat  $B_1$  overwegend oudere mensen bevat,  $B_2$  daarentegen merendeels jongere. Dit doet het onderscheidingsvermogen van de toets afnemen t.o.v. een gelijk-

matige verdeling. Bovendien is de procedure dan erg gevoelig voor afwijking van de modelveronderstellingen.

Merk nog op dat in geval  $\delta$  in (3.9) nul is, beide criteria samenvallen.

Ook valt te overwegen de verdeling van de populatie in twee delen te laten geschieden op grond van het feit of de  $x(b)$  al dan niet groter is dan de mediaan van de  $x(b_i)$  van de personen in eenzelfde leeftijdsgroep. Dit criterium eist dus een dergelijke indeling in leeftijdsgroepen, zie fig 3.6. Deze klassificatie is niet afhankelijk van een veronderstelling van een lineair verband tussen  $x$  en  $l$ .

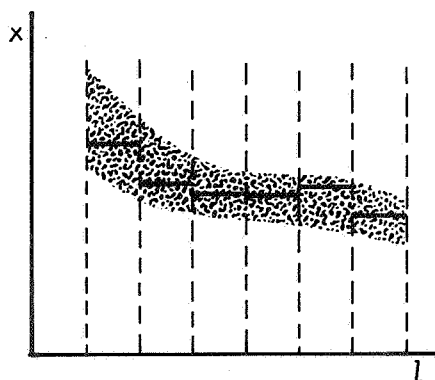


fig. 3.6 indeling naar mediaan van de leeftijdsgroep

Daar echter in iedere leeftijdsgroep slechts 10 waarnemingen verricht zijn, heeft de mediaan een grote spreiding, hetgeen de splitsing een ietwat onnauwkeuriger karakter gegeven zou hebben dan de splitsing met behulp van een regressielijn.

#### 4. BEZWAREN TEGEN DE GEBRUIKTE METHODE, ALTERNATIEVEN

##### 4.1. Afgrenzing van de populatie

Het in §2 genoemde criterium "geregistreerd als lijdende aan een hart- of vaatziekte", dat gebruikt werd om de populatie te bepalen, heeft een subjectief karakter, zodat nogal vaag is op welke mensen de resultaten betrekking hebben.

#### 4.2. Afhankelijkheid van toetsen

Zoals bij zoveel onderzoeken is getracht zoveel mogelijk informatie uit de steekproef te halen, meer zelfs dan de proefopzet strikt genomen toelaat, hetgeen het risico van onverantwoorde uitspraken verhoogt.

De gebruikte toetsingsprocedure heeft een onbetrouwbaarheid van 0.05, dat wil zeggen dat, in lange reeksen onafhankelijke toepassingen van de toets in gevallen waar de nulhypothese waar is, met frekwentie van 1 op de 20 keer het foutieve resultaat "verwerp de nulhypothese" gegeven wordt.

In het onderhavige geval hebben we evenwel te maken met reeksen herhalingen van dezelfde toets die zeker niet onafhankelijk zijn. Eerstens zijn er reeksen toetsen betreffende eenzelfde variabele van de 2e categorie, en verschillende van de 1e categorie, reeksen voor één variabele van de 1e en verschillende van de 2e categorie, terwijl daarenboven variabelen van de 2e categorie meestal niet onderling onafhankelijk zijn. (Naar samenhang van variabelen van de 1e categorie is apart gekeken, zie §5.4).

Dit kan tot gevolg hebben dat, wanneer men ten onrechte, door ongelukkig toeval, een nulhypothese verwerpt, dezelfde fout doorwerkt bij het toetsen van volgende nulhypothesen, zodat ook die ten onrechte verworpen worden, terwijl men licht geneigd is aan te nemen dat waar zoveel hypothesen verworpen worden er toch zeker een groot aantal terecht verworpen wordt.

#### 4.3. Niet longitudinale proefopzet

Door de niet longitudinale proefopzet (§2) wordt naast het leeftijds-effekt een daarvan niet te onderscheiden geboortedatum-effekt geïntroduceerd. Ook andere effecten, zoals selectie door sterfte of verwijdering uit het waarnemingsmateriaal wegens hartziekte e.d., kunnen schijneffecten teweeg brengen.

De noodzaak hiermee rekening te houden blijkt uit het bekende voorbeeld dat de lengte van oudere mensen in Nederland gemiddeld minder is dan die van jongere. Het is niet juist hieruit te konkluderen, dat naarmate een mens ouder wordt, hij ook korter wordt. (Voor een discussie van zulke effecten zie Hemelrijk (1972)). Het begrip "verloop met de leeftijd" moet dus met de nodige omzichtigheid worden gehanteerd.

In het voorbijgaan kan nog worden opgemerkt dat bovenstaande overwegingen een extra aansporing zijn het criterium (3.8) te verkiezen boven (3.10)

#### 4.4. Apologie

De eerder geschetste werkwijze is, ondanks bezwaren, aangehouden, daar het onderzoek gebruikt werd als detektiemiddel. Sommige van de resultaten zullen dan ook de toets van nader onderzoek wellicht niet kunnen doorstaan. Daarbij dient dan het uit te voeren onderzoek, inclusief de statistische analysemethoden, van te voren nauwkeurig te zijn beschreven.

#### 4.5. Alternatieven

Om de in §4.2 besproken afhankelijkheid van toetsen te ondervangen, ligt een multivariate analysemethode voor de hand, waarin immers expliciet rekening wordt gehouden met de afhankelijkheid van de behandelde problemen. Door het grote aantal variabelen lijkt een dergelijke analyse als detektiemiddel minder geschikt te zijn.

In plaats van slechts te onderscheiden, of een persoon grote dan wel kleine x-waarde voor zijn leeftijd heeft, kan men ook trachten rekening te houden met hoe groot die x-skore dan wel is. Dat zou kunnen leiden tot een model van de vorm

$$\underline{y}(b) = \alpha' + \beta' l(b) + \gamma' q(b) + \underline{\eta}(b)$$

met  $q$  gedefinieerd als in (3.7), waarbij de hypothese  $\gamma' = 0$ .

Zonder verder nog in te gaan op voor- en nadelen van dit model, vermelden wij dat besloten werd tot (3.5), omdat de resultaten daarbij eenvoudig, zelfs grafisch, weer te geven zijn, terwijl er bovendien geen reden was om aan te nemen dat de genoemde alternatieve methode tot essentieel andere conclusies zou hebben geleid.

### 5. EEN PAAR KANTTEKENINGEN

Tijdens het onderzoek werd nog aandacht besteed aan een aantal zaken die enig licht werpen op de adekwaatheid van de modellen (3.5) en (3.9), en

op de relatie tussen variabelen van de 1e categorie.

### 5.1. Knippunt in de regressielijn

Daar het niet onmogelijk leek dat voor sommige variabelen de regressielijn op de leeftijd zou verlopen als in figuur 5.1, d.w.z. met een knik ergens tussen 20 en 70 jaar, werd bekeken of zo'n knik kon worden aangetoond op het (arbitraire) punt van 45 jaar. Zonder in te gaan op een precieze formulering van het gebruikte model vermelden we dat in geen der gevallen een knik kan worden aangetoond, behoudens bij de variabelen "cardiothoracale index" en "RRsyst bij inspanning" (zie ZUIDERVELD(1974)). Tijdens het verdere onderzoek werden de modellen (3.5) resp. (3.9) aangehouden, resultaten voor deze twee variabelen moeten echter met argwaan worden beschouwd.

Als bijproduct van deze analyse werden voor variabelen  $y$  van de 2e categorie schattingen voor  $\alpha$  en  $\beta$  in (5.1) verkregen:

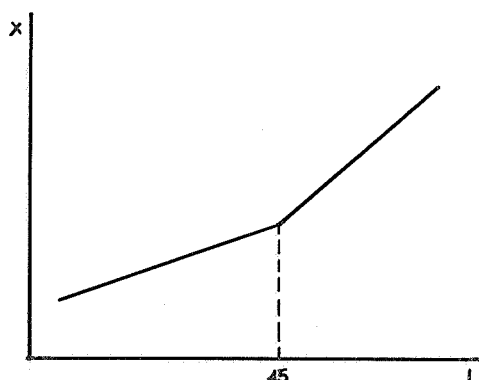


fig. 5.1 regressielijn met knik

$$(5.1) \quad \left. \begin{array}{l} y(p_j) = \alpha + \beta I_j + \theta_j \\ \theta_j \sim N(0, \sigma_\theta^2), (\theta_i, \theta_j) \text{ stoch. onafh.} \end{array} \right\} i, j = 1, \dots, 100; i \neq j$$

Dit geeft interessant vergelijkingsmateriaal met de schattingen voor  $\alpha_1, \beta_1, \alpha_2$  en  $\beta_2$  in (3.5).

### 5.2. Verlopende variantie

In alle modellen werd de veronderstelling gemaakt dat de restterm een

variantie bezit, die niet afhangt van de waarde van 1(b). Is dit toch het geval dan zijn de gebruikte toetsen niet geheel korrekt.

Om een indruk te krijgen werd voor elke variabele de steekproef-varian-  
tie in leeftijdsklassen van 10 aansluitende jaren berekend.

Met de toets van Hartley (DE JONGE(1960)) werd de gelijkheid van variantie nagegaan. Met een onbetrouwbaarheid van 0.01 werd in twee gevallen de hypo-  
these van gelijke variantie verworpen. Het betreft de variabelen (zie ZUIDERVELD(1974)) "diastolische bloeddruk 's avonds" en "BCG ratio". De resultaten van toetsing van regressie-koefficiënten is dus verdacht in deze gevallen.

### 5.3. Samenhang van variabelen van de 1e categorie

Om na te gaan of er veel verband is tussen de splitsingen teweegge-  
bracht door de verschillende variabelen van de eerste categorie, is de ma-  
trix van partiële korrelatie-koefficiënten onder konstant houden van de  
leeftijd berekend.

Daaruit blijkt dat de splitsing naar gewicht en naar overgewicht sterk sa-  
menhangen, evenals die naar de beide bloeddrukken.

#### Partiele korrelatie-matrix van variabelen van de 1e categorie onder konstanthouden van leeftijd.

naam van de variabele

| (zie ZUIDERVELD (1974))              | (a)   | (q)   | (o)   | (tc) | (RRS) | (RRD) | (r)  |
|--------------------------------------|-------|-------|-------|------|-------|-------|------|
| aktiviteitsnivo (a)                  | 1,00  |       |       |      |       |       |      |
| gewicht (q)                          | 0,12  | 1,00  |       |      |       |       |      |
| overgewicht (o)                      | 0,12  | 0,82  | 1,00  |      |       |       |      |
| totaal cholesterolgehalte (tc)       | -0,18 | 0,05  | 0,13  | 1,00 |       |       |      |
| bloeddruk in rust, systolisch (RRS)  | -0,26 | 0,09  | 0,33  | 0,27 | 1,00  |       |      |
| bloeddruk in rust, diastolisch (RRD) | -0,37 | 0,22  | 0,35  | 0,32 | 0,73  | 1,00  |      |
| roken (r)                            | -0,20 | -0,08 | -0,27 | 0,09 | -0,27 | -0,11 | 1,00 |

## LITERATUUR

- [1] J. HEMELRIJK, *Statistiek te pas en te onpas* (1972). Hoofdstuk 12
- [2] H. DE JONGE, *Inleiding tot de medische statistiek*, deel 2 (1960)
- [3] U. ZUIDERVELD, *Honderd gezonde mannen, een onderzoek naar het cardio-vasculaire verouderingsproces met een non-invasieve meetmethoden* (1974)