

De Discussie-rubriek

Een noot bij Wilmink's clique detektie

Bij het artikel van F.W.Wilmink over een (nieuwe) procedure voor clique detektie kunnen de volgende kanttekeningen geplaatst worden.

- Recentelijk zijn efficiëntere algorithmen gepubliceerd, waarvan vooral dat van Bron en Kerbosch (ACM algorithm 457) de aandacht verdient. Ook Bierstone's algoritme wordt door Wilmink niet vermeld.
 - in 1972 is door mij in het kader van een onderzoeksproject naar overeenkomst in stemgedrag (Roll Call Analysis) een Algol programma voor clique detektie ontwikkeld. Dit 1-clique algoritme volgens Peay (1970) is door F.J.A.M. van Veen uitgebreid voor n-cliques en geconverteerd naar Fortran. Zie verder Program Bulletin no. 18, revisie B 1976 van het Technisch Centrum FSW.
- Het clique detektie deel hiervan is volgens mij in essentie gelijk aan zijn procedure. Er aan vooraf gaat een eventuele inkrimping van de similarity matrix (T), als er geïsoleerde punten optreden.
- de permutatiestap is een extra stap waarvan bepaalde resultaten zijn te verwachten. Echter twee overwegingen spelen hierbij een rol.
 1. zoals opgemerkt worden de cliques eerder gevormd als de nullen ("de informatie" over welke elementen niet bij elkaar in een clique kunnen) in de bovenste rijen voorkomen. De cliques worden eerder gevormd.
 2. het aantal cliques gevormd tijdens het proces (door ons ontwerpcliques genoemd) is echter in het algemeen groter. Het aantal ontwerpcliques is - zoals destijds door metingen is vastgesteld en ook door Wilmink wordt aangegeven met de parameter DIFMAX - evenzeer bepalend voor de efficiëntie van het proces. In alle stappen moeten ontwerpcliques uit vorige stappen onderzocht worden, tenminste op het erin voorkomen van element i (in de i-de stap). Bovendien neemt het aantal ontwerpcliques zo sterk toe met n (orde n^2) dat het beperken van dit aantal aan te bevelen is en belangrijker schijnt dan het effect van 1.
- Destijds is echter gemeten dat extra eliminatiestappen meer tijd vragen dan het meenemen van teveel (deel)cliques.

Nu is het dus zo dat het aantal ontwerpcliques groter wordt naarmate meer nullen in de bovenste rijen voorkomen. Na het lezen van Wilmink's artikel was mijn hypothese dan ook dat 'permutatie' het proces zou kunnen versnellen door juist de nullen in de onderste rijen te laten voorkomen; dat wil zeggen T sorteren volgens afnemende rijtotalen, net omgekeerd als bij Wilmink. Deze veronderstelling heb ik kunnen schragen met enige experimenten met ons programma Cliqan, bestaande uit

- a. generatie van een n-orde symmetrische, aselecte 0/1 matrix met verhouding $n(0):n(1) = .777:1$ (door Wilmink opgegeven verhouding nullen en enen)
- b. sorteren van deze matrix volgens
 - oplopende rijtotalen (A van ascending)
 - afnemende rijtotalen (D van descending)
- c. clique detektie op matrix A, D en de ongesorteerde vorm (O)

In onderstaande tabel zijn de gemeten rekentijden in seconden gegeven

sort	n	30	40	60	60	60	60	60	60
A		7.6	31.5	185	179	236	181	152	149
D		7.3	27.3	164	150	175	149	138	127
O		7.5	28.5	173	154	190	154	146	148
A-D		.04	.15	.12	.19	.34	.21	.10	.17

Toegegeven dat het onderzoek beperkt is in de zin dat

- slechts één 0/1 verhouding is onderzocht
 - géén eliminatie tijdens de iteraties heeft plaats gevonden
- kunnen wij hieruit konkluderen dat in de gegeven omstandigheden:

1. $D \leq O \leq A$ zodat het effect van sorteren, als het bestaat, optreedt bij aflopende rijtotalen, en niet bij oplopende zoals Wilmink veronderstelt.
2. de verschillen nogal variëren en de efficiëntie dus nogal afhankelijk is van de input data, de T matrix van de onderzoeker.

Samenfattend zouden wij willen stellen dat clique detectie een vrij gecompliceerd proces is waarin zoveel factoren bepalend zijn voor de efficiëntie, dat slechts een uitgebreid onderzoek van betrokken algorithmen een zinvolle vergelijking mogelijk maakt.

17.10.77

Willem H. van Hoboken

De onbruikbaarheid van vereenvoudigde gewichten in
multipele regressieanalyses^{x)}.

J. Hagnaars

Th. van der Net

Een eerste, vluchtige lezing van het artikel van Pijl en Nijssse leverde een soort reactie op van: onzin!, vermoedelijk voortkomend uit het idee: dat kan (mag?) niet waar zijn en uit de wat eenzijdige preoccupatie van Pijl/Nijssse met "voorspellen".

Een hernieuwde grondige lezing van het artikel, riep eerder een soort reactie op van: dus toch! Immers, een dermate grote vereenvoudiging van gewichten als toegepast door Pijl en Nijssse, die toch ongeveer dezelfde multipele korrelatie oplevert, dat zegt weer eens alles over onze verfijnde meet- en analysetechnieken.

Dit kommentaar is weer het resultaat van een derde fase, waarbij de strekking van het artikel van Pijl en Nijssse serieus wordt genomen, terwijl er een aantal kanttekeningen gemaakt worden die de bruikbaarheid van deze vereenvoudigde procedure toch in twijfel trekken.

Een eerste opmerking heeft te maken met een zekere diskrepan-
tie tussen de korrelaties zoals vermeld in tabel 2 en de
resultaten van tabel 3.

Zonder alles na te rekenen valt het op dat

$r_{7-15} = r_{12-15} = .61$, terwijl de partiele gestandaardiseerde regressiegewichten ongelijk zijn aan elkaar (resp. .50 en .37 bedragen). Eveneens inkonsistent is het feit dat

$r_{7-15} < r_{13-15}$, terwijl wederom uitgedrukt in gestandaardiseerde regressiegewichten het effect van variabele 7 groter

x) Kommentaar op: S.J. Pijl en M. Nijssse, "Een empirisch onderzoek naar de bruikbaarheid van vereenvoudigde gewichten in multipele regressie", MDN, jrg. 2, nr. 3, 1977, blz. 99-107.

is dan van variabele 13. (Ook alle multipele korrelaties voor de eerste en derde opzet zijn hierdoor (logisch) onmogelijk, n.l. te hoog).

Is hier slechts sprake van een drukfout in tabel 2 (b.v. r_{7-15} moet in feite .68 zijn) of is er meer aan de hand?

Uitgaande van het feit dat de weergegeven resultaten (minstens in grote lijnen) korrekt zijn, kan als belangrijk bezwaar tegen de gevolgde werk- en redeneerwijze ingebracht worden dat deze inductief van aard is. In feite gaat het om de redenering dat in een aantal gevallen het empirisch zo blijkt te zijn dat het niet zo veel verschil maakt of men nu ongelijke, dan wel gelijke gewichten gebruikt. Maar om praktisch bruikbaar te zijn moet een onderzoeker in zijn konkrete geval weten of het verantwoord is om gelijke gewichten te gebruiken. Zo is de aanbevolen methode niet bruikbaar in het geval er suppressorvariabelen optreden. Dat Pijl en Nijssen vermelden dat deze niet zoveel voorkomen (waar halen ze dat vandaan) is voor een bepaalde onderzoeker een geringe troost. Zijn vraag is: komen ze bij hem voor, en om dat te weten zal hij toch meestal de "gewone" regressie moeten gebruiken. Zolang niet vaststaat met een redelijke precisie onder welke omstandigheden deze vergroving geen kwaad kan, heb je er praktisch niet zo veel aan. Werkt het ook bij relatief onafhankelijke variabelen (tabel 3 suggereert dat het dan in mindere mate het geval is); werkt het ook bij een regressievergelijking met 7 of meer onafhankelijke variabelen waarvan er diverse zeer lage regressiegewichten hebben, die echter tesamen nog een significant deel van de verklaarde variantie voor hun rekening nemen? Totdat een onderzoeker met een grote mate van waarschijnlijkheid kan zeggen dat in zijn bepaalde geval gelijke gewichten "hetzelfde" resultaat oplevert als de "gewone" gewichten, tot zolang is deze methode niet als praktisch bruikbaar te bestempelen.

Het doet zeer vreemd (en onverwacht) aan dat de regressiegewichten, afgerond op 1 decimaal een minder nauwkeurig beeld opleveren dan de gelijke-gewichten-methode. De door Pijl en Nijssse gegeven verklaring dat "de predictoren, welke gewogen worden met een waarde net onder de .05, wegvallen in de regressie-vergelijking, terwijl deze predictoren, gewogen met de oorspronkelijke gewichten, wel degelijk een behoorlijk deel aan het percentage verklaarde variantie bijdragen" (blz. 105) voldoet niet erg. Dit alles geldt immers a fortiori voor de vergelijking met gelijke gewichten.

Het lijkt dan ook waarschijnlijk (hoewel in dit specifieke geval tegengesproken) dat i.h.a. afronding op 1 decimaal nauwkeuriger is dan de gelijke gewichten.

Een aspekt waar Pijl en Nijssse niet op wijzen is de moeilijkheid met behulp van de gelijke gewichtenmethode de oorspronkelijke ongestandaardiseerde regressiecoëfficiënten terug te rekenen, terwijl deze toch vaak van belang zijn. Immers indien de gelijke gewichten weer worden omgerekend, dan kunnen onvermoede en onwenselijke zaken passeren indien de varianties onderling sterk verschillen.

Tenslotte, Pijl en Nijssse merken op: "de vrij grote accurateid bij het gebruik van gelijke gewichten zet een vraagteken bij het toepassen en vooral het interpreteren van gestandaardiseerde kleinste kwadraten regressiegewichten" (blz. 106). Uitspraken over de grootte van de effecten van de ene variabele ten opzichte van de effecten van de andere (zoals die ook voorkomen in b.v. padanalyse) worden als uiterst bedenkelijk gekwalificeerd.

Dit nu lijkt duidelijk overtrokken, of liever stelt het probleem verkeerd. Het blijft immers waar dat indien b.v.

$p_{41} = .20$ en $p_{42} = .40$ de "beste" schatting van de

gestandaardiseerde effecten van de variabelen X_1 en X_2 op X_4 .20 resp. .40 is, en dus ook de "beste" schatting van het effect van X_1 twee maal zo groot is als de "beste" schatting van het effect van X_2 . Het feit dat een spaarzamer model, met gelijke regressiegewichten, eenzelfde aanpassing aan de data geeft, maakt dit eenvoudiger model nog niet beter! Analooq hieraan kan het b.v. bij padanalyse blijken dat een theoretisch model redelijk goed past bij de data. Meestal echter zijn er diverse concurrerende modellen op te stellen, die veelal "eenvoudiger" zullen zijn en die evengoed bij de data passen. Dit is echter nog geen bewijs voor de juistheid van de eenvoudigere modellen. In de uitspraken van Pijl en Nijssse zit impliciet de hypothese dat de enige relevante hypothese die getoetst moeten worden is: er is geen verschil tussen padcoëfficiënten.

Wordt deze hypothese niet verworpen, dan wordt gekonkludeerd dat het verschil niet relevant was. Misschien dat de toetsing: de ene coëfficiënt is veel groter dan de andere (b.v. 3 maal zo groot) eveneens een goed resultaat oplevert! Welke (alternatieve) hypothese getoetst moet worden is een theoretische kwestie, en kan nooit automatisch luiden: geen verschil.

De "werkelijkheid" gedraagt zich niet noodzakelijkerwijs volgens de idealen van "parsimony".

Onze konklusie n.a.v. het artikel van Pijl en Nijssse: zeker niet om nu de gelijke gewichten methode te gaan gebruiken of aan te gaan bevelen. Wel: alvorens uitspraken te doen over het sterker dan wel zwakker zijn van effecten t.o.v. elkaar, nagaan of de alternatieve hypothese dat ze "gelijk" (dan wel "ongelijk" met een bepaald hoeveelheid) zijn serieuzer te nemen dan tot nu toe, en meer routinematig procedures zoals b.v. beschreven in Johnston J., "Econometric Methods". 5-6 Linear Restrictions blz. 155 e.v., toe te gaan passen.

VEREENVOUDIGDE GEWICHTEN IN MULTIPELE REGRESSIEANALYSES: BRUIKBAAR EN GEBRUIKSKLAAR; EEN ANTWOORD AAN HAGENAARS EN VAN DER NET.

M. Nijssse

S.J. Pijl

De eerste opmerking van HAGENAARS en VAN DER NET, welke betrekking heeft op het bestaan van discrepanties tussen correlaties en de ermee corresponderende partiële gestandaardiseerde regressiegewichten, komt voort uit het over het hoofd zien van kolom 3 in Tabel 3. Hierin wordt vermeld op welke aantallen proefpersonen de verschillende analyses zijn uitgevoerd. De beschrijvende statistische maten, zoals $\bar{X}$, s en r zullen daardoor van analyse tot analyse enigszins van elkaar verschillen.

De gegevens van Tabel 1 en 2 geven een totaalbeeld van de beschikbare gegevens; zij zijn berekend m.b.v. een 'missing data' voorziening, waardoor de beschikbare gegevens optimaal zijn benut.

Uiteraard zijn de regressieanalyses gebaseerd op proefpersonen die op alle opgenomen variabelen een score hebben. Naar onze mening zou het te ver gevoerd hebben om voor de 19 door ons uitgevoerde analyses alle standaarddeviaties en intercorrelaties per analyse te vermelden. Dat de gepubliceerde eindresultaten meer dan slechts in grote lijnen correct zijn, moet de lezer dus in goed vertrouwen aannemen. Voor wie desondanks toch twijfels houdt ligt de computeruitvoer ter beschikking.

Met ons artikel hebben wij niet willen betogen dat het toepassen van sterk vereenvoudigde gewichten à la WAINER in alle onderzoekssituaties de gebruikelijke regressieanalyse kan vervangen. Aan het slot van ons artikel is dit ook met nadruk opgemerkt. Wel geloven wij dat ons materiaal tamelijk representatief is voor onderzoeken waarbij regressieanalyses worden uitgevoerd: predictoren, welke vrij hoog met een criterium correleren, maar waarvan de intercorrelaties eveneens tamelijk hoog zijn. In dergelijke situaties zal men zelden meer dan 6 predictoren nodig hebben om een adequate beschrijving van de data te geven. Het toevoegen van meer predictoren levert dan ook meestal geen significante bijdrage aan het percentage verklaarde variantie. In de gevallen waarin dit wel zo is, is dit vermoedelijk te danken aan het grote aantal proefpersonen, waardoor lage F-waarden toch significant zijn; de toename van het percentage verklaarde variantie is dan meestal absoluut gezien erg klein.

Suppressorvariabelen (variabelen die nul of zelfs negatief met het criterium correleren, maar die desondanks een aanzienlijke toename van R^2 voor hun rekening nemen, omdat de correlatie met één of meer andere predictoren vrij hoog is) komen volgens NUNNALLY (1967, blz. 167) in de praktijk weinig voor. In ons materiaal correleert predictor variabele 9 laag (.04) met het criterium, maar de correlaties met de andere mogelijke predictoren zijn eveneens laag (zie Tabel 2). Wanneer variabele 17 als criterium wordt genomen, vinden wij alleen voor predictor-variabele 10 een lage correlatie (-.04). Ook hier geldt dat de correlaties met vrijwel alle overige mogelijke predictoren laag zijn. Men zou kunnen overwegen variabele 14 als predictor in de laatste twee analyses op te nemen (deze correleert .31 met variabele 10 en -.28 met het criterium). Het aantal beschikbare personen daalt dan echter tot 52 (gemiddeld 13 personen per predictor variabele). Regressie-analyse wordt dan zinloos, omdat de standaarddeviaties van de regressiegewichten in een dergelijk geval zeer groot zullen zijn; over de bijdrage van variabele 14 aan R^2 bestaat dan grote onzekerheid.

De constatering van HAGENAARS en VAN DER NET dat de regressiegewichten, welke op 1 decimaal nauwkeurig worden afgerond, een minder nauwkeurig beeld opleveren dan wanneer van gelijke gewichten gebruik wordt gemaakt, is op zich juist maar tevens nietszeggend. Er worden immers geheel verschillende criteria gehanteerd voor het laten wegvallen van een variabele, te weten: het ongestandaardiseerde gewicht heeft een lagere waarde dan .05 resp. de correlatie met het criterium is niet significant.

De opmerking over de onmogelijkheid om de ongestandaardiseerde gewichten te kunnen terugrekenen is niet juist; er geldt immers: $b_i = \beta_i \left(\frac{s_y}{s_i} \right)$. Een consequentie van onze uitkomsten is dat het in het algemeen niet zo belangrijk is of men voor β_i het partieel gestandaardiseerde gewicht dan wel een sterk vereenvoudigd gewicht kiest.

De strekking van de rest van het commentaar van HAGENAARS en VAN DER NET, dat over padcoëfficiënten en het kiezen tussen concurrerende modellen gaat, is ons niet erg duidelijk geworden. We hebben opgemerkt dat wij de consequenties, welke het prefereren van gelijke gewichten boven differentiële gewichten voor bijv. pad-analyses heeft, op dit moment niet overzien. Overigens lijkt het ons evident dat wanneer een eenvoudig model de data even goed blijkt te beschrijven als een meer ingewikkeld model, het eenvoudige model altijd de voorkeur verdient. Wanneer een eenvoudig model minder bevredigend past bij een bepaalde theorie dan een meer ingewikkeld model, terwijl beide modellen een adequate beschrijving van de data opleveren, dan wordt het tijd om de theorie te herzien en niet het eenvoudige model.

Tot slot: wij hebben nooit beweerd dat het WAINER-recept zich leent voor critiekloze toepassing op alle soorten materiaal. Wij geloven wel dat in veel gevallen gelijke gewichten even goed zullen blijken te werken als differentiële. (De titel van onze reactie is dus enigszins overtrokken). Van een eenzijdige preoccupatie met 'voorspellen' kan men ons evenmin beschuldigen. Het tegendeel is het geval: wij geloven juist dat de gelijke gewichten-methode belangrijke consequenties heeft voor de theorievorming. Belangstellenden wordt aangeraden de eerstkomende nummers van Psychological Bulletin in de gaten te houden; er is nl. een reactie op het artikel van WAINER onderweg. Verder zij verwezen naar een artikel van B.F. GREEN Jr. in Multivariate Behavioral Research, 1977, 12, 263 - 187, getiteld Parameter Sensitivity in Multivariate Methods.

Literatuur: NUNNALLY, J.C. (1967). Psychometric theory. New York:Mc Graw-Hill.

Reakties op "Het schatten van factorscores"

W.E.Saris

Naar aanleiding van het artikel van J.Mulder e.a "Het schatten van factorscores zijn drie reakties binnen gekomen waarvoor we zeer erkentelijk zijn.

De reactie van P.Vijn is opgenomen in dit nummer als een zelfstandig artikel. De reactie van J.ten Berge is hierna opgenomen. Zijn opmerkingen leveren een verbetering van de Generalized Anderson en Rubin procedure op. We hebben deze verbeteringen ook geïmplementeerd in het programma.

Verder ontvingen we van G.J. van Driel een mededeling dat hij in zijn proefschrift het zelfde probleem had aangeroerd en onder andere ook een schatter had afgeleid die voldoet aan dezelfde voorwaarden als de G.A.R methode. Wij waren van zijn werk niet op de hoogte en aangezien misschien ook andere mensen hierin geïnteresseerd zijn zullen we de volledige titel vermelden: G.J.van Driel Enige factor-analytische technieken, toegepast op examenresultaten. Universitaire pers Rotterdam verschenen in 1975.

Reaktie op Mulder c.s.: Het schatten van faktorscores

J.M.F. ten Berge, Subfakulteit Psychologie, R.U. Groningen.

Mulder c.s. (1977) introduceren de GAR-methode voor het schatten van faktorscores. Zowel in de afleiding als in de resulterende oplossing zijn verbeteringen aan te brengen. Het probleem is, in de notatie van Mulder c.s., een matrix A te vinden waarvoor de functie

$$g(A) = \text{tr}(A' \Sigma A - 2 A' \Lambda \Phi + \Phi) \quad (1)$$

minimaal is onder de randvoorwaarde

$$A' \Sigma A = \Phi. \quad (2)$$

Dankzij (2) is minizering van (1) equivalent aan maximering van

$$h(A) = \text{tr} A' \Lambda \Phi. \quad (3)$$

$$\text{Laat } \Sigma = P P' \quad (4)$$

$$\text{en } \Phi = Q Q' \quad (5)$$

benedendriehoeksontbindingen zijn, en laat

$$P^{-1} \Lambda \Phi Q = U D V' \quad (6)$$

met $U'U = V'V = VV' = I$ en D diagonaal, $d_{ii} \geq 0$.

een singuliere waarden- of Eckart-Young ontbinding zijn. Dan

is die A waarvoor (3) maximaal is onder (2) te vinden als

$$A = A^* \stackrel{\text{def}}{=} (P')^{-1} U V' Q'. \quad (7)$$

Bewijs: Herschrijven van $h(A)$ levert

$$h(A) = \text{tr} A' \Lambda \Phi = \text{tr} A' P P^{-1} \Lambda \Phi Q Q^{-1} = \text{tr} A' P U D V' Q^{-1} \quad (8)$$

hetgeen via cyclische permutatie kan worden omgezet in

$$h(A) = \text{tr} \left\{ (D^{\frac{1}{2}} V' Q^{-1} A' P) (U D^{\frac{1}{2}}) \right\}. \quad (9)$$

Voor de spoorfunctie in (9) geldt de ongelijkheid van Schwarz:

$$h(A) \leq \left\{ \text{tr} D^{\frac{1}{2}} V' Q^{-1} A' P P' A (Q')^{-1} V D^{\frac{1}{2}} \right\}^{\frac{1}{2}} \left\{ \text{tr} U D U' \right\}^{\frac{1}{2}} \quad (10)$$

hetgeen te vereenvoudigen is tot

$$h(A) \leq \left\{ \text{tr} D \right\}^{\frac{1}{2}} \left\{ \text{tr} D \right\}^{\frac{1}{2}} = \text{tr} D. \quad (11)$$

Aangezien $\text{tr} D$ konstant is onder variatie van A is het voldoende voor een maximum van (3) over een A te beschikken waarvoor $h(A) = \text{tr} D$ terwijl (2) vervuld is. Het is gemakkelijk te verifiëren dat $h(A^*) = \text{tr} D$ en dat $A^{*'} \Sigma A^* = \Phi$, zodat A^* uit (7) de gezochte oplossing is.

Oplossing (7) is, in vergelijking met de oplossing van Mulder c.s. aanzienlijk minder bewerkelijk. Bovendien staat de hier gepresenteerde bewijsvoering er voor garant dat met (7) ook inderdaad de maximale funktiewaarde van $h(A)$ wordt bereikt. Zo'n garantie ontbreekt bij de oplossing van Mulder c.s. Weliswaar voldoet deze oplossing aan

een noodzakelijke voorwaarde voor een maximum (cf. de normaalvergelijking A6) maar er zijn meerdere oplossingen die aan die voorwaarde voldoen. Nemen we bijv.

$$A = (P')^{-1} U T V' Q' \quad (12)$$

waarin T een willekeurige tekenmatrix is, dan is voor elke keuze van T aan de normaalvergelijking voldaan, en wel met bijbehorende funktiewaarde $h(A, T) = \text{tr } T D$. Bij elke keuze van T behalve $T = I$ zal slechts een lokaal maximum of zelfs een minimum worden bereikt, aangenomen dat D geen gelijke diagonaalelementen bevat. Het is allerm minst duidelijk of de oplossing van Mulder c.s. van dit soort problemen gevrijwaard is.

Hakstian (1975) schreef in de kontekst van factor-matching over Φ -constrained factor transformation. De overeenkomst tussen Φ -constrained en struktuurbewarende transformaties is niet louter taalkundig van aard: Hakstian lost een probleem op wat een speciaal geval is (nl. met $\Sigma = I$) van het maximeren van (3) onder randvoorwaarde (2). Hakstian's oplossing maximeert inderdaad de betreffende funktie, voornamelijk dankzij stap 4 op pag. 251. De noodzaak van deze stap wordt overigens door Hakstian niet aangetoond.

Literatuur

- Hakstian, A.R. Procedures for Φ -constrained confirmatory factor transformation. Multivariate Behavioral Research, 10, 2, 1975, 245-253.
- Mulder, J. de Pijper, W.M., en Saris, W.E. Het schatten van factor-scores. MDN 3, 2, juni 1977, p. 56-66.