

HOOFSTUK 1

Inleiding en Agtergrond

By empiriese navorsing is vergelyking van groepe baie keer van belang, of andersins die bepaling van verbande tussen veranderlikes wat gemeet word. Daar word dan gevra na betekenisvolheid tussen gemiddeldes of betekenisvolheid van 'n verband. Onder “betekenisvolheid” word gewoonlik verstaan dat 'n sg. “nulhipotese” van geen verskil in gemiddeldes (of geen verband), verwerp word op 'n voorafbepaalde betekenispeil (gewoonlik 5%). Anders gestel: dat die sg. “p-waarde” kleiner is as 0,05. Hierdie “betekenisvolheid”, of ook genoem “statistiese betekenisvolheid”, beteken eintlik maar net dat die waarskynlikheid klein (bv. $\leq 0,05$) is dat die nulhipotese verkeerdelik verwerp kan word. Dus dui dit daarop dat die verband of verskille wat die waarskynlikheidsteekproef of -steekproewe oplewer, nie maar 'n toevalligheid is nie, want die kans vir so 'n toevalligheid is klein (sê 5%). Wat dit egter *nie* sê *nie* is *hoe belangrik* die verskille of verband wel is. Om uitspraak oor die belangrikheid van verskille of verbande te gee, word van *effekgrootte-indekse* gebruik gemaak.

Dit is die doel van hierdie handleiding om effekgrootte-indekse te gee en te bespreek by die meeste van gevalle waarmee navorsers in empiriese navorsing te doen kry.

1.1 Praktiese betekenisvolheid

Soos met statistiese betekenisvolheid, waarby die nulhipotese verwerp moet word, is die vraag: wanneer is 'n verskil of verband groot genoeg om belangrik te wees? Hiervoor kan effekgrootte-indekse gebruik word in die sin dat sulke

indekse se grootte direk eweredig is aan die belangrikheid van 'n verskil in gemiddeldes of 'n verband tussen veranderlikes. As 'n indeks groot genoeg is, kan ons die resultaat as *prakties betekenisvol* bestempel. Hierdie is 'n algemene term wat in verskillende kontekste gebruik kan word. By kliniese proewe staan dit bv. bekend as “klinies betekenisvol” en as die Opvoedkundige dit gebruik, kan daarna verwys word as bv. “opvoedkundig betekenisvol”.

Pogings om riglynwaardes te koppel aan effekgrootte-indekse, om as sg. afsnypunte te dien vir “klein”, “medium” en “groot” effekte is deur skrywers soos Cohen (1969, 1977, 1988), aangewend. Weens die feit dat sulke waardes arbitrêr kan wees, is daar ook baie kritiek daarteen. Daar sal deurgaans in hierdie handleiding sulke riglynwaardes gegee word, met nodige motivering daarby, maar tesame met die kritiese evaluering daarvan.

1.2 Wanneer is effekgrootte-indekse nodig?

Gestel dat die gemiddelde diastoliese bloeddruk van 25 hipertensiewe pasiënte met bv. 10 mmHg verlaag na 'n sekere behandeling. Gestel hierdie verlaging was statisties betekenisvol op 'n 1%-peil. Die navorser kan op grond hiervan sê dit was ook 'n prakties betekenisvolle verlaging, want hy/sy beoordeel dit op die bekende mmHg-skaal waarby dit bekend is dat 10 eenhede groot genoeg is om belangrik te wees. Hier is 'n effekgrootte-indeks nie nodig nie. 'n Ander voorbeeld is die verkryging van 'n hoogsbetekenisvolle korrelasie ($p < 0,0001$) van 0,8 tussen 'n nuwe psigometriese toets en 'n standaardtoets om bv. depressie te meet, op 'n proefgroep van 200 persone, wat daarop dui dat die nuwe toets geldig is vir die proefgroep, want 'n 0,8-korrelasie is uit die sielkundige se ervaring groot genoeg om op geldigheid te dui.

By beide die voorbeelde was kennis van die skaal of vorige ervaring genoegsaam om sondermeer 'n uitspraak te kon gee oor praktiese

betekenisvolheid. Die bepaling van 'n effekgrootte-indeks en evaluering daarvan t.o.v. riglynwaardes was dus nie nodig nie.

Daar is egter gevalle waar effekgrootte-indekse wel nodig is om te bepaal, ten einde uitspraak oor praktiese betekenisvolheid al dan nie, te kan gee. Hier volg 'n lys van sulke gevalle:

Geval 1:

Die *skaal* waarop verskille in gemiddeldes gemeet word, is *onbekend*. Bv. in 'n ongestandaardiseerde vraelys waar vrae op 'n 4-punt-Likertskaal gevra word, is dit nie duidelik wat 'n verskil in gemiddeldes van 0,3 beteken nie. By gestandaardiseerde skale soos bv. stanage- of stienskaal ken die navorser vooraf die standaardafwyking van die skaal en kan verskille in die konteks daarvan beoordeel word, sodat effekgrootte-indekse oorbodig kan wees. Werk die navorser met die sg. rou-tellings voor standaardisasie, is die skaal gewoonlik minder bekend.

Geval 2:

Variasie van metings op 'n bepaalde skaal wissel afhangende van situasie of subjekte. So kan 'n geselekteerde groep studente se standaardafwyking van IK's heelwat kleiner wees as bv. 'n groep uit die breë samelewing. Omdat effekgrootte-indekse kompenseer vir die variasie in metings, sou die gebruik daarvan in hierdie geval verkieslik wees om praktiese betekenisvolheid te bepaal.

Geval 3:

Waarskynlikheidsteekproewe (soos ewekansige- of gestratifiseerde steekproewe) word uit populasies getrek en 'n statisties-betekenisvolle resultaat word verkry, maar onsekerheid oor die belangrikheid van verskille of verbande bestaan. Hier word effekgrootte-indekse bereken as 'n *tweede stap* nadat vir statisties-betekenisvolheid getoets is. Indien met *groot* steekproewe gewerk word, is die resultaat baie keer statisties betekenisvol, sodat die effekgrootte dan

daarop gaan dui of dit ook prakties belangrik is. By *klein* steekproewe is 'n statisties-betekenisvolle verskil (of verband) gewoonlik ook 'n belangrike verskil (of verband) en is die bepaling van 'n effekgrootte-indeks meesal onnodig. Lewer dit egter 'n statisties nie-betekenisvolle resultaat op, kan die effekgrootte wel op 'n belangrike verskil of verband dui. Dit kan twee dinge beteken: Eerstens dat die gerealiseerde resultaat maar toevallig so sterk is, of, tweedens dat daar wel 'n groot verskil of verband bestaan, maar dat die steekproewe te klein is om dit te bevestig. Lg. sou aanleiding kon gee tot herontwerp van die navorsing – groter steekproewe en dalk akkurater metings. Omdat klein steekproewe baie keer by *nuwe navorsing* in *loodsstudies* gebruik word, gee effekgroottes aanduidings of die navorsing voortgesit moet word of nie. Die volgende tabel gee die potensiele probleem wanneer gevolgtrekkings uit data gemaak moet word as a funksie van die effekgrootte en betekenispeil (Rosenthal et. al., 2000:4).

Tabel 1: Potensiele inferensiële probleme as 'n funksie van effekgroottes en betekenispeil

	Effekgrootte: "Aanvaarbaar" (groot genoeg)	Effekgrootte: "Onaanvaarbaar" (te klein)
Betekenispeil: "Aanvaarbaar" (klein genoeg)	Geen probleem	Verwar statistiese betekenisvolheid met praktiese belangrikheid
Betekenispeil: "Onaanvaarbaar" (te groot)	Onvermoë om praktiese belangrikheid van nie-betekenisvolle resultate raak te sien	Geen probleem

Geval 4:

Wanneer *meta-analise* gedoen word, is effekgroottes nodig om resultate van verskillende studies saam te voeg. Omdat bepaling van praktiese betekenisvolheid nie die oogmerk by meta-analise is nie, gaan ons nie verder op hierdie gebruik van effekgroottes in nie. Die boeke van Rosenthal (1991),

Hedges & Olkin (1985) en Hunter & Schmidt (2004) kan hiervoor geraadpleeg word.

Geval 5:

Wanneer die *gerealiseerde onderskeidingsvermoë* ("power") van 'n statistiese toets na afloop van die eksperiment (d.i. post-hoc) bepaal wil word, is die effekgrootte nodig. Cohen (1969, 1977, 1988) gee hiervoor tabelle bykans by elke soort statistiese toets waarby die effekgrootte o.a. benodig word, as inset.

Geval 6:

By die beplanning van 'n studie kan die *steekproefgrootte* bepaal word wat nodig is om 'n sekere onderskeidingsvermoë (gewoonlik 80%) te bereik by 'n sekere betekenispeil. Hiervoor is die effekgrootte nodig en as dit nie bekend is nie, moet dit beraam word m.b.v. 'n loodsstudie. Hier kan weereens van Cohen se tabelle gebruik gemaak word.

Geval 7:

By *volledige opnames* (sensusse) waar volledige populasies betrek word, is bepaling van effekgroottes soms al manier om praktiese betekenisvolheid te beoordeel. In die praktyk kom volledige opnames baie voor. 'n Paar voorbeelde is (Steyn, 1999):

- (a) 'n Studie om verkragters en gewapende rowers uit 3 beskikbare gevangnisse te vergelyk. Albei die populasies van verkragters en gewapende rowers was so klein, dat almal betrek kon word.
- (b) Eerstejaar Psigologiestudente ondergaan 'n psigometriesse toets om die geldigheid daarvan te bepaal. In plaas daarvan om 'n ewekansige steekproef te trek en net dié studente te toets, is dit in die praktyk makliker om die hele klas te toets. Die klas is nou 'n studiepopulasie wat volledig getoets word. Al sou 'n ewekansige steekproef gebruik word, sou slegs veralgemeen kon word na hierdie

studiepopulasie. 'n Volledige toetsing gee dus die presiese resultate aangaande dié populasie.

- (c) By die uitstuur van 'n vraelys aan 'n waarskynlikheidsteekproef uit 'n sekere teikenpopulasie, is die respons so swak (sê 20%) dat waar die steekproef verteenwoordigend was van die populasie, is die respondente dit nie meer nie, want hulle het hulleself gekies toe hulle gerespondeer het. Die navorser word nou gedwing om die respondente as 'n sub-populasie van die teikenpopulasie te beskou. Nou vorm die respondente 'n volledige opname of sensus van die sub-populasie van respondente en kan nie meer veralgemeen word na die teikenpopulasie nie, al was dit oorspronklik die plan.
- (d) In 'n studie waar jongmense met aanpassingsversteurings (eksperimentele groep) met 'n kontrolegroep vergelyk word, was dit moeilik om persone vir die eksperimentele groep te kry en is al die beskikbare persone oor 'n sekere tydperk wat by 'n aantal instansies werkzaam is, psigometries getoets – dus 'n volledige populasie. Kontrole persone uit dieselfde werksplekke en ouderdomme as eksperimentele persone, is egter ewekansig gekies.
- (e) Die volledige kliënte-databasis van 'n bank is elektronies beskikbaar en daarmee word op grond van korrelasies, risiko-faktore geïdentifiseer. Met huidige rekenaar-tegnologie tot 'n navorser se beskikking, is dit geen probleem om selfs die miljoene kliënte se data te ontleed nie. Al sou die studie beperk word tot sê 'n ewekansige steekproef van 10 000 kliënte, is dié steekproef so groot dat enige korrelasie statisties betekenisvol gaan wees. Bowendien behoort so 'n groot ewekansige steekproef die populasie baie goed te verteenwoordig, sodat daarmee gehandel kan word asof dit maar die populasie is.
- (f) Voorbeeld B in Hoofstuk 3 gee 'n verdere voorbeeld van 'n volledige populasie.

Om hierdie laaste geval van volledige populasies duidelik te onderskei, sal in hierdie handleiding in die besprekings van effekgrootte-indekse, telkens eers die indeks vir 'n populasie gegee word en daarna beramers daarvan, wanneer met waarskynlikheidsteekproewe gewerk word.

1.3 Vereistes vir goeie effekgrootte maatstawwe

Preacher & Kelley (2011) stel die volgende vereistes waaraan 'n goeie effekgrootte maatstaf behoort te voldoen:

1. Op 'n toepaslike skaal te wees. Sonder 'n interpreteerbare skaal is dit moeilik om te besluit of die effekgrootte groot genoeg is om sinvol te wees. So is die gestandaardiseerde verskil in gemiddeldes (verskil in gemiddeldes deur die standaardafwyking gedeel) onafhanklik van die skaal van meting en kan bv. gestandaardiseerde verskille in gemiddelde IK's en bloeddrukke op dieselfde wyse geïnterpreteer word. 'n Korrelasiekoëffisiënt as effekgrootte is ook 'n voorbeeld hiervan en is maklik om te interpreteer.
2. 'n Vertrouensinterval beskikbaar. As 'n effekgrootte uit 'n steekproef bepaal word, sal dit in waarde verskil van die effekgrootte van die populasie waaruit die steekproef getrek is. 'n Vertrouensinterval gee dan 'n goeie aanduiding van waardes waartussen die populasie-waarde daarvan, met 'n hoë waarskynlikheid, kan varieer.
3. Onafhanklik van steekproefgrootte. As 'n steekproef gebruik word om die populasie-waarde van die effekgrootte te beraam, behoort dit geen funksie te wees van die steekproefgrootte nie. Waardes van effekgroottes gebaseer op verskillende steekproefgroottes behoort sodoende op dieselfde wyse geïnterpreteer te kan word.
4. Onsydig, konsekwent en doeltreffend te wees. Onsydigheid behels dat die verwagte effekgrootte gebaseer op 'n steekproef gelyk is aan die populasie effekgrootte, terwyl konsekwentheid die konvergensie van die effekgrootte van die steekproef na die populasie-waarde toe verseker

indien die steekproefgrootte vergroot word. Doeltreffendheid vereis sinvolle lae variasie van die steekproef-effekgrootte, wat laer word met toenemende steekproefgrootte.

In die besprekings vanaf Hoofstukke 4 tot 8 word effekgroottes beskou wat gewoonlik aan die bogenoemde vereistes voldoen.

1.4 Gebruik van effekgrootte in toepassingsvelde van Statistiek

Weens die feit dat daar oor jare 'n debat in psigologiese tydskrifte was oor die gebruik van statistiese betekenistoetse, het die *American Psychological Association* (APA) die *Task Force on Statistical Inference* (TFSI) in die lewe geroep. Hulle verslag (Wilkinson & TFSI, 1999) lewer aanbevelings betreffende data-ontledings en sekere van die hoofaanbevelings volgens Kline (2004a:13) is:

1. Gebruik net die nodige en eenvoudigste statistiese ontledings.
2. Moenie statistiese resultate van rekenaar-uitsette rapporteer sonder kennis van die betekenis daarvan nie.
3. Dokumenteer aannames aangaande populasie-effekgroottes, steekproefgroottes of metings onderliggend aan a priori beraming van statistiese onderskeidingsvermoë ("power"). Gebruik eerder 'n vertrouensinterval betreffende waargenome resultate as om 'n post-hoc onderskeidingsvermoë te gee.
4. Rapporteer waargenome effekgroottes vir primêre uitkomst of wanneer p-waardes gegee word. Dit bevorder beter navorsing en verskaf inligting vir meta-analise later.
5. Rapporteer vertrouensintervalle van effekgroottes.
6. Bied tot 'n redelike mate bewyse aan dat statistiese aannames wat gemaak word, geld.

Let op dat effekgroottes 'n belangrike aspek van hulle aanbevelings vorm. Die vyfde uitgawe van die *APA's Publication Manual* (APA, 2001: 21-26) gee o.a. die volgende aanbevelings n.a.v. die TFSI-verslag (kyk Kline, 2004a: 13):

1. Rapporteer toepaslike beskrywende statistiek, soos gemiddeldes, variansies, groepgroottes en saamgevoegde binne-groepe variansies en kovariansies vir vergelykende studies, of 'n korrelasie matriks in 'n regressie-ontleding. Hierdie inligting is nodig vir meta-analises en verdere ontledings deur navorsers.
2. Effekgroottes behoort amper altyd gerapporteer te word. Die afwesigheid daarvan word selfs as 'n voorbeeld van 'n gebrekkige studie voorgehou.
3. Die gebruik van vertrouensintervalle word sterk aanbeveel.

Die sesde uitgawe van die *APA's Publication Manual* (APA, 2010: 33) stel verder: “.. complete reporting of all tested hypotheses and estimates of appropriate effect sizes and confidence intervals are the minimum expectations for all APA journals”.

Kline maak na aanleiding van die bg. aanbevelings van die TFSI en *APA Publication Manual* ook sy volgende aanbevelings:

1. Navorsers behoort nie 'n statisties betekenisvolle resultaat as sulks informatief te beskou nie – bv. dat dit outomaties dui op beduidendheid en herhaalbaarheid.
2. 'n Statisties nie-betekenisvolle resultaat behoort nie sondermeer geïgnoreer te word nie – bv. die nie-verwerping van die nulhipotese beteken nie noodwendig 'n populasie-effek wat nul is nie. Só kan moontlike voordelige effekte in navorsing misgekyk word.
3. Effekgroottes behoort altyd gerapporteer te word, en indien moontlik, tesame met vertrouensintervalle daarvan. Dit beteken dat effekgroottes nie net in opvolging van statisties betekenisvolle resultate bereken word

nie. Die klem behoort op effekgroottes as sulks gelê te word, sodat dit nie net gerapporteer word nie, maar ook geïnterpreteer word.

Huberty (2002) gee 'n lys van 19 tydskrifte in opvoedkundige psigologie waarin die gebruik van effekgroottes aangemoedig word. Bruce Thompson gee 'n opsomming van die vereistes aangaande effekgroottes van 9 tydskrifte op sy webbladsy by <http://www.coe.tamu.edu/~bthompson/index.htm>.

Bartlett (1997) stel in 'n redaksionele artikel dat effekgroottes veels te dikwels afwesig is in navorsingsartikels in sport- en oefen-wetenskappe. Effekgroottes behoort volgens hom as 'n belangrike deel van statistiese toetsing beskou te word. Die redakteur van *Research Quarterly for Exercise and Sport* moedig outeurs aan om na aanleiding van 'n artikel van Thomas, et.al. (1991) effekgroottes (of statistieke wat die berekening daarvan moontlik maak) in te sluit in artikels. Fern & Monroe (1996) gee aan lesers van die *Journal of Consumer Research* 'n oorsig oor die gebruik van effekgroottes in verskillende toepassings, die interpretasie daarvan en verbande tussen verskillende effekgroottes.

Hoewel bostaande 'n oorsig oor gebruik van effekgroottes in sekere toepassingsvelde van statistiek is, is geensins gepoog om 'n volledige lys te gee nie. Dit gee egter 'n aanduiding van hoe sterk die gebruik daarvan aangemoedig word deur die tydskrifte. Die gebruik van effekgroottes in navorsing sukkel nog om pos te vat. Daarvan getuig Thompson (2001) as hy na 11 empiriese studies verwys, waarin die gebruik van effekgroottes in 23 tydskrifte na 1994 ondersoek is. Kirk (1996) bevind ook dat lae persentasies artikels in 4 psigologie tydskrifte waarin statistiese inferensie gedoen is, effekgroottes bevat het. Meer onlangs het Cumming et al. (2007) ondersoek ingestel of psigoloë hul statistiese praktyke verander het sedert 1998, nadat die APA Task Force in Statistical Inference (TFSI) hul aanbevelings gemaak het waarin verbeterde statistiese praktyke aangemoedig word, insluitende die reporting van effekgroottes en vertrouensintervalle.

1.5 Effekgroottes in statistiese literatuur en- rekenaarpakkette

Min standaardhandboeke oor statistiese metodes bevat enige materiaal betreffende effekgroottes. Die uitsondering is Sheskin (2000) se *Handbook of parametric and nonparametric statistical procedures* wat 'n verskeidenheid effekgrootte-indekse bespreek. Verder gee die 4de uitgawe van Tabachnick & Fidell (2001) se *Using multivariate statistics* aandag aan meerveranderlike effekgrootte-indekse en so ook Huberty (1994) se *Applied discriminant analysis*. Omdat so min statistiese handboeke effekgroottes behandel, het dit ook nie inslag gevind by 'n bekende statistiese rekenaarpakket soos SAS (SAS Institute, Inc.2002-2003) nie, terwyl SPSS (SPSS Inc., 2007) en STATISTICA (StatSoft Inc., 2011) se uitvoer beperk is tot die rapportering van eta-kwadraat by ANOVA en MANOVA.

Met die doel om 'n onderskeidingsvermoë en steekproefgrootte te kan bepaal by aanwending van 'n verskeidenheid statistiese betekenistoetse, bespreek Cohen (1969, 1977, 1988) effekgroottes breedvoerig. Omdat sy doel nie die gebruik van effekgrootte is ten einde praktiese betekenisvolheid te bepaal nie, gee hy geen aandag aan beraming van effekgroottes nie. Vir meta-analise is ook effekgroottes nodig en vir daardie doel kan die boeke van Hedges & Olkin (1985) Rosenthal (1991), en Hunter & Schmidt (2004) geraadpleeg word. In die tydskrif *Statistics in Medicine* gee D'Agostino (1999) in 'n redaksionele artikel wel riglyne betreffende praktiese betekenisvolheid (hy noem dit "kwantitatiewe of kliniese betekenisvolheid"). Dit was na aanleiding van die artikel van Feinstein (1999) waar baie breedvoerig na praktiese betekenisvolheid by die vergelyking van twee groepe gekyk is.

Uit bostaande bespreking is dit duidelik dat selfs in die statistiese literatuur effekgroottes as 'n "vreemde" konsep beskou word en dit dus geen wonder is dat dit min of geen inslag gevind het in Toegepaste Statistiek nie. Sover my kennis strek, is effekgroottes en praktiese betekenisvolheid dan ook nooit deel van 'n

inleidende statistiek-kursus nie, hetsy vir eerstejaar Statistiek-studente of as dienskursus. 'n Uitsondering is die kursus STTN124 aan die Noordwes-Universiteit se Potchefstroomkampus. Dis tog jammer dat hierdie belangrike aspek van Statistiek nie bevorder word deur statistici nie, maar dat gebruikers daarvan - veral in die Psigologie en Opvoedkunde - die vlagdraers daarvan moet wees!

1.6 Doelstellings en struktuur van handleiding

- Juis omdat daar in die statistiese literatuur min oor effekgroottes geskryf is, het hierdie handleiding ten doel om die onderwerp vir statistiese konsultante te ontsluit, siende dat baie van die bronne “weggesteek” is in bv. die psigologiese literatuur.
- Vir die navorser wat statistiese metodes as hulpmiddel gebruik in navorsing, wil hierdie handleiding poog om die verskeidenheid effekgroottes wat by verskillende navorsingsprobleme pas, in een bundel saam te voeg.
- Daar word op die agtergrond en interpretasie van effekgroottes gekonsentreer en nie soseer op die statistiese teorie daaragter nie. Die metodes word gegee met genoeg voorbeelde om dit te illustreer.
- Die berekening van meeste effekgrootte-indekse is eenvoudig as aanvaar word dat standaard beskrywende statistieke soos rekenkundige gemiddeldes, standaardafwykings (SA's) en korrelasiekoëffisiënte reeds bereken is (met behulp van rekenaarpakette soos EXCEL, STATISTICA, SAS of SPSS). Die doel is dus om die leser te help om effekgrootte-indekse te bereken uit resultate wat met behulp van die rekenaar verkry is.
- Baie van die vertrouensintervalle van effekgroottes is te ingewikkeld om self te bereken en daarvoor word SAS-programme wat spesiaal ontwikkel is, op die webbladsy van die handleiding geplaas. Die doel van die

handleiding is ook om navorser leiding te gee oor die gebruik van vertrouensintervalle en dit vir hom/haar maklik te maak om te bereken.

Hoofstuk 2 gee 'n literatuuroorsig oor effekgrootte-indekse, nadat eers aandag gegee is aan metingskale in en aannames wat gemaak word by navorsing. Omdat ruimskoots van voorbeelde in empiriese navorsing gebruik gemaak word deur die hele handleiding heen, word in hoofstuk 3 'n verskeidenheid van voorbeelde gegee met die doel om verder aan te gebruik. Hoofstuk 4 bespreek effekgrootte-indekse as gestandaardiseerde verskille tussen twee groepgemiddeldes, terwyl Hoofstuk 5 effekgroottes gee vir 'n verskeidenheid van verbande tussen veranderlikes. Die vergelyking tussen meer as twee groepe het ook effekgrootte-indekse tot gevolg en is die onderwerp van Hoofstuk 6. Hoofstuk 7 brei die gevalle van vergelyking tussen twee of meer groepgemiddeldes uit na die meerveranderlike geval terwyl Hoofstuk 8 oor effekgrootte-indekse vir groep-oorvleueling handel.

Daar bestaan na my wete nog nie effekgrootte-indekse in onder andere gevalle waar getoets word vir normaliteit en homogeniteit van variansies nie en is die onderwerpe vir verdere navorsing.

In meer gespesialiseerde onderwerpe bestaan daar wel effekgrootte-indekse en ek volstaan met 'n lys daarvan met verwysings:

1. Effekgroottes in meerkfaktor ontwerpe – as uitbreiding van Hoofstuk 6: Kline (2004a: Hoofstuk 7), Fidler & Thompson (2001), Olejnik & Algina (2000).
2. Effekgroottes by inter-waarnemer ooreenstemming: Shoukri (2004).