

Analyse van surveydata met SPSS Complex Samples

**Met toepassingen op de SCV-survey
en de survey van de stadsmonitor**

Jan Pickery

Studiedienst van de Vlaamse Regering

Vlaamse overheid



Analyse van surveydata met SPSS Complex Samples

Met toepassingen op de SCV-survey en de survey van de stadsmonitor

Jan Pickery



Samenstelling

Diensten voor het Algemeen Regeringsbeleid
Studiedienst van de Vlaamse Regering (SVR)

Jan Pickery

Leescomité

Geert Loosveldt
Geert Molenberghs
Edwin Pelfrene
Dries Verlet

Verantwoordelijke uitgever

Josée Lemaître
Administrateur-generaal
Boudewijnlaan 30 bus 23
1000 Brussel

Lay-out cover

Diensten voor het Algemeen Regeringsbeleid
Communicatie
Patricia Van Dichel

Druk

Agentschap voor Facilitair Management

Depotnummer

D/2014/3241/114

<http://www.vlaanderen.be/svr>

Inleiding	2
1. Een korte situering van SPSS-Complex Samples.....	2
2. Analyse SCV2013	3
2.1. SCV-steekproef	3
2.2. Complex Samples Plan voor SCV-survey	3
2.3. Voorbeeldanalyses SCV2013	4
2.3.1. Betrouwbaarheidsintervallen	4
2.3.2. Significantietsen	7
3. Analyse Stadsmonitor 2011	11
3.1. Steekproef survey stadsmonitor	11
3.2. Complex Samples Plan voor survey van de stadsmonitor	12
3.3. Voorbeeldanalyses Stadsmonitor 2011	13
3.3.1. Betrouwbaarheidsintervallen	13
3.3.2. Significantietsen	14
3.3.3. Analyse van een deel van het databestand	16
Slotbeschouwingen.....	21
Bibliografie	21
Bijlagen	23
Bijlage 1 Invoeren van het Complex Samples plan voor SCV2013 met SPSS.....	23
Bijlage 2 Opvragen van een Complex Samples-frequentietabel met betrouwbaarheidsintervallen	28
Bijlage 3 Invoeren van het Complex Samples plan voor de survey van de stadsmonitor 2011	32

Inleiding

Statistische formules voor betrouwbaarheidsintervallen en significantietoetsen gaan doorgaans uit van enkelvoudige aselechte steekproeven (*Simple Random Sample*). Zo'n steekproef is ook de assumptie bij de 'default-berekening' in statistische software zoals SPSS. De courante surveypraktijk is echter anders. Het trekken van een steekproef verloopt vaak in verschillende stappen waarbij de populatie opgedeeld wordt in strata, waarna clusters getrokken worden en pas nadien respondenten (uit de clusters). Deze praktijk verschilt duidelijk van enkelvoudige aselechte steekproeven. De meeste surveys voorzien ook gewichten, die (proberen te) corrigeren voor ongelijke selectiekansen, differentiële non-respons en/of een over- of ondervertegenwoordiging van bepaalde categorieën in de steekproef in verhouding tot de populatie. Stratificatie, clustering en ongelijke gewichten hebben een impact op de inferentiële statistische uitspraken (veralgemeningen naar de populatie: bepalen van betrouwbaarheidsintervallen en uitvoeren van statistische toetsen) die gedaan kunnen worden op basis van een steekproef. De formules die gelden bij enkelvoudige aselechte steekproeven zijn dan niet meer geldig.

Er bestaan verschillende methoden om toch op een correcte manier significantietoetsen uit te voeren en/of schattingen te verkrijgen van standaardfouten bij analyses van data die niet het resultaat zijn van enkelvoudige aselechte steekproeven. Rust (1985) geeft een overzicht en een beschrijving van zulke methoden, die reeds langer gekend zijn, maar eerder recentelijk ruimer ingang gevonden hebben. Ruwweg kunnen de methoden ingedeeld worden in linearisatiemethoden ('Taylor series linearisation') en replicatieve methoden (onder andere 'Balanced Repeated Replication' en de 'jackknife method'). SPSS past de 'Taylor series linearisation'-methode toe in de module 'Complex Samples'. Deze tekst illustreert het gebruik van die module bij de analyse van de SCV-survey en de survey voor de Stadsmonitor. Met een aantal eenvoudige voorbeelden wordt stilgestaan bij de verschillende resultaten van 'default'-analyses en Complex Samples-analyses. Met 'default' bedoelen we eigenlijk elke andere analyse dan Complex Samples die menugewijs wordt uitgevoerd na eerst gewichten te hebben gedefinieerd met 'Weight cases'. Dat zijn analyses die de complexiteit van het steekproefdesign niet in rekening brengen. De tekst is opgevat als een gebruikershandleiding. Het technisch-statistische niveau is beperkt. Een aantal tabellen worden in de oorspronkelijke SPSS-layout weergegeven om de resultaten sneller zelf terug te kunnen vinden. Bovendien bevatten de bijlagen screenshots van (voorbereidingen tot) analyses in SPSS, waardoor het makkelijk is om een aantal analyses zelf te repliceren.

1. Een korte situering van SPSS-Complex Samples

SPSS Complex Samples maakt het enerzijds mogelijk om steekproeven te trekken (met stratificatie en clustering) en anderzijds om bij de analyse van surveydata rekening te houden met de kenmerken van de steekproef.

Om een *steekproef* te kunnen *trekken* met SPSS Complex Samples moet je een lijst hebben van de volledige populatie, het steekproefkader. Je kan voor de steekproef clusteren en stratificeren en in verschillende stappen werken. Eigenlijk kan dit allemaal ook met syntax van het SPSS-basispakket. Met Complex Samples verloopt het misschien iets vlotter en kan je ineens gewichten berekenen en het steekproefplan ook omzetten in een analyseplan. Een *analyse van surveydata* die rekening houdt met de afwijkingen van de enkelvoudige aselechte steekproef gebeurt immers via een zogenaamde 'plan-file'. Eens gemaakt kan dat bestand voor elke analyse opgeroepen worden en het kan ook ter beschikking gesteld worden van andere onderzoekers. Dat plan bevat de steekproefinformatie die nodig is om de geschatte standaardfouten en varianties (en dus ook betrouwbaarheidsintervallen) te berekenen rond statistische parameters. Zoals de twee voorbeelden zullen verduidelijken, is dat doorgaans informatie over de gebruikte clusters, strata en gewichten.

De beschikbare analyses nemen bij elke nieuwe versie van SPSS toe. In SPSS 20 is het bijvoorbeeld al mogelijk om Complex Samples logistische regressie en Complex Samples Cox regressie uit te voeren. Meer informatie over Complex Samples valt te lezen in IBM (2011).

2. Analyse SCV2013

2.1. SCV-steekproef

Omdat het analyseplan rekening moet houden met de kenmerken van het steekproefdesign, staan we nog even stil bij de voornaamste kenmerken van dat design van de SCV-survey. De populatie omvat de inwoners van het Vlaamse en Brusselse Hoofdstedelijke Gewest die ouder zijn dan 18 jaar en de Nederlandse taal machtig zijn. Zij moeten geen Belg zijn en ook niet Nederlandstalig zijn, maar het interview moet wel kunnen plaatsvinden in het Nederlands. Om het aantal Nederlandsonkundige geselecteerde personen in het Brusselse Hoofdstedelijke Gewest (enigszins) te beperken, wordt in Brussel getrokken op Nederlandstalige adressen. Het Rijksregister van natuurlijke personen vormt het steekproefkader.

De steekproef kan omschreven worden als een gestratificeerde tweetrapssteekproef met clustering op het niveau van de postsectoren. Concreet wordt de populatie opgedeeld in 6 strata (de 5 Vlaamse provincies en Brussel). In eerste instantie worden dan binnen die 6 strata clusters getrokken, gelokaliseerd in postsectoren (postcodes). Hoeveel clusters getrokken worden in een postsector is afhankelijk van het toeval waarbij de kans proportioneel is aan de bevolkingsomvang. In één postsector kunnen dus ook meerdere clusters getrokken worden. Voor elke cluster wordt gemikt op 10 interviews. Om non-respons te ondervangen is het aantal ter beschikking gestelde adressen natuurlijk hoger. Dat aantal varieert tussen postsectoren in functie van de responscijfers bij vorige SCV-surveys. Om de ongelijke selectiekansen die hieruit voortvloeien te compenseren, worden gewichten berekend. Die gewichten worden daarna nog aangepast op basis van een non-responsanalyse en om bepaalde steekproefverdelingen te conformeren aan gekende populatieverdelingen. Uiteindelijk worden er twee gewichten ter beschikking gesteld. Het eerste telt op tot de populatieomvang. Het tweede, herschaalde, gewicht heeft een gemiddelde gelijk aan 1, zodanig dat gewogen en ongewogen steekproefomvang gelijk zijn. De details van deze steekproeftrekking en van de berekening van de gewichten worden steeds uitgebreid gedocumenteerd in de basisdocumentaties van de survey (zie bijvoorbeeld Carton, e.a., 2014). Het is duidelijk dat er bij de schatting van de gewichten een mate van onzekerheid speelt. Die gewichten mogen daarom bij de berekening van bijvoorbeeld standaardfouten niet als constanten beschouwd worden (Kalton & Flores-Cervantes, 2003). Het is dat wat een default gewogen analyse in SPSS juist wel doet.

2.2. Complex Samples Plan voor SCV-survey

Van het steekproefdesign van de SCV-survey onthouden we de verschillende stappen, de stratificatie en de clustering en de gewichten. De steekproeftheorie en de ervaring leert echter dat bij zo'n steekproefdesign de grootste bijdrage in de variantie komt van de eerste stap van de steekproeftrekking, het trekken van de clusters, en veel minder van de volgende stap, het trekken van personen binnen de clusters. De meest gebruikte methoden van variantieschattingen brengen daarom niet alle stappen van het steekproefdesign in rekening maar alleen de eerste. Het opnemen van strata, clusters en gewichten en de bepaling dat de clusters in die eerste stap getrokken werden met teruglegging (zelfs als dat laatste niet het geval is), zorgt bij de berekening van de varianties voor een goede benadering (zie bijvoorbeeld Brogan, 2005, p. 451).

Dat maakt het analyseplan voor de SCV-survey eigenlijk zeer eenvoudig. We hebben slechts drie variabelen nodig: een stratum- en een clusteridentificatie en het gewicht. Het ingeven van die variabelen in het Complex Samples Plan gebeurt menugestuurd en wijst zichzelf uit (zie screenshots in bijlage 1). In een tweede stap moeten we alleen nog de steekproeftrekking met teruglegging (*With Replacement*) aanvinken. Bij SCV is dat ook correct. Postsectoren kunnen meerdere malen getrokken worden. Maar zelfs als dat niet het geval was, krijgen we met deze specificatie dus een goede benadering van de berekening van de varianties en standaardfouten. Tot slot kan de eindigheidscorrectie (Finite Population Correction) nog aangevinkt of uitgevinkt worden. Aanvinken mag in dit geval alleen als de gewichten optellen tot de populatieomvang. Het mag dus wel met het oorspronkelijke gewicht, maar niet met het herschaalde gewicht. Gegeven de verhouding steekproef – populatie (± 1.500 respondenten op

$\pm 5.200.000$ Vlamingen die tot de populatie behoren) is een eindigheidscorrectie in dit geval eigenlijk zinloos, maar als we het niet-herschaalde gewicht gebruiken, mogen we dit vakje wel aangevinkt laten.

2.3. Voorbeeldanalyses SCV2013

2.3.1. Betrouwbaarheidsintervallen

We tonen twee verschillende analyses. In de eerste analyse schatten we een percentage met het bijhorende betrouwbaarheidsinterval. In de tweede analyse toetsen we een verschil in percentages. Voor de eerste analyse kijken we naar de verhuisintenties van de Vlamingen.

In SCV2013 werd de eenvoudige ja/nee-vraag gesteld 'Bent u van plan om in de komende vijf jaar te verhuizen?'. 'Weet niet' en 'geen antwoord' waren beschikbaar voor de interviewer op de draagbare computer, maar werden niet als antwoordalternatieven aangeboden aan de respondenten. Die antwoorden (39 in totaal) laten we buiten beschouwing zodat we 1.476 geldige antwoorden hebben. De ongewogen verdeling daarvan bevindt zich in tabel 1.

Tabel 1 Verhuisintentie – ongewogen

	Aantal	Percentage
0 – niet	1.149	77,85
1 – wel	327	22,15
Totaal	1.476	100,00

SPSS heeft geen eenvoudige manier om betrouwbaarheidsintervallen op te vragen voor categorische variabelen. Een omweg is mogelijk door de dichotome categorische variabele te hercoderen tot een 0/1-variabele, die variabele als metrisch te beschouwen en er 'Descriptive Statistics' voor op te vragen via 'Explore'. Zo wordt er eigenlijk een betrouwbaarheidsinterval berekend rond een gemiddelde.

We kunnen ook zelf een betrouwbaarheidsinterval berekenen, bijvoorbeeld met de formule die vertrekt van een benadering gebaseerd op de geëigende t-verdeling of, voor grotere steekproeven, op de standaardnormale verdeling (z-verdeling):

$$\hat{p} \pm t_{(v, 1-\frac{\alpha}{2})} \sqrt{\frac{\hat{p}(1-\hat{p})}{n}}$$

waarbij

- $\hat{p}$ = steekproefschatter voor de populatieproportie π
= p , de waargenomen steekproefproportie
- t = waarde uit de t-verdeling, die overeenkomt met het gewenste betrouwbaarheidsniveau; bij grotere steekproeven vrijwel gelijk aan de z-waarde uit de standaardnormaalverdeling
- v = aantal vrijheidsgraden -> voor selectie van de voor de steekproef geëigende t-verdeling in de familie van t-verdelingen
- α = vooropgestelde onbetrouwbaarheidsdrempel
- n = steekproefomvang

Als we deze formule toepassen, krijgen we als ondergrens van het 95%-betrouwbaarheidsinterval 20,03% en als bovengrens 24,27%. (Als we ons bij een steekproef van 1.476 eenheden 5% kans geven om ons te vergissen ($\alpha = 0,05$), is t gelijk aan 1,962.)

Als we de dummy verhuisintentie beschouwen als een metrische variabele en daarop de Explore-procedure loslaten, krijgen we een vrijwel identiek betrouwbaarheidsinterval met

eveneens 20,03% als ondergrens en 24,28% als bovengrens (zie tabel 2). De metrische omweg levert dus bijna exact dezelfde resultaten op als de categorische berekening.

Tabel 2 Output van de Explore-procedure in SPSS voor verhuisintentie – ongewogen

Descriptives			Statistic	Std. Error
verhuisintentie	Mean		,2215	,01081
	95% Confidence Interval for Mean	Lower Bound	,2003	
		Upper Bound	,2428	
	5% Trimmed Mean		,1906	
	Median		,0000	
	Variance		,173	
	Std. Deviation		,41543	
	Minimum		,00	
	Maximum		1,00	
	Range		1,00	
	Interquartile Range		,00	
	Skewness		1,342	,064
	Kurtosis		-,198	,127

Zoals hierboven al gesteld, werkt de SCV-survey echter met gewichten. Er zit variabiliteit in die gewichten: niet elke respondent heeft hetzelfde gewicht. Het gebruik van die gewichten leidt zo tot andere parameterschattingen. Omdat er in die gewichten een mate van onzekerheid zit, mogen ze niet als constanten beschouwd worden. Dat heeft ook een impact op betrouwbaarheidsintervallen. Om de effecten na te gaan, gebruiken we die gewichten in eerste instantie op de defaultwijze in SPSS (via menu: Data → Weight cases).

Het gewogen percentage verschilt een klein beetje van het ongewogen percentage: 21,68% plant een verhuis binnen de 5 jaar (tegenover 22,15% ongewogen). Voor de schatting van dat percentage maakt het niet uit of we het oorspronkelijke gewicht (tabel 3) gebruiken of het herschaalde gewicht (tabel 4).

Tabel 3 Verhuisintentie – gewogen met niet-herschaalde gewicht (SCV_gewicht13)

		Aantal	Percentage
0	- niet	3.980.822	78,32
1	- wel	1.102.159	21,68
	Totaal	5.082.980	100,00

Tabel 4 Verhuisintentie – gewogen met herschaalde gewicht (SCV_gewicht_herschaald13)

		Aantal	Percentage
0	- niet	1.153	78,32
1	- wel	319	21,68
	Totaal	1.473	100,00

Voor het 95%-betrouwbaarheidsinterval zijn er natuurlijk wel verschillen. Als we werken met het niet-herschaalde gewicht is dat betrouwbaarheidsinterval onrealistisch klein (21,65% – 21,72%). Als we het herschaalde gewicht gebruiken, krijgen we een interval 19,58% – 23,79%. Een beetje verderop in deze tekst zullen we alle berekende betrouwbaarheidsintervallen nog eens oplijsten, maar uit de eerste vergelijking hier blijkt dat de ongewogen resultaten en de resultaten gewogen met het herschaalde gewicht vergelijkbaar zijn: het

geschatte percentage is wat lager, maar het interval is ongeveer even breed. Default-gebruik van het gewone, niet-herschaalde, gewicht levert natuurlijk een onrealistisch klein en fout betrouwbaarheidsinterval op.

We voeren ook enkele Complex Samples-analyses uit. In eerste instantie gebruiken we daarvoor het volledige plan, dat is opgesteld volgens de bepalingen in sectie 2.2. Dat Complex Samples Plan houdt dus rekening met de stratificatie, de clustering en de gewichten én met de impact van die elementen op de geschatte standaardfouten. Met Complex Samples kunnen wel eenvoudig betrouwbaarheidsintervallen opgevraagd worden voor categorische variabelen. Dat gebeurt allemaal relatief eenvoudig menugestuurd (bijlage 2 toont de screenshots). De resultaten bevinden zich in tabel 5. Het geschatte percentage dat een verhuis plant binnen de 5 jaar is hetzelfde als het default gewogen percentage: 21,68%. Maar het betrouwbaarheidsinterval is nu wel wat groter. Vooral de bovengrens is verhoogd tot 24,02%, zodat het interval nu ook niet meer perfect symmetrisch is.

Tabel 5 Output van de frequentieanalyse 'verhuisintentie' met SPSS Complex Samples

verhuisintentie						
		Estimate	Standard Error	95% Confidence Interval		Unweighted Count
				Lower	Upper	
Population Size	,00	3980821,62	113605,73	3756271,37	4205371,87	1149
	1,00	1102158,71	69683,72	964423,61	1239893,82	327
	Total	5082980,33	137991,41	4810229,96	5355730,71	1476
% of Total	,00	78,32%	1,14%	75,98%	80,49%	1149
	1,00	21,68%	1,14%	19,51%	24,02%	327
	Total	100.0000%	0,00%	100,00%	100,00%	1476

Om de oorzaak van het verschil te duiden, hebben we ook een aantal andere analyseplannen gecreëerd in Complex Samples. We hebben de eindigheidscorrectie eens laten vallen en het gebruik van clustering en/of stratificatie niet opgenomen in het plan. We veranderen dus de mate waarin de analyse rekening houdt met de verschillende elementen van het design van de survey, niet het design zelf natuurlijk. Het maakt allemaal weinig tot geen verschil. Het is duidelijk dat het voornamelijk de gewichten zijn die ervoor zorgen dat het betrouwbaarheidsinterval groter wordt.

In tabel 6 lijsten we alle verschillende analyses met bijhorende betrouwbaarheidsintervallen op. De verschillen tussen de Complex Samples-analyses onderaan zijn minimaal. Dat maakt duidelijk dat zowel de negatieve impact van het clusteren als de positieve impact van het stratificeren te verwaarlozen zijn. Het zijn vooral de gewichten die maken dat de correct geschatte betrouwbaarheidsintervallen groter zijn dan deze die het resultaat zijn van de default-analyses in SPSS. De getoonde verschillen ogen misschien niet zo spectaculair. Maar de breedte van het interval is toch gestegen van 4,21 procentpunten naar 4,51 procentpunten. Als we bij de analyse rekening houden met het design van de steekproef en (vooral) met de gewichten, is (in dit geval) het betrouwbaarheidsinterval 7% groter dan als we dat niet doen.

Tabel 6 Overzicht van de verschillende analyses met bijhorende betrouwbaarheidsintervallen

Analyse	Parameter- schatting	95%-betrouwbaarheidsinterval		
		Ondergrens	Bovengrens	Breedte
Default SPSS-analyses				
ongewogen, zelf berekend met formule	22,15	20,03	24,27	4,24
ongewogen, metrisch (explore)	22,15	20,03	24,28	4,25
gewogen, niet-herschaald gewicht, zelf berekend	21,68	21,65	21,72	0,07
gewogen, niet-herschaald gewicht, metrisch (explore)	21,68	21,65	21,72	0,07
gewogen, herschaald gewicht, zelf berekend	21,68	19,58	23,79	4,21
gewogen, herschaald gewicht, metrisch (explore)	21,68	19,58	23,79	4,21
Complex Samples-analyses				
volledige plan	21,68	19,51	24,02	4,51
plan dat geen rekening houdt met eindigheidscorrectie	21,68	19,51	24,02	4,51
plan dat geen rekening houdt met clustering	21,68	19,52	24,01	4,49
plan dat geen rekening houdt met stratificatie	21,68	19,50	24,03	4,53
plan dat geen rekening houdt met clustering of stratificatie (alleen gewichten)	21,68	19,52	24,01	4,49

2.3.2. Significantietoetsen

In een volgende analyse gaan we na wat de impact is van het al dan niet rekening houden met de steekproef en de gewichten op een toets voor verschillen. We analyseren de deelname aan erediensten volgens opleidingsniveau. Vraag 31 van SCV2013 luidde: 'Mensen nemen wel eens deel aan kerkelijke of religieuze plechtigheden naar aanleiding van een huwelijk, begrafenis en dergelijke. Als we deze NIET meetellen, hoe vaak neemt u dan deel aan kerkelijke of godsdienstige erediensten?'. De ongewogen antwoordverdeling op deze vraag wordt getoond in tabel 7. Bijna de helft van de respondenten neemt nooit deel aan erediensten.

We hebben deze vraag ook gedichotomiseerd: 'maandelijks of vaker versus niet regelmatig' en laten verder ook de item non-respons buiten beschouwing (zie tabel 8). Van de 1.514 respondenten die de vraag beantwoordden, nemen er 194 regelmatig deel aan erediensten.

Tabel 7 Deelname aan erediensten – ongewogen

	Aantal	Percentage
1 – nooit	753	49,7
2 – zeer zelden	427	28,2
3 – enkel op religieuze of godsdienstige feestdagen	140	9,2
4 – maandelijks	71	4,7
5 – meerdere keren per maand	38	2,5
6 – wekelijks	75	5,0
7 – meerdere keren per week	10	0,7
88 – geen antwoord	1	0,1
Totaal	1.515	100,0

Tabel 8 Deelname aan erediensten, dichotoom – ongewogen

	Aantal	Percentage
0 – niet regelmatig	1.320	87,2
1 – maandelijks of vaker	194	12,8
Totaal	1.514	100,0

In onze analyse werken we verder met deze dichotome variabele. We gaan na in welke mate de deelname aan erediensten verschilt volgens opleidingsniveau. Tabel 9 toont een eenvoudige kruistabel (ongewogen) met het hoogst behaalde diploma.

Tabel 9 Deelname aan erediensten volgens opleidingsniveau – ongewogen

		Deelname aan erediensten		
Opleidingsniveau		Niet regelmatig	Maandelijks of vaker	Totaal
Geen/lager onderwijs	Aantal	90	25	115
	rijpercentage	78,3%	21,7%	
Lager secundair onderwijs	Aantal	265	42	307
	rijpercentage	86,3%	13,7%	
Hoger secundair onderwijs	Aantal	499	73	572
	rijpercentage	87,2%	12,8%	
Niet-universitair hoger onderwijs	Aantal	303	44	347
	rijpercentage	87,3%	12,7%	
Universitair onderwijs	Aantal	163	10	173
	rijpercentage	94,2%	5,8%	
Totaal	Aantal	1.320	194	1.514
	rijpercentage	87,2%	12,8%	

$$\chi^2=16,074, p = 0,00292$$

Uit tabel 9 blijkt duidelijk de samenhang tussen opleidingsniveau en deelname aan erediensten. Voor het laagste opleidingsniveau bedraagt het aantal regelmatige deelnemers 21,7%; voor het hoogste opleidingsniveau amper 5,8%. De middelste drie niveaus situeren zich daartussen, dicht bij het algemene gemiddelde van 12,8%. De berekende Chi-kwadraatwaarde voor deze tabel is gelijk aan 16,074; met 4 vrijheidsgraden is de bijhorende kans gelijk aan 0,00292. De kans op een even grote chi-kwadraatwaarde als de nulhypothese juist is, oftewel de kans om dergelijke verschillen te vinden in een steekproef van deze omvang als er in de populatie geen samenhang zou zijn tussen beide kenmerken (onafhankelijkheid) is dus kleiner dan 3 op 1.000. Omdat die kans zo klein is, verwerpen we de nulhypothese van onafhankelijkheid en kunnen we besluiten dat er wel een samenhang is (en dus een verschil in deelname aan erediensten volgens opleidingsniveau).

We voeren dezelfde analyse uit, maar nu gewogen met het herschaalde gewicht. De univariate frequentietabel bevindt zich in tabel 10. Tabel 11 kruist de deelname aan erediensten met het opleidingsniveau.

Tabel 10 Deelname aan erediensten, dichotoom – gewogen met herschaalde gewicht

	Aantal*	Percentage
0 – niet regelmatig	1.296	85,6
1 – maandelijks of vaker	219	14,4
Totaal	1.514	100,0

* Bemerk dat de aantallen niet exact optellen tot het totaal. Dat komt door de afronding van de aantallen. Eventueel zouden ook aantallen met decimalen weergegeven kunnen worden. Door het gebruik van de gewichten zijn het immers geen reële aantallen respondenten.

Tabel 11 Deelname aan erediensten volgens opleidingsniveau – gewogen met herschaalde gewicht

		Deelname aan erediensten		
Opleidingsniveau		Niet regelmatig	Maandelijks of vaker	Totaal
Geen/lager onderwijs	Aantal	209	63	272
	rijpercentage	76,8%	23,2%	
Lager secundair onderwijs	Aantal	218	34	252
	rijpercentage	86,5%	13,5%	
Hoger secundair onderwijs	Aantal	474	72	546
	rijpercentage	86,8%	13,2%	
Niet-universitair hoger onderwijs	Aantal	255	40	295
	rijpercentage	86,4%	13,6%	
Universitair onderwijs	Aantal	140	9	149
	rijpercentage	94,0%	6,0%	
Totaal	Aantal	1.296	218	1.514
	rijpercentage	85,6%	14,4%	

$\chi^2=26,379$, $p = 0,00003$

Bemerk dat het gewogen aantal vooral voor de laagstgeschoolden sterk afwijkt van het ongewogen aantal. De gewichten zijn mede het resultaat van de ondervertegenwoordiging van lagergeschoolden. Gewogen vinden we zo ook een hoger percentage regelmatige deelnemers aan erediensten. Voor de populatie wordt dat geschat op 14,4%; 1,5 procentpunten hoger dan het ongewogen percentage bij de respondenten.

De conditionele percentages in de gewogen kruistabel verschillen eigenlijk niet zo heel veel van deze in de ongewogen tabel. Bij het laagste opleidingsniveau vinden we een regelmatige deelname van 23,2% (tegenover 21,7% in tabel 9); bij het hoogste opleidingsniveau is dat 6% (tegenover 5,8%). De drie middelste niveaus situeren zich daartussen en zijn vrijwel gelijk aan elkaar. Ook hier besluiten we tot significante samenhang tussen opleidingsniveau en deelname. De kans dat die samenhang te wijten is aan toeval is zelfs nog veel kleiner: 3 op 100.000. De reden voor dit verschil is betrekkelijk eenvoudig. De groep waarvoor de conditionele verdeling het meest afwijkt van de marginale verdeling, de laagstgeschoolden, is gewogen veel groter dan ongewogen. Zo lijkt die afwijking gewogen robuuster en dus signifikanter dan ongewogen. Dat is op zich geen resultaat van de SCV-survey, maar wel van de toepassing van de gewichten. Vanuit deze overweging is het waarschijnlijk dat de ongewogen toets beter aansluit bij de waarheid dan de gewogen toets.

Een default analyse gewogen met het oorspronkelijke gewicht levert hier weinig bijkomende inzichten op. De percentages zijn dezelfde als in tabellen 10 en 11, maar de aantallen zijn veel groter, en de p -waarde bijgevolg nog veel kleiner ($p < 1E-16$). Daarom tonen we in tabel 12

ineens de Complex Samples-analyse. Het uitvoeren van zo'n analyses verloopt ook eenvoudig menugestuurd.

Tabel 12 Output van de kruistabel met SPSS Complex Samples

opleid13_eak * erediensten2			erediensten2		Total
opleid13_eak			,00	1,00	
1) geen/lo	Population Size	Estimate	720225,734	217585,648	937811,382
	% within opleid13_eak	Estimate	76,8%	23,2%	100,0%
		Standard Error	4,6%	4,6%	0,0%
	95% Confidence Interval	Lower	66,6%	15,4%	100,0%
		Upper	84,6%	33,4%	100,0%
	Unweighted Count		90	25	115
2) lager sec	Population Size	Estimate	754022,120	117691,497	871713,616
	% within opleid13_eak	Estimate	86,5%	13,5%	100,0%
		Standard Error	2,0%	2,0%	0,0%
	95% Confidence Interval	Lower	82,0%	10,0%	100,0%
		Upper	90,0%	18,0%	100,0%
	Unweighted Count		265	42	307
3) hoger sec	Population Size	Estimate	1636954,177	249646,704	1886600,882
	% within opleid13_eak	Estimate	86,8%	13,2%	100,0%
		Standard Error	1,5%	1,5%	0,0%
	95% Confidence Interval	Lower	83,5%	10,6%	100,0%
		Upper	89,4%	16,5%	100,0%
	Unweighted Count		499	73	572
4) nuho	Population Size	Estimate	878561,370	138630,092	1017191,462
	% within opleid13_eak	Estimate	86,4%	13,6%	100,0%
		Standard Error	2,0%	2,0%	0,0%
	95% Confidence Interval	Lower	81,9%	10,1%	100,0%
		Upper	89,9%	18,1%	100,0%
	Unweighted Count		303	44	347
5) unief	Population Size	Estimate	481638,532	31325,801	512964,332
	% within opleid13_eak	Estimate	93,9%	6,1%	100,0%
		Standard Error	1,9%	1,9%	0,0%
	95% Confidence Interval	Lower	88,9%	3,3%	100,0%
		Upper	96,7%	11,1%	100,0%
	Unweighted Count		163	10	173
Total	Population Size	Estimate	4471401,932	754879,741	5226281,673
	% within opleid13_eak	Estimate	85,6%	14,4%	100,0%
		Standard Error	1,1%	1,1%	0,0%
	95% Confidence Interval	Lower	83,2%	12,3%	100,0%
		Upper	87,7%	16,8%	100,0%
	Unweighted Count		1320	194	1514

De output in tabel 12 toont dezelfde percentages als de default gewogen analyse in tabel 11, zie bijvoorbeeld de 23,2% regelmatige deelname bij het laagste opleidingsniveau. Verder zijn er schattingen van de aantallen in de populatie en kunnen er eenvoudig betrouwbaarheidsintervallen opgevraagd worden per categorie, zowel voor de aantallen als voor de percentages. Wij tonen alleen de intervallen voor de percentages. In principe kunnen deze betrouwbaarheidsintervallen ook berekend worden voor de percentages in tabel 11. Maar omdat Complex Samples een correctere schatting geeft van de standaardfout, zijn ook de betrouwbaarheidsintervallen in tabel 12 correcter. Zo blijkt dat het 95%-betrouwbaarheids-interval voor de laagstgeschoolden liefst 3 keer zo groot is als dat voor de mensen met ten hoogste een diploma hoger secundair (18 procentpunten tegenover 6 procentpunten). Dit hangt natuurlijk samen met het werkelijke aantal respondenten in die categorie, en dat is het ongewogen aantal. Zo combineert SPSS Complex Samples de gewogen percentages met de ongewogen aantallen. De procedure levert ook een significantietoets. Dat is geen gewone chi-

kwadraattoets, maar een aangepaste F-toets (zie tabel 13). Die toont dat de kans om in een steekproef als de onze een samenhang te vinden die even sterk is of nog sterker als er in de populatie geen samenhang is, gelijk is aan 0,0023 (of 0,0024 afhankelijk van de toets – verschil is doorgaans verwaarloosbaar). Met een kans gelijk aan ± 2 op 1.000 besluiten we ook hier tot een significante samenhang tussen scholingsniveau en deelname aan erediensten. Deze p -waarde bevindt zich inderdaad dicht bij die van de ongewogen toets (tabel 9) dan bij die van de gewogen toets (tabel 11).

Tabel 13 Output van de significantietoets van SPSS Complex Samples bij de kruistabel

Tests of Independence						
		Chi-Square	Adjusted F	df1	df2	Sig.
opleid13_eak *	Pearson	26,208	4,830	3,099	446,281	,0022712775
erediensten2	Likelihood Ratio	26,020	4,796	3,099	446,281	,0023839728

Net zoals bij de betrouwbaarheidsintervallen toont dit voorbeeld dat het gebruik van SPSS Complex Samples geen overbodige luxe is. Hier is de inhoudelijke conclusie steeds dat er een samenhang is tussen opleidingsniveau en het al dan niet regelmatig deelnemen aan erediensten. Maar de zekerheid waarmee we die conclusie kunnen trekken, verschilt sterk. Bij toetsen waarbij de p -waarde cirkelt rond de gangbare grenzen van 0,05 of 0,01 kan het al dan niet rekening houden het steekproefdesign en de gewichten wel degelijk tot andere besluiten leiden. De informatie die nodig is om zo'n analyse te doen is uiteindelijk zeer beperkt: een stratum- en clusteridentificatie en een gewicht. Die variabelen zijn opgenomen in het bestand zodat ook niet-SPSS-gebruikers er in hun analyse rekening mee kunnen houden. Ook andere softwarepakketten zoals Stata of SAS bieden die mogelijkheden (en soms zelfs nog meer).

Bij de SCV-survey is het belangrijkste element om rekening mee te houden veruit het gewicht. De stratificatie en clustering hebben amper een impact op de precisie. Waarschijnlijk is de provincie als stratificatiecriterium niet meer zo relevant. Omdat wij veel, kleine clusters gebruiken, is de negatieve impact van het clusteren ook beperkt.

3. Analyse Stadsmonitor 2011

3.1. Steekproef survey stadsmonitor

Ook voor deze survey staan we even stil bij de steekproefprocedure. De survey voor de stadsmonitor betreft een twee- à driejaarlijkse schriftelijke enquête bij de inwoners ouder dan 15 jaar van de 13 centrumsteden (Aalst, Antwerpen, Brugge, Genk, Gent, Hasselt, Kortrijk, Leuven, Mechelen, Oostende, Roeselare, Sint-Niklaas en Turnhout). De steekproefgrootte varieert van stad tot stad. Dat is gedeeltelijk historisch gegroeid. De grotere steden 'kregen wat extra respondenten', maar nu ook weer niet zoveel dat er van een proportionele verdeling kan gesproken worden. Als vergelijken van de steden de belangrijkste doelstelling zou zijn, zou een gelijke steekproefgrootte overigens aangewezen zijn omdat de steekproefomvang (aantal respondenten) voor de precisie veel belangrijker is dan de steekproeffractie (aandeel van de respondenten in verhouding tot de populatieomvang). Verder zijn er steden die zelf meer respondenten willen om wijken/districten/stadsdelen te kunnen vergelijken en daar ook zelf een financiële bijdrage voor over hebben. Voor die steden is vergelijken dus een expliciete doelstelling, wat dan weer wel resulteerde in een min of meer gelijke steekproefomvang per wijk/district/stadsdeel. Een uitzondering vormen de wijken met zeer weinig inwoners. Het aantal respondenten kan daar wat kleiner zijn omdat een eindigheidscorrectie er wel iets uitmaakt. In 2011 kozen Aalst, Antwerpen, Genk en Turnhout ervoor om meer respondenten te hebben om intrastedelijke vergelijkingen te kunnen maken. Turnhout, dat in de voorbeeldanalyses aan bod zal komen, onderscheidde bijvoorbeeld 7 stadsdelen. Vooral 'Stadsbos en Noorden' is zeer klein met minder dan 1.000 inwoners in de onderzochte doelpopulatie (15+). Het aantal geselecteerde inwoners van dat stadsdeel was bijgevolg wat lager dan bij de andere stadsdelen omdat hier een eindigheidscorrectie mogelijk en aangewezen is.

De steekproeven werden via impliciete stratificatie getrokken uit de bevolkingsregisters van de centrumsteden. De bevolkingslijsten werden geordend volgens nationaliteit, geslacht en leeftijd waarna met een vaste sprong het benodigde aantal inwoners werd geselecteerd. Op die manier worden de populatieverdelingen volgens nationaliteit, geslacht en leeftijd alvast in de getrokken steekproef volledig weerspiegeld. Binnen elke stad (of voor de steden met extra respondenten binnen elk stadsdeel) heeft iedere inwoner ook dezelfde kans om getrokken te worden.

Na het binnenlopen van de laatste vragenlijsten, werden er gewichten berekend voor alle respondenten. Die gewichten houden rekening met de populatieomvang van de stad of het stadsdeel en met de verdelingen volgens leeftijd en geslacht in elke stad of elk stadsdeel. Er werden alleen gewichten ter beschikking gesteld die sommeren tot het populatietotaal. Voor deze tekst worden nog een aantal herschaalde gewichten berekend, maar die werden in eerste instantie niet bezorgd aan de onderzoekers die de survey analyseren. Meer methodologische informatie over de survey van de stadsmonitor kan gevonden worden in Schelfaut (2009).

3.2. Complex Samples Plan voor survey van de stadsmonitor

Van dit steekproefdesign onthouden we vooral de disproportionele stratificatie. Het aantal respondenten in de verschillende steden is niet in verhouding tot het aandeel van die stad in het totaal van alle inwoners van de 13 centrumsteden. Bij de steden die zelf extra interviews bekostigden, is er ook zo'n onevenredige vertegenwoordiging per wijk/district/stadsdeel. De gewichten corrigeren hiervoor en ook voor de afwijkingen ten opzichte van de verdeling volgens leeftijd en geslacht in elke stad of wijk/district/stadsdeel. Die gewichten moeten dus opgenomen worden in het Complex Samples Plan. In sommige wijken is het totale aantal inwoners relatief klein, zodanig dat een eindigheidscorrectie zinvol kan zijn. Daarom zullen we in dit plan ook de bevolkingstotalen opnemen.

Dat maakt ook dit analyseplan eigenlijk redelijk eenvoudig. We hebben opnieuw drie variabelen nodig: een stratumidentificatie, een populatietotaal voor elk stratum en het gewicht. Het ingeven van de stratumidentificatie en het gewicht in het Complex Samples Plan gebeurt opnieuw menugestuurd (zie screenshots in bijlage 3). In een volgend venster kiezen we nu 'Equal WOR': Equal probability sampling WithOut Replacement. Binnen elk stratum heeft ieder individu gelijke selectiekansen en de selectie gebeurt ook effectief zonder teruglegging. Door de systematische steekproeftrekking kan een inwoner niet tweemaal geselecteerd worden. In het daaropvolgende venster geven we aan dat de populatietotalen gehaald kunnen worden uit de gelijknamige variabele.

Eén element van het steekproefdesign wordt zo niet opgenomen in het plan: de impliciete stratificatie die volgt uit de systematische steekproeftrekking. Dat is niet ongevoel. De Taylor Series Linearisatiemethode die SPSS Complex Samples toepast, biedt niet de mogelijkheid om impliciete stratificatie als gevolg van systematische steekproeftrekking als designfactor op te nemen. Bovendien wordt de fout die aldus gemaakt wordt, doorgaans getolereerd omdat ze in ons nadeel speelt (zie bijvoorbeeld Goedemé, 2010). De impliciete stratificatie zou de precisie wat moeten verhogen. Als we er geen rekening mee houden, zijn de berekende betrouwbaarheidsintervallen wat te groot, net zoals de p -waarden bij significantietoetsen. Zo'n fout is aanvaardbaarder dan een fout in de omgekeerde richting.

3.3. Voorbeeldanalyses Stadsmonitor 2011

3.3.1. Betrouwbaarheidsintervallen

Ook voor de survey van de Stadsmonitor 2011 analyseren we een tot twee categorieën herleide categorische variabele. We bekijken de tevredenheid met de stad. Aan alle respondenten werd de eenvoudige vraag gesteld 'In welke mate ben je tevreden over de stad waar je woont?', die beantwoord kon worden met vijf gradaties van tevredenheid. Tabel 14 toont de ongewogen verdeling van die variabele.

Tabel 14 Tevredenheid met de stad – ongewogen

	Aantal	Percentage
Zeer ontevreden	511	2,7
Eerder ontevreden	1.661	8,8
Noch tevreden, noch ontevreden	3.074	16,3
Eerder tevreden	9.631	51,2
Zeer tevreden	3.832	20,4
Geen antwoord	119	0,6
Totaal	18.828	100,0

Deze variabele wordt dus gedichotomiseerd: eerder tevreden en zeer tevreden versus de andere antwoorden. De mensen die de vraag onbeantwoord laten, worden buiten beschouwing gelaten. Zo vinden we dat net geen 72% van de respondenten tevreden is met de stad waarin ze wonen (zie tabel 15).

Tabel 15 Tevredenheid met de stad, dichotoom – ongewogen

	Aantal	Percentage
0 – niet	5.246	28,04
1 – wel	13.463	71,96
Totaal	18.709	100,00

De gewogen verdeling ziet er toch behoorlijk anders uit. Het geschatte percentage voor de populatie (alle inwoners van de 13 centrumsteden) ligt een stuk hoger dan het ongewogen percentage van alle respondenten (zie tabel 16). De voornaamste verklaring, die uit de analyse achteraf blijkt, is dat steden met relatief tevreden inwoners ondervertegenwoordigd waren in de survey.

Tabel 16 Tevredenheid met de stad, dichotoom – gewogen

	Aantal	Percentage
0 – niet	303.463	23,55
1 – wel	985.382	76,45
Totaal	1.288.845	100,00

In tabel 17 presenteren we opnieuw verschillende betrouwbaarheidsintervallen rond het geschatte percentage tevreden met de stad. Die intervallen zijn ofwel default berekend met SPSS (ongewogen of gewogen met een herschaald gewicht, gemiddeld gelijk aan 1) ofwel met SPSS Complex Samples (met twee verschillende plannen). Voor de default analyses beperken we ons tot die waarbij de dummy als metrische variabele beschouwd wordt (in SPSS via Descriptives – Explore). De verschillen met de categorische berekening bleken in het eerste voorbeeld immers minimaal. De resultaten gewogen met het oorspronkelijke, niet-herschaalde, gewicht laten we achterwege omdat dat sowieso tot onrealistisch kleine betrouwbaarheidsintervallen leidt.

Tabel 17 Overzicht van de verschillende analyses van de tevredenheid met de stad, met bijhorende betrouwbaarheidsintervallen

Analyse	Parameter- schatting	95%-betrouwbaarheidsinterval		
		Ondergrens	Bovengrens	Breedte
Default SPSS-analyses				
ongewogen	71,96	71,32	72,60	1,28
gewogen, herschaald gewicht	76,45	75,85	77,06	1,21
Complex Samples-analyses				
volledige plan	76,45	75,65	77,24	1,59
plan dat geen rekening houdt met stratificatie (alleen gewichten)*	76,45	75,63	77,26	1,63

* Dit impliceert ook een andere eindigheidscorrectie. Er wordt slechts één algemeen populatietotaal ingegeven: 1.297.665, het totaal aantal personen in de doelpopulatie.

Voor de breedte van het betrouwbaarheidsinterval zijn de verschillen tussen de default analyses beperkt, net zoals de verschillen tussen de Complex Samples-analyses. Maar de verschillen tussen de default analyses en de Complex Samples-analyses zijn wel aanzienlijk. Ook hier blijkt de positieve impact van het stratificeren zeer beperkt. Of we in het plan rekening houden met de strata of niet, maakt immers amper iets uit. Correct omgaan met de gewichten is wel belangrijk. Dat vergroot de betrouwbaarheidsintervallen. In absolute aantallen ogen de getoonde verschillen misschien niet heel spectaculair. Maar de breedte van het interval stijgt toch van 1,21 procentpunten tot 1,59 procentpunten. Dat is een stijging met meer dan 30%!

Dat grote relatieve verschil is een gevolg van de zeer diverse gewichten, die zelf het resultaat zijn van het steekproefdesign. Zo zijn er in Gent 1.021 respondenten voor een populatie van 204.037 personen. In Turnhout zijn er 2.343 respondenten voor 34.602 personen. Het is dus logisch dat dat uitmondt in zeer verschillende gewichten die grotere standaardfouten met zich meebrengen en de precisie van schattingen voor de volledige populatie verminderen. Als de stadsmonitor uitsluitend schattingen voor die volledige populatie tot doel had, was een proportioneel gestratificeerde steekproef de beste keuze. Maar omdat de precisie van schattingen per stad (en voor Turnhout ook per stadsdeel) eveneens een expliciete doelstelling was, is het afwijken van die proportionaliteit wel geoorloofd en zelfs noodzakelijk. Bij een volledig proportioneel gestratificeerde steekproef van dezelfde omvang zou het onmogelijk zijn om betrouwbare uitspraken te doen voor Turnhout, laat staan voor stadsdelen uit Turnhout. Maar de disproportionele verdeling verlaagt dus wel de precisie van schattingen voor de volledige populatie. De optimale verdeling van de steekproefeenheden over de verschillende steden en stadsdelen varieert dus naargelang de onderzoeksvraag. Er is niet één ideale verdeling.

3.3.2. Significantietoetsen

We analyseren dezelfde variabele (tevredenheid met de stad) en gaan na of er een verschil is in tevredenheid tussen Belgen en niet-Belgen. Dat onderzoeken we met een eenvoudige 2x2-kruistabel en bijhorende χ^2 -toets. Tabel 18 toont de ongewogen resultaten.

Tabel 18 Tevredenheid met de stad volgens nationaliteit – ongewogen

Opleidingsniveau		Tevreden met stad		Totaal
		Niet	Wel	
Niet-Belg	Aantal	259	811	1.070
	rijpercentage	24,2%	75,8%	
Belg	Aantal	4.987	12.652	17.639
	rijpercentage	28,3%	71,7%	
Totaal	Aantal	5.246	13.463	18.709
	rijpercentage	28,0%	72,0%	

$$\chi^2=8,270, p = 0,004$$

We merken een duidelijk verschil tussen Belgen en niet-Belgen onder de respondenten. Het aandeel tevredenen bij de niet-Belgen ligt 4 procentpunten hoger dan bij de Belgen. Op basis van de ongewogen toets besluiten we dat dit een significant verschil is ($p = 0,004$). Maar uit de vorige pagina's weten we al dat de gewogen schatting van het aandeel tevredenen in de populatie een stuk hoger ligt. De gewogen kruistabel (tabel 19) ziet er dan ook heel anders uit dan de ongewogen kruistabel. In die tabel wordt gewogen met een herschaald gewicht (gemiddelde gelijk aan 1).

Tabel 19 Tevredenheid met de stad volgens nationaliteit – gewogen met herschaalde gewicht

Opleidingsniveau		Tevreden met stad		Totaal
		Niet	Wel	
Niet-Belg	Aantal	309	949	1.258
	rijpercentage	24,6%	75,4%	
Belg	Aantal	4.094	13.348	17.442
	rijpercentage	23,5%	76,5%	
Totaal	Aantal	4.403	14.297	18.700
	rijpercentage	23,5%	76,5%	

$$\chi^2=0,775, p = 0,379$$

Het aandeel tevredenen is inderdaad hoger, maar alleen bij de Belgen. Hierboven gaven we de verklaring dat steden met relatief tevreden inwoners ondervertegenwoordigd waren. Hier kunnen we eraan toevoegen dat vooral de tevreden Belgen ondervertegenwoordigd zijn. De achterliggende verklaring is dat de relatie tussen nationaliteit en tevredenheid ook verschillen vertoont tussen de steden. In sommige steden zijn Belgen meer tevreden dan niet-Belgen, in andere steden is het omgekeerd. De verdere uitwerking van die relatie valt buiten het bestek van deze tekst, maar het resultaat is alvast dat we op basis van de gewogen analyse voor de populatie niet meer kunnen besluiten tot een significante samenhang tussen tevredenheid en nationaliteit. Gewogen is het aandeel tevredenen bij de Belgen zelfs iets hoger, maar dat verschil kan louter het gevolg zijn van toeval ($p = 0,379$).

Een analyse met Complex Samples (volledig plan) geeft dezelfde resultaten voor de percentages. Daarnaast kunnen ook schattingen van de aantallen in de populatie opgevraagd worden en betrouwbaarheidsintervallen rond de aantallen en rond de percentages. Tabel 20 toont die resultaten. Zo kunnen we voor de Belgen met 95% zekerheid zeggen dat het werkelijke aandeel tevreden zich bevindt tussen 75,7% en 77,3%. Bij de niet-Belgen is het betrouwbaarheidsinterval logischerwijze veel groter, omdat het aantal niet-Belgen in de steekproef veel kleiner is.

Tabel 20 Output van de kruistabel met SPSS Complex Samples

nationaliteitsgroep_2011 * tevredenheid_stad					
nationaliteitsgroep_2011			tevredenheid_stad		
			,00	1,00	Total
Niet-Belg	Population Size	Estimate	21316,572	65437,108	86753,680
	% within nationaliteitsgroep_2011	Estimate	24,6%	75,4%	100,0%
		Standard Error	2,0%	2,0%	0,0%
	95% Confidence Interval	Lower	20,8%	71,3%	100,0%
		Upper	28,7%	79,2%	100,0%
		Unweighted Count	259	811	1070
Belg	Population Size	Estimate	282146,593	919944,987	1202091,580
	% within nationaliteitsgroep_2011	Estimate	23,5%	76,5%	100,0%
		Standard Error	0,4%	0,4%	0,0%
	95% Confidence Interval	Lower	22,7%	75,7%	100,0%
		Upper	24,3%	77,3%	100,0%
		Unweighted Count	4987	12652	17639
Total	Population Size	Estimate	303463,164	985382,096	1288845,260
	% within nationaliteitsgroep_2011	Estimate	23,5%	76,5%	100,0%
		Standard Error	0,4%	0,4%	0,0%
	95% Confidence Interval	Lower	22,8%	75,6%	100,0%
		Upper	24,4%	77,2%	100,0%
		Unweighted Count	5246	13463	18709

Voor ons is de bijhorende significantietoets relevanter. Net zoals bij de default gewogen toets vinden we dat het verband niet significant is. De kans dat het gevonden verschil louter het gevolg is van toeval, is zelfs nog een stuk hoger ($p = 0,590$ tegenover $p = 0,379$), zie tabel 21.

Tabel 21 Output van de significantietoets van SPSS Complex Samples bij de kruistabel

Tests of Independence							
		Chi-Square	Adjusted F	df1	df2	Sig.	
nationaliteitsgroep_2011 * tevredenheid_stad	Pearson	,790	,293	1	18789		,588
	Likelihood Ratio	,782	,290	1	18789		,590

The adjusted F is a variant of the second-order Rao-Scott adjusted chi-square statistic. Significance is based on the adjusted F and its degrees of freedom.

Net zoals bij de betrouwbaarheidsintervallen toont dit voorbeeld dat meer statistisch voorbehoud nodig is bij de uitspraken die wij doen op basis van deze data dan default-analyses in SPSS doen uitschijnen. De variatie in de gewichten heeft een impact op de inferentiële uitspraken die mogelijk zijn. Het is noodzakelijk om te corrigeren voor die differentiële gewichten.

3.3.3. Analyse van een deel van het databestand

Bij de survey van de stadsmonitor is er nog een ander argument om je te wenden tot Complex Samples-analyses dan correcte betrouwbaarheidsintervallen en significantietoetsen. Het vergemakkelijkt ook analyses van een deel van het databestand. Vaak zijn onderzoekers geïnteresseerd in één bepaalde stad. Het selecteren van respondenten heeft bij default-analyses zeer onwenselijke effecten op de relatieve grootte van de gewichten. Als je bijvoorbeeld uitsluitend de gegevens voor Turnhout wil analyseren, kan je geen gebruik maken van de gewichten die standaard beschikbaar gesteld worden voor de volledige dataset, ook niet de herschaalde.

Kijken we naar de analyse voor Turnhout van dezelfde variabele 'tevredenheid met de stad' in tabel 22. Het ongewogen percentage, het percentage onder de respondenten, is gelijk aan

57,05% met een 95%-betrouwbaarheidsinterval dat net groter is dan 4,02 procentpunten. Als we het oorspronkelijke gewicht gebruiken, merken we dat het geschatte percentage tevreden in de populatie toch iets hoger is: 58,30%. We weten dat het betrouwbaarheidsinterval van 1,04 procentpunt natuurlijk te klein is. Herschalen van het gewicht dan maar? Als we de herschaling toepassen die geldt voor het volledige bestand (en ook gebruikt werd in tabellen 17 en 19), houden we amper 498 respondenten over voor Turnhout. Het bijhorende betrouwbaarheidsinterval (bijna 8,7 procentpunten) is dan ook veel te groot. Als we ons beperken tot default-analyses, is de enige optie in dit geval een nieuwe herschaling, specifiek voor Turnhout. In die nieuwe herschaling moeten we er dus voor zorgen dat het gemiddelde van de gewichten *gelijk is aan 1 voor de respondenten uit Turnhout*. Bij default-analyses hebben we dus eigenlijk een nieuwe berekening van de gewichten nodig voor elke deelverzameling van het totale bestand die we willen analyseren. Erg praktisch is dat niet. Voor Complex Samples-analyses van slechts een deel van het databestand moeten er geen nieuwe gewichten worden berekend. Bovendien blijkt het default geschatte betrouwbaarheidsinterval ook met het herschaalde gewicht iets te klein. Met Complex Samples wordt het geschat op 4,61 procentpunten. Ook voor een analyse van alleen Turnhoutse respondenten verlagen de gewichten de precisie en vergroten ze de betrouwbaarheidsintervallen.

Tabel 22 Overzicht van de verschillende analyses van de tevredenheid met de stad voor Turnhout, met bijhorende betrouwbaarheidsintervallen

Analyse	N	Parameter-schatting	95%-betrouwbaarheidsinterval		
			Ondergrens	Bovengrens	Breedte
Default SPSS-analyses					
ongewogen, metrisch (explore)	2.326	57,05	55,04	59,06	4,02
gewogen, niet-herschaald gewicht	34.602	58,30	57,78	58,82	1,04
gewogen, herschaald gewicht (volledig bestand)	498	58,30	53,96	62,65	8,69
gewogen, herschaald gewicht (specifiek Turnhout)	2.325	58,30	56,30	60,31	4,01
Complex Samples-analyse					
volledige plan	2.326	58,30	55,98	60,59	4,61

Tot slot tonen we nog een significantietoets op de Turnhoutse deelpopulatie. We gaan daarbij na in welke mate de tevredenheid verschilt volgens stadsdeel. De ongewogen tabel 23 toont duidelijke verschillen. In het centrum is 61,3% van de respondenten tevreden, in Stadsbos en Noorden is dat duidelijk minder dan de helft (47,3%). Volgens de chi-kwadraattoets blijken die verschillen ook net significant ($p = 0,048$).

Tabel 23 Tevredenheid met de stad volgens stadsdeel in Turnhout – ongewogen

Opleidingsniveau		Tevreden met stad		Totaal
		Niet	Wel	
Blijkhoeft	Aantal rijpercentage	153 44,3%	192 55,7%	345
Centrum	Aantal rijpercentage	181 38,7%	287 61,3%	468
Schorvoort	Aantal rijpercentage	141 40,1%	211 59,9%	352
Stadsbos en Noorden	Aantal rijpercentage	88 52,4%	80 47,6%	168
Stedelijk Wonen Oost	Aantal rijpercentage	149 43,6%	193 56,4%	342
Stedelijk Wonen West	Aantal rijpercentage	138 41,9%	191 58,1%	329
Zevendonk	Aantal rijpercentage	149 46,3%	173 53,7%	322
Totaal	Aantal rijpercentage	999 42,9%	1.327 57,1%	2.326

 $\chi^2=12,706, p = 0,048$

Een default gewogen toets kan terug met verschillende gewichten. Bij gebruik van het oorspronkelijke gewicht zal alles natuurlijk significant zijn. Het gebruik van het op basis van de volledige steekproef herschaalde gewicht, zorgt voor een veel te kleine steekproefomvang in Turnhout. De enige min of meer aanvaardbare optie is het specifiek voor Turnhout herschaalde gewicht. Tabel 24 toont die optie. Gewogen zijn de verschillen tussen de twee uitersten iets kleiner (13 procentpunten verschil tussen Centrum en Stadsbos en Noorden versus ongewogen 13,7 procentpunten), maar ze zijn vooral niet meer significant ($p = 0,377$). De grootste afwijking van het globale percentage, die van Stadsbos en Noorden, lijkt minder robuust door het veel kleinere gewogen aantal respondenten (65 tegenover 168). De chi-kwadraatwaarde daalt dus omdat het stadsdeel dat het sterkst verschilt van het totaal voor Turnhout kleiner wordt na weging.

Tabel 24 Tevredenheid met de stad volgens stadsdeel in Turnhout – gewogen met voor Turnhout herschaalde gewicht

Opleidingsniveau		Tevreden met stad		Totaal
		Niet	Wel	
Blijkhoef	Aantal	87	112	199
	rijpercentage	43,7%	56,3%	
Centrum	Aantal	306	472	778
	rijpercentage	39,3%	60,7%	
Schorvoort	Aantal	70	104	174
	rijpercentage	40,2%	59,8%	
Stadsbos en Noorden	Aantal	34	31	65
	rijpercentage	52,3%	47,7%	
Stedelijk Wonen Oost	Aantal	207	276	483
	rijpercentage	42,9%	57,1%	
Stedelijk Wonen West	Aantal	214	300	514
	rijpercentage	41,6%	58,4%	
Zevendonk	Aantal	52	61	113
	rijpercentage	46,0%	54,0%	
Totaal	Aantal	970	1.365	2.326
	rijpercentage	41,7%	58,3%	

$$\chi^2=6,425, p = 0,377$$

Logischerwijze is geen van beide toetsen (die bij tabel 23 en bij tabel 24) correct. Omdat de ongewogen toets geen rekening houdt met de impact van de gewichten op de precisie, heeft hij een te hoog statistisch onderscheidingsvermogen. Die toets zal te snel de nulhypothese verwerpen. Bovendien toetst hij verschillen tussen percentages die niet de geschatte populatiepercentages zijn. De toets op basis van de voor Turnhout herschaalde gewichten, houdt geen rekening met de werkelijke aantallen in de verschillende categorieën. Complex Samples houdt rekening met de gewichten en de werkelijke aantallen en biedt dus de oplossing. Tabellen 25 en 26 tonen de resultaten van de Complex Samples-analyse. De verschillen in geschatte percentages zijn (uiteraard) dezelfde als die in tabel 24. De significantietoets geeft aan dat het gevonden verschil aan toeval te wijten kan zijn ($p = 0,297$). De toets bij tabel 23 heeft een nog veel lagere p -waarde. Dat komt omdat die andere percentages vergelijkt, maar voornamelijk omdat die geen rekening houdt met de impact van de gewichten op de precisie. Die impact is ook als we ons beperken tot Turnhout aanzienlijk. Desalniettemin verhoogt het design, waarvan de gewichten een uitloper zijn, het statistisch onderscheidingsvermogen voor toetsen voor verschillen tussen stadsdelen. De significantie van de Complex Samples-toets ligt immers wel in de lijn van die bij tabel 24, maar hetzelfde verschil in procenten wordt toch duidelijk minder waarschijnlijk geacht, gegeven de nulhypothese van geen verschil in de populatie ($p = 0,297$ tegenover $p = 0,377$). Dit is een illustratie dat ons steekproefdesign met oversampling van de kleine stadsdelen toch meer statistisch onderscheidingsvermogen heeft dan een proportionele verdeling (die zou aansluiten bij de verdeling in tabel 24). Het verschil is misschien wel kleiner dan verwacht. Het statistisch onderscheidingsvermogen zou nog toenemen als het kleinste stadsdeel nog meer oversampled zou worden.

Tabel 25 Output van de kruistabel met SPSS Complex Samples

stratum_gewichten * tevredenheid_stad						
stratum_gewichten				tevredenheid_stad		
				,00	1,00	Total
Turnhout-Blijkhoeve	% within stratum_gewichten	Estimate		43,8%	56,2%	100,0%
		Standard Error		2,5%	2,5%	0,0%
		95% Confidence Interval	Lower	38,9%	51,2%	100,0%
			Upper	48,8%	61,1%	100,0%
		Unweighted Count		153	192	345
Turnhout-Centrum	% within stratum_gewichten	Estimate		39,3%	60,7%	100,0%
		Standard Error		2,3%	2,3%	0,0%
		95% Confidence Interval	Lower	34,9%	56,0%	100,0%
			Upper	44,0%	65,1%	100,0%
		Unweighted Count		181	287	468
Turnhout-Schorvoort	% within stratum_gewichten	Estimate		40,2%	59,8%	100,0%
		Standard Error		2,5%	2,5%	0,0%
		95% Confidence Interval	Lower	35,4%	54,9%	100,0%
			Upper	45,1%	64,6%	100,0%
		Unweighted Count		141	211	352
Turnhout-Stadsbos en Noorden	% within stratum_gewichten	Estimate		51,8%	48,2%	100,0%
		Standard Error		3,7%	3,7%	0,0%
		95% Confidence Interval	Lower	44,6%	41,0%	100,0%
			Upper	59,0%	55,4%	100,0%
		Unweighted Count		88	80	168
Turnhout-Stedelijk Wonen Oost	% within stratum_gewichten	Estimate		42,8%	57,2%	100,0%
		Standard Error		2,7%	2,7%	0,0%
		95% Confidence Interval	Lower	37,7%	51,8%	100,0%
			Upper	48,2%	62,3%	100,0%
		Unweighted Count		149	193	342
Turnhout-Stedelijk Wonen West	% within stratum_gewichten	Estimate		41,6%	58,4%	100,0%
		Standard Error		2,7%	2,7%	0,0%
		95% Confidence Interval	Lower	36,4%	52,9%	100,0%
			Upper	47,1%	63,6%	100,0%
		Unweighted Count		138	191	329
Turnhout-Stedelijk Wonen West	% within stratum_gewichten	Estimate		46,1%	53,9%	100,0%
		Standard Error		2,5%	2,5%	0,0%
		95% Confidence Interval	Lower	41,1%	48,9%	100,0%
			Upper	51,1%	58,9%	100,0%
		Unweighted Count		149	173	322
Totaal	% within stratum_gewichten	Estimate		41,7%	58,3%	100,0%
		Standard Error		1,2%	1,2%	0,0%
		95% Confidence Interval	Lower	39,4%	56,0%	100,0%
			Upper	44,0%	60,6%	100,0%
		Unweighted Count		999	1327	2326

Tabel 26 Output van de significantietoets van SPSS Complex Samples bij de kruistabel

Tests of Independence						
		Chi-Square	Adjusted F	df1	df2	Sig.
stratum_gewichten *	Pearson	6,185	1,225	4,248	9924,004	,297
tevredenheid_stad	Likelihood Ratio	6,143	1,217	4,248	9924,004	,301

The adjusted F is a variant of the second-order Rao-Scott adjusted chi-square statistic. Significance is based on the adjusted F and its degrees of freedom.

Ook de voorbeeldanalyses op de survey van de stadsmonitor tonen duidelijk aan dat het gebruik van SPSS Complex Samples nodig is. Die survey heeft verschillende doelstellingen: schattingen van betrouwbare parameters mogelijk maken voor steden (en/of stadsdelen), maar ook het vergelijken van die steden (en/of stadsdelen). Het resultaat is een complex design waarvan onder andere gewichten een logische uitloper zijn. Dat heeft een impact op de

precisie die door default-analyses genegeerd wordt. Het voorbeeld van de stadsmonitor toont ook dat het gebruik van Complex Samples niet noodzakelijk een 'last' betekent. Default-analyses van delen van het databestand vereisen telkens nieuwe gewichten en voldoen dan nog niet. Daarentegen blijft hetzelfde Complex Samples Plan steeds bruikbaar, ook als slechts één stad of stadsdeel of een aantal steden samen geanalyseerd worden. Voor Complex Samples-analyses blijven altijd dezelfde gewichten van toepassing. Dat 'gebruiksgemak' compenseert de inspanning om je te moeten wenden tot een bijkomende module in SPSS.

Slotbeschouwingen

Algemene statistische tekstboeken vertrekken voor inferentiële formules meestal van de assumptie van een enkelvoudige aselechte steekproef. Bij surveyonderzoek is die assumptie vaker niet dan wel realistisch. Stratificatie, clustering en gewichten noodzaken andere berekeningswijzen van standaardfouten, betrouwbaarheidsintervallen en significantietoetsen. Die andere berekeningswijzen zijn ondertussen ook voldoende gedocumenteerd in de survey en steekproefliteratuur (zie bijvoorbeeld Cochran, 1977; Chambers & Skinner, 2003), maar nog niet voldoende doorgedrongen in de algemene handboeken. De grotere statistische softwareprogramma's bieden tegenwoordig ook de mogelijkheden om bij de meest courante analyses de stratificatie, clustering en ongelijke gewichten wel in rekening te brengen. Deze tekst had als bedoeling de koudwatervrees die op dit vlak heerst, wat af te zwakken. Softwarematig focuste de tekst volledig op SPSS en werd aangetoond hoe bij analyses van de SCV-survey en de survey voor de stadsmonitor correcte betrouwbaarheidsintervallen en significanties bekomen kunnen worden met SPSS Complex Samples. Andere programma's zoals Stata, SAS en R hebben dezelfde of nog meer mogelijkheden, maar vallen buiten het bestek van deze tekst.

De voorbeelden tonen aan dat default-analyses vaak tot verkeerde conclusies kunnen leiden en dat de Complex Samples-analyses dus geen overbodige luxe betekenen. In die zin blijft het op z'n minst ongelukkig dat correcte inferentiële statistiek in SPSS een bijkomende, dure, module vereist.

Bibliografie

Brogan, D. (2005). Sampling error estimation for survey data. In: United Nations. Department of Economic and Social Affairs Statistics Division (red.). *Household Sample Surveys in Developing and Transition Countries. Studies in Methods F. 96*. New York: United Nations.

Carton, A., Vander Molen, T. & Pickery, J. (2013). *Sociaal-culturele verschuivingen in Vlaanderen 2012. Methodologisch-technisch rapport en procesevaluatie van de dataverzameling. SVR-Methoden en Technieken 2013/5*. Brussel: Studiedienst van de Vlaamse Regering.

Chambers, R. & Skinner, C. (2003). *Analysis of Survey Data*. Chichester: Wiley.

Cochran, W. (1977), *Sampling Techniques*. New York: Wiley.

Goedemé, T. (2010). *The standard error of estimates based on EU-SILC. An exploration through the Europe 2020 poverty indicators. CSB Working Paper 10/09*. Antwerpen: Universiteit Antwerpen/CSB.

IBM (2011). *IBM SPSS Complex Samples 20*. Chicago: IBM Corporation.

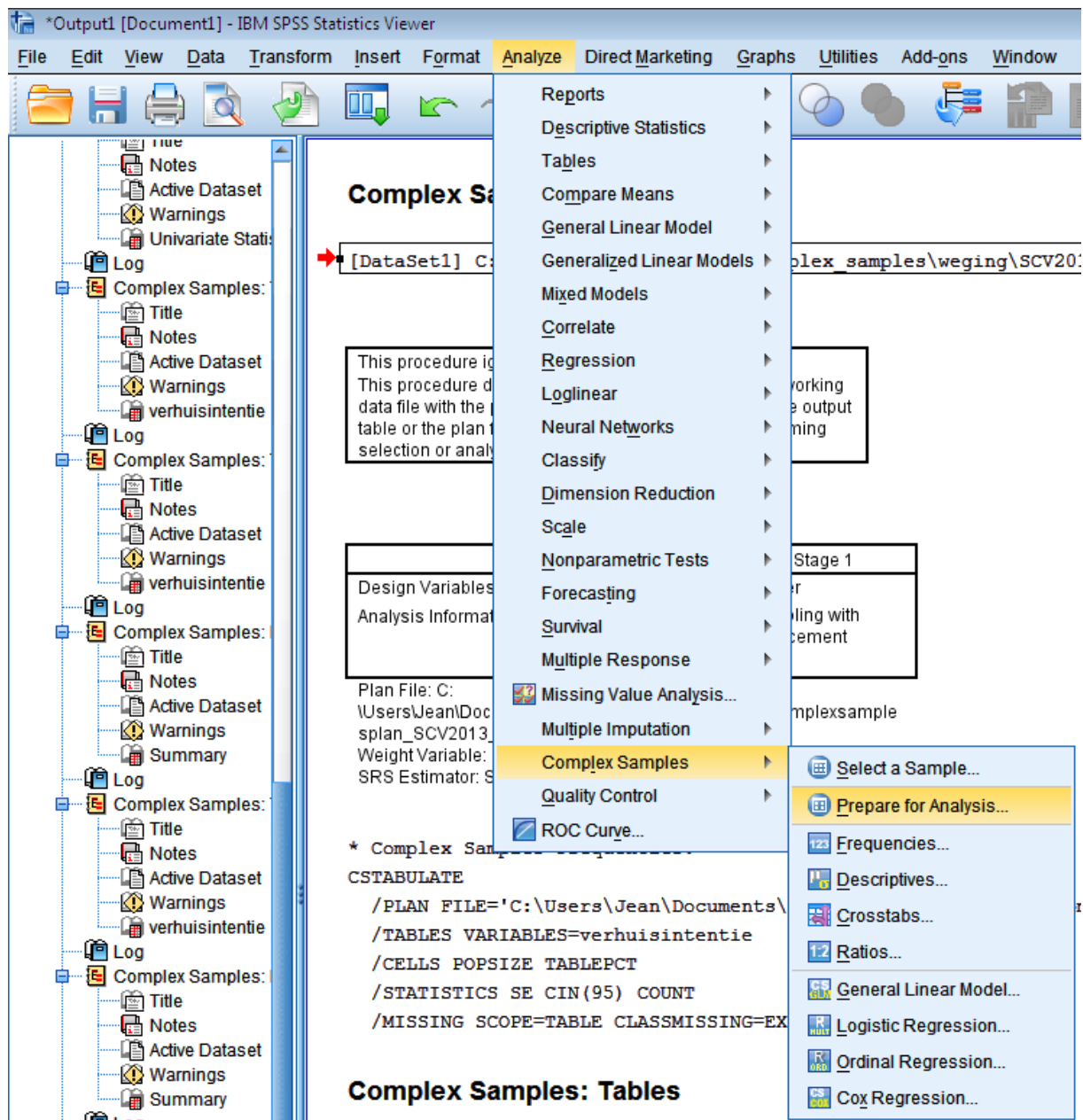
Kalton, G. & Flores-Cervantes, I. (2003). Weighting Methods. In: *Journal of Official Statistics*, 19 (2), 81-97.

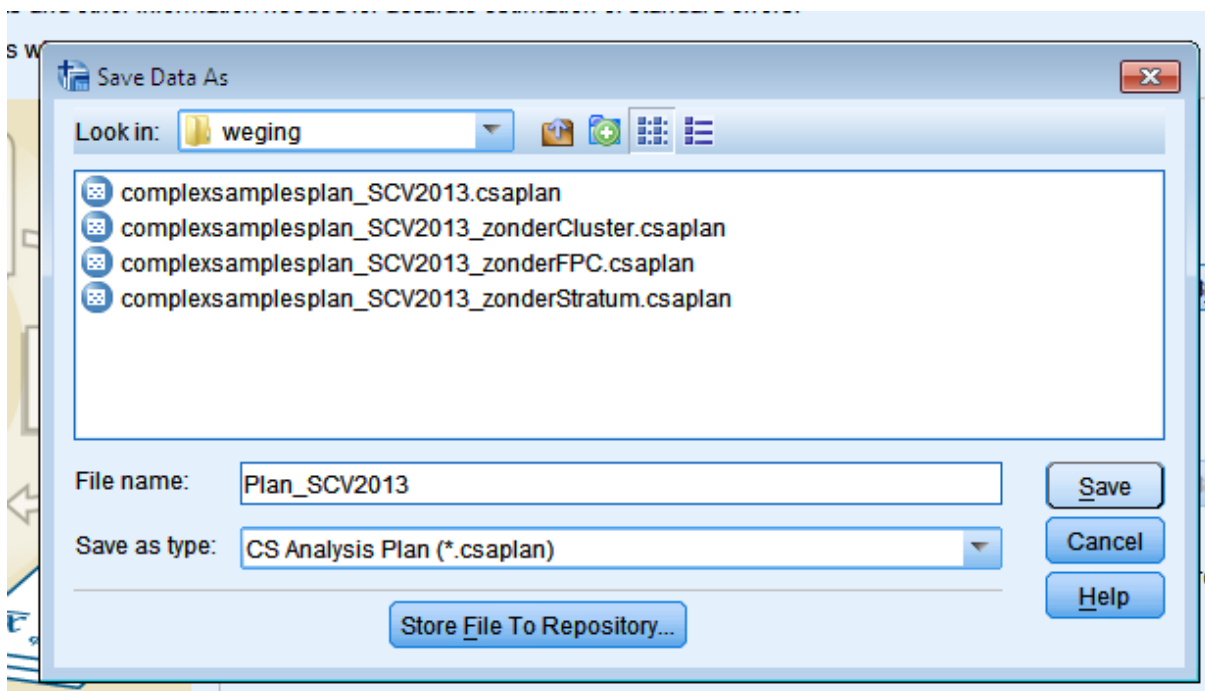
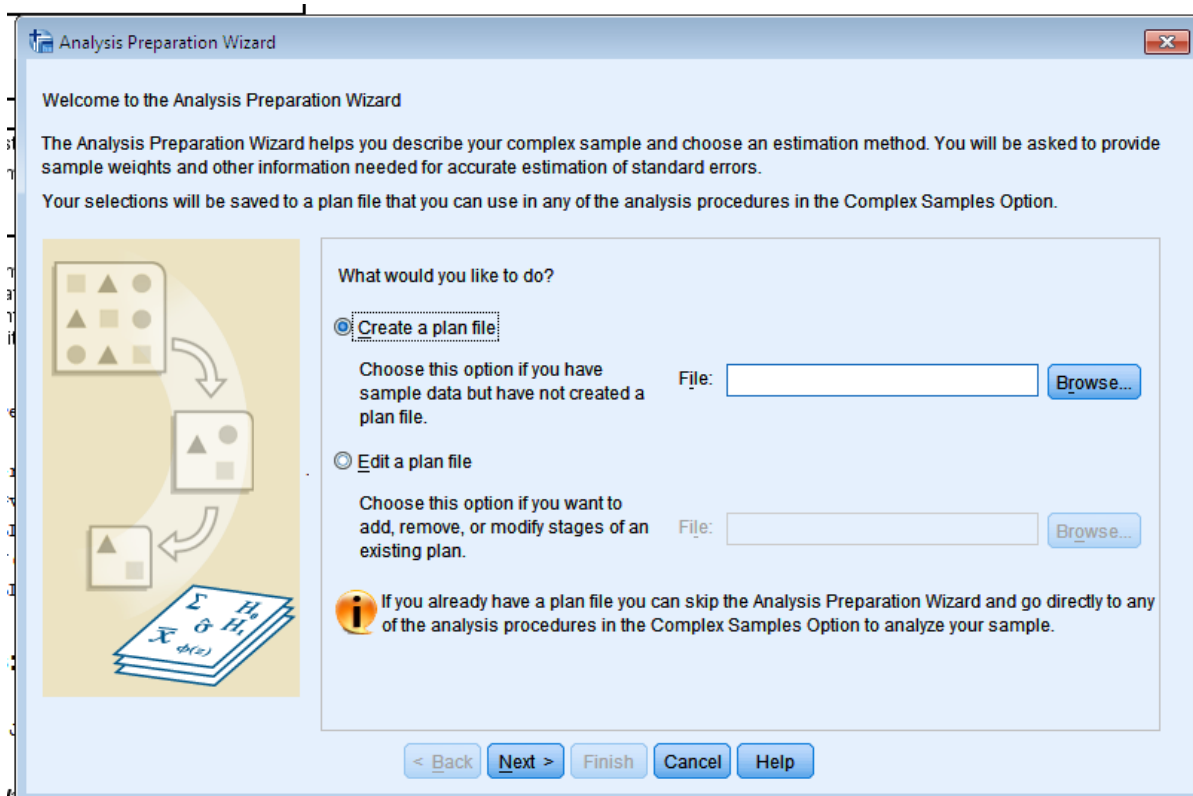
Rust, K. (1985). Variance Estimation for Complex Estimators in Sample Surveys. In: *Journal of Official Statistics*, 1 (4), 381-397.

Schelfaut, H. (2009). *Survey Stadsmonitor "Thuis in de stad 2008". Methodologisch Rapport. SVR-Technisch Rapport 2009/1*. Brussel: Studiedienst van de Vlaamse Regering.

Bijlagen

Bijlage 1 Invoeren van het Complex Samples plan voor SCV2013 met SPSS





Analysis Preparation Wizard

Welcome to the Analysis Preparation Wizard

The Analysis Preparation Wizard helps you describe your complex sample and choose an estimation method. You will be asked to provide sample weights and other information needed for accurate estimation of standard errors.

Your selections will be saved to a plan file that you can use in any of the analysis procedures in the Complex Samples Option.



What would you like to do?

☒ **Create a plan file**

Choose this option if you have sample data but have not created a plan file.

File:

☐ **Edit a plan file**

Choose this option if you want to add, remove, or modify stages of an existing plan.

File:

 If you already have a plan file you can skip the Analysis Preparation Wizard and go directly to any of the analysis procedures in the Complex Samples Option to analyze your sample.

Analysis Preparation Wizard

Stage 1: Design Variables

In this panel you can select variables that define strata or clusters. A sample weight variable must be selected in the first stage.

You can also provide a label for the stage that will be used in the output.

Welcome

Stage 1

Design Variables

Estimation Method

Summary

Completion

Variables:

- ☒ mytype
- ☒ gezinstype
- ☒ oversamplingfactor
- ☒ selectiekans
- ☒ designgewicht13
- ☒ respons_nieuw
- ☒ woningtype_4cat
- ☒ burgstaat2
- ☒ Predicted probability (PR...
- ☒ responsquintiel
- ☒ nonresponsgewicht13
- ☒ gewicht_na_stap2
- ☒ poststratificatie_gewest...
- ☒ gewicht_voor_poststratif...
- ☒ opl_eak_gesl
- ☒ poststratificatie_eak13
- ☒ scv_gewicht_herschaal...
- ☒ verhuisintentie

Strata:

☒ stratum

Clusters:

☒ cluster

Sample Weight:

☒ scv_gewicht13

Stage Label:

Analysis Preparation Wizard

Stage 1: Estimation Method

In this panel you select a method for estimating standard errors.

The estimation method depends on assumptions about how the sample was drawn.

Which of the following sample designs should be assumed for estimation?

☒ **WR (sampling with replacement)**

If you choose this option you will not be able to add additional stages. Any sample stages after the current stage will be ignored when the data are analyzed.

☒ Use finite population correction (FPC) when estimating variance under simple random sampling assumption

☐ **Equal WOR (equal probability sampling without replacement)**

The next panel will ask you to specify inclusion probabilities or population sizes.

☐ **Unequal WOR (unequal probability sampling without replacement)**

Joint probabilities will be required to analyze sample data. This option is available in stage 1 only.

< Back Next > Finish Cancel Help

Analysis Preparation Wizard

Stage 1: Plan Summary

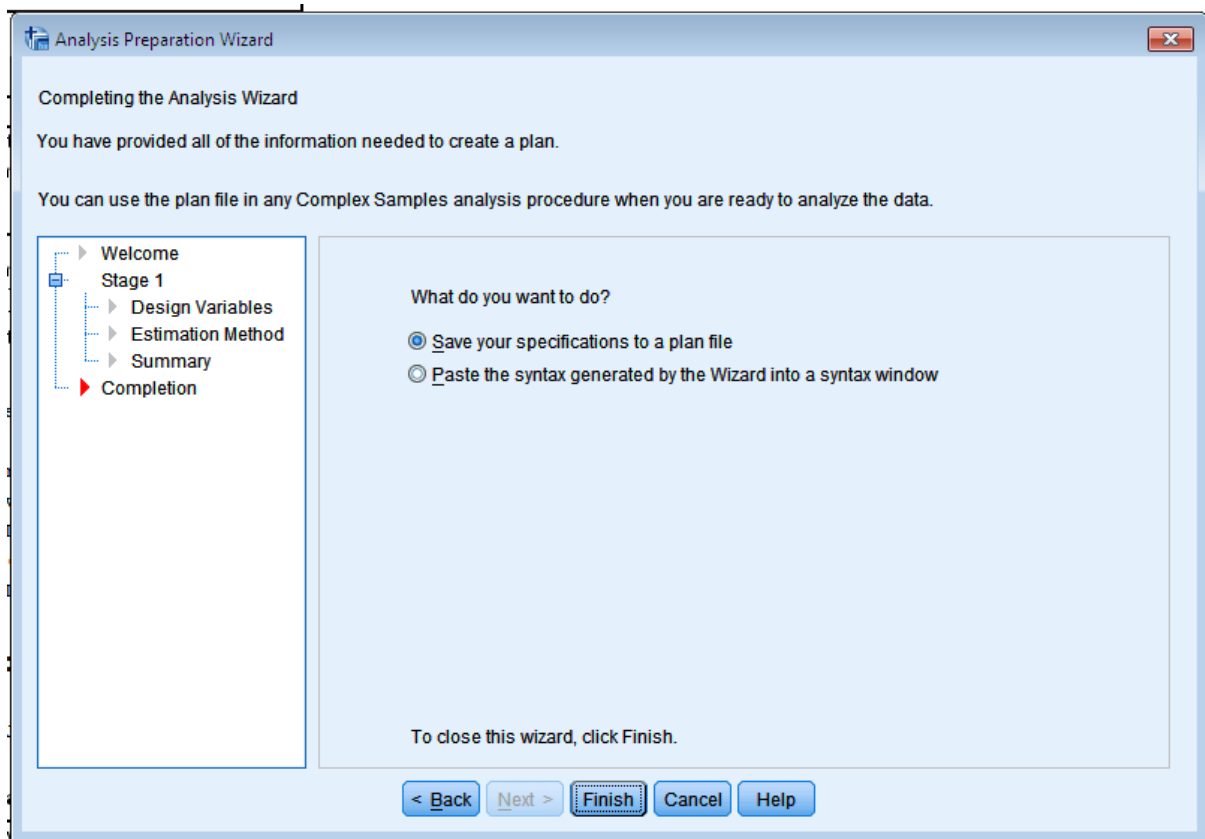
This panel summarizes the plan so far. The next step is the Completion panel.

Summary:

Stage	Label	Strata	Clusters	Weights	Size	Method
1	(None)	stratum	cluster	scv_gewich t13	(n/a)	WR

File: C:\Users\Uean\Documents\complex_samples\weging\Plan_SCV2013.csaplan

< Back Next > Finish Cancel Help



Bijlage 2 Opvragen van een Complex Samples-frequentietabel met betrouwbaarheidsintervallen

The screenshot shows the IBM SPSS Statistics Viewer interface. The 'Analyze' menu is open, and the path 'Complex Samples' is highlighted. The 'Complex Samples: Plan' sub-menu is also open, showing options like 'Select a Sample...', 'Prepare for Analysis...', 'Frequencies...', 'Descriptives...', 'Crosstabs...', 'Ratios...', 'General Linear Model...', 'Logistic Regression...', 'Ordinal Regression...', and 'Cox Regression...'. The 'Frequencies...' option is selected.

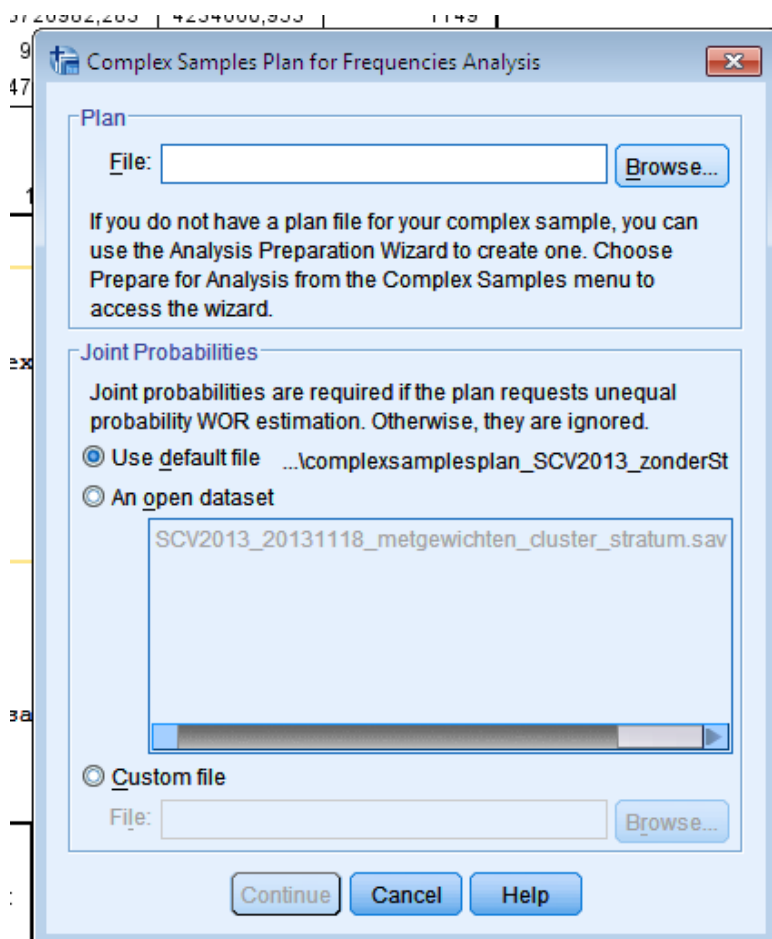
The main window displays the 'Complex Samples: Plan' dialog box. The 'Population Size' is set to 345, and the '% of Total' is set to 4%. The 'Analysis Plan' section shows the following commands:

```

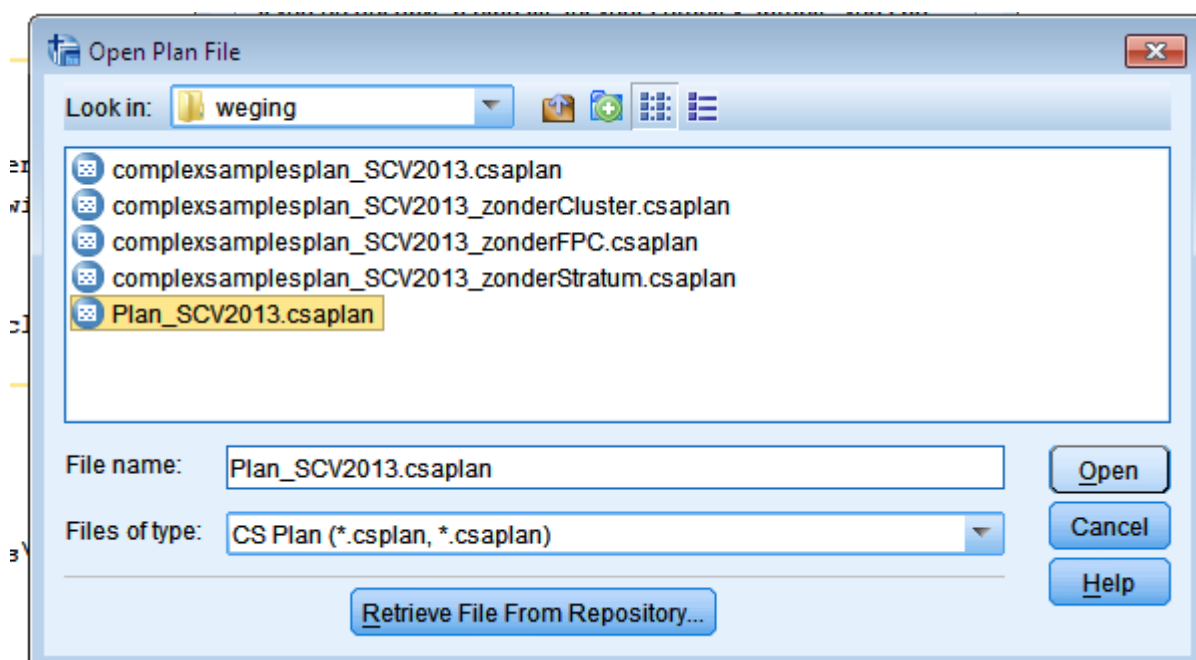
* Analysis Plan
CSPLAN ANALYSIS
/PLAN FILE=
/PLANVARS=
/SRSESTIMATE=
/PRINT PLAN
/DESIGN STRATA=stratum CLUSTER=cluster
/ESTIMATOR TYPE=WR.
  
```

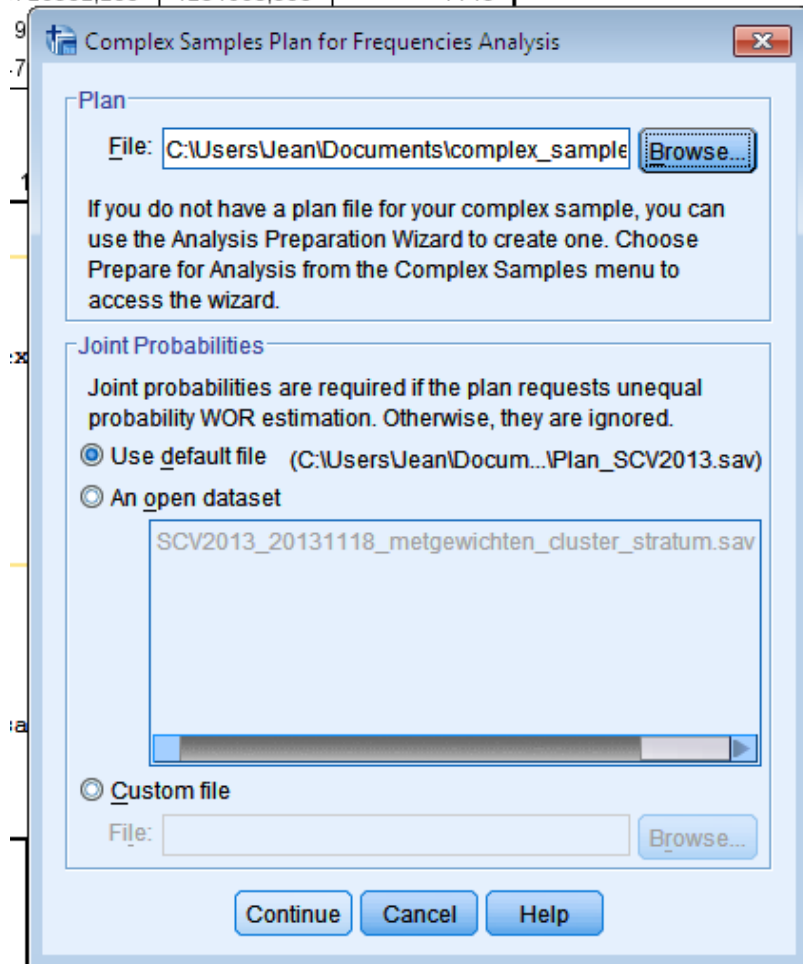
The 'Warnings' section shows the following table:

	95% Confidence Interval	
	Lower	Upper
Population Size	3726982,283	4234660,953
% of Total	965266,264	1239051,164
	4781948,497	5384012,168
	75,9660%	80,4965%
	19,5035%	24,0340%
	100,0000%	100,0000%

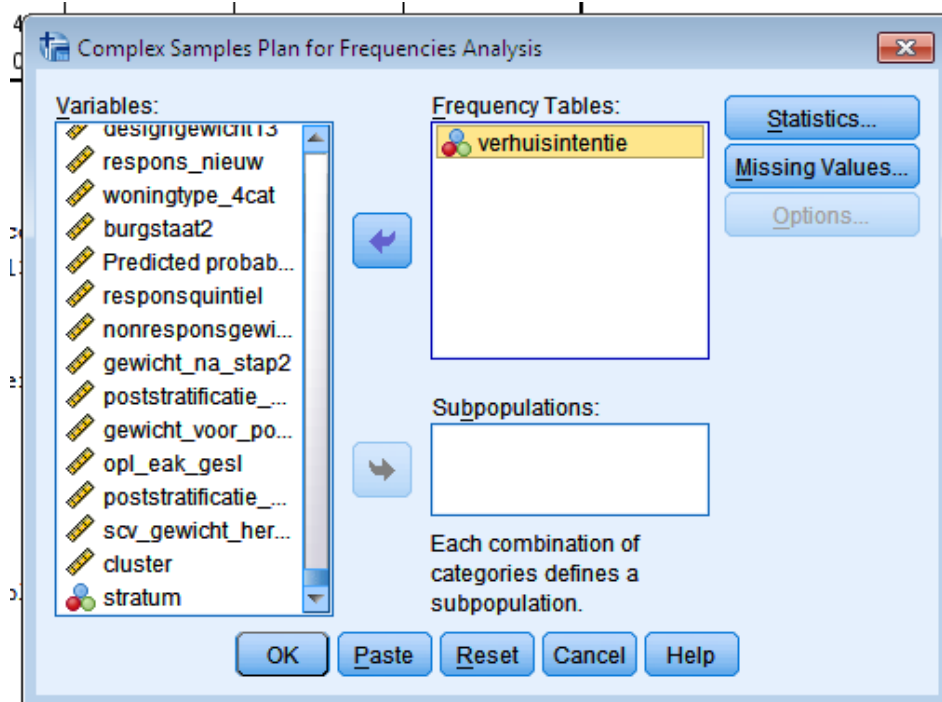


Je moet een plan ingeven, te kiezen via Browse...

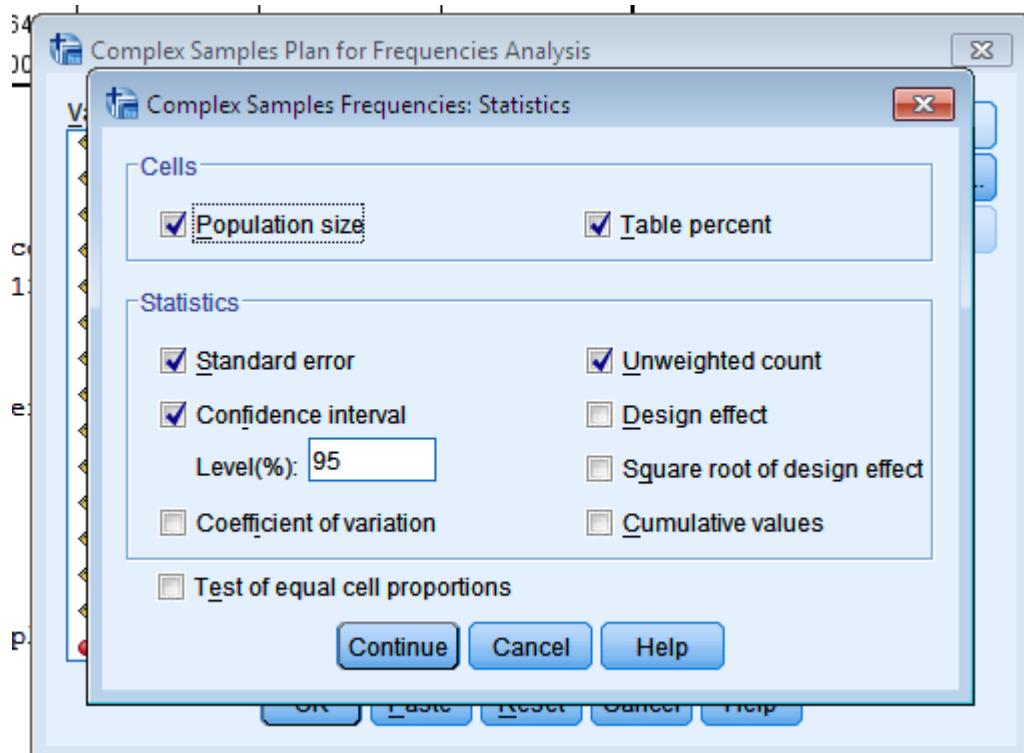




Nu een variabele selecteren...



Tot slot: aanvinken wat je allemaal wil te zien krijgen.



Bijlage 3 Invoeren van het Complex Samples plan voor de survey van de stadsmonitor 2011

*data_stadsmonitor_2008_2011_vrderden_2011_eigen_analyse.sav [DataSet1] - IBM SPSS Statistics Data Editor

File Edit View Data Transform **Analyze** Direct Marketing Graphs Utilities Add-ons Window Help

Reports
Descriptive Statistics
Tables
Compare Means
General Linear Model
Generalized Linear Models
Mixed Models
Correlate
Regression
Loglinear
Neural Networks
Classify
Dimension Reduction
Scale
Nonparametric Tests
Forecasting
Survival
Multiple Response
Missing Value Analysis...
Multiple Imputation
Complex Samples
Quality Control
ROC Curve...

	Name	Decimals	Label	Value Labels
1	surveyjaar	0		None
2	iden_gemeente_gewicht	2		{1,00,
3	stratum_gewichten_2008	2		None
4	stratum_gewichten_2011	2		{11,00
5	basisgewicht	2		None
6	respons	2		None
7	nonresponsgewicht	2		None
8	definitiefgewicht	2		None
9	jaartal	0		None
10	unieknummer2008	0		None
11	barcode_vragenlijst_2011	0	Barcode vragen...	None
12	identificatienummer_2011	0		None
13	gemeente_2011	2		None
14	gemeente_beide	2		{1,00,
15	POSTCODE_beide	0	POSTCODE	None
16	WOONSTAD	0	WOONSTAD	{1, Mir
17	woonstadduurbin	2		{1,00,
18	WOONWONING_2011	0	WOONWONING	{1, Mir
19	TEVR_BUU			Zee
20	TEVR_STA			Zee
21	GBOUW_B			Hel
22	STRTEN_B			Hel
23	VGROEN_B			Hel
24	VSPLVRZ_B			Hel
25	VJUGDBTN			Hel
26	VACTPENSI			Hel
27	AANGPRTN			Hel
28	VHALTES_B			Hel
29	VBUSSEN_B			Hel
30	BEREIKCENTR_B			Hel

Select a Sample...
Prepare for Analysis...
Frequencies...
Descriptives...
Crosstabs...
Ratios...
General Linear Model...
Logistic Regression...
Ordinal Regression...
Cox Regression...

Analysis Preparation Wizard

Stage 1: Design Variables

In this panel you can select variables that define strata or clusters. A sample weight variable must be selected in the first stage.

You can also provide a label for the stage that will be used in the output.

Welcome
 Stage 1
 Design Variables
 Estimation Method
 Summary
 Completion

Variables:

- diploma2008 [dipl2008]
- diploma2011 [dipl_2011]
- dipl2008tris
- dipl2011tris
- diploma_beide
- opleidingbis_beide
- V45 [INKMN_2008]
- INKMN_2011
- INKOMEN [INKMN_2011...]
- iden_gemeente_compl...
- tevreden_met_buurt_dic...
- stratu_gewicht_vergelijk...
- tevredenheid_stad
- definitiefgewicht_hersch...
- gemeente_2011 = 13 (F...
- definitiefgewicht_hersch...
- populatie_totaal
- stratum_gewichten

Strata:

iden_weegstratum_complexsamples

Clusters:

Sample Weight:

definitiefgewicht

Stage Label:

< Back Next > Finish Cancel Help

Analysis Preparation Wizard

Stage 1: Estimation Method

In this panel you select a method for estimating standard errors.

The estimation method depends on assumptions about how the sample was drawn.

Welcome
 Stage 1
 Design Variables
 Estimation Method
 Size
 Summary
 Add Stage 2
 Completion

Which of the following sample designs should be assumed for estimation?

☐ **WR (sampling with replacement)**
 If you choose this option you will not be able to add additional stages. Any sample stages after the current stage will be ignored when the data are analyzed.
☒ Use finite population correction (FPC) when estimating variance under simple random sampling assumption

☒ **Equal WOR (equal probability sampling without replacement)**
 The next panel will ask you to specify inclusion probabilities or population sizes.

☐ **Unequal WOR (unequal probability sampling without replacement)**
 Joint probabilities will be required to analyze sample data. This option is available in stage 1 only.

! = incomplete section

< Back Next > Finish Cancel Help

Analysis Preparation Wizard

Stage 1: Size

In this panel you specify inclusion probabilities or population sizes for the current stage.

You can provide a size that is fixed across strata or specify sizes on a per-stratum basis.

Welcome

Stage 1

Design Variables

Estimation Method

Size

Summary

Add Stage 2

Completion

Variables:

- surveyjaar
- iden_gemeente_gewicht
- stratum_gewichten_2008
- stratum_gewichten_2011
- basisgewicht
- respons
- nonresponsgewicht
- jaartal
- unieknummer2008
- Barcode vragenlijst [barcode_...]
- identificatienummer_2011
- gemeente_2011
- gemeente_beide
- POSTCODE [POSTCODE_bei...]
- WOONSTAD [WOONSTAD]
- woonstadduurbin
- WOONWONING [WOONWONI...]

Units: **Inclusion Probabilities**

☒ Value:

☐ Unequal values for strata:

Define...

☐ Read values from variable:

= incomplete section

< Back Next > Finish Cancel Help

Analysis Preparation Wizard

Stage 1: Size

In this panel you specify inclusion probabilities or population sizes for the current stage.

You can provide a size that is fixed across strata or specify sizes on a per-stratum basis.

Welcome

Stage 1

Design Variables

Estimation Method

Size

Summary

Add Stage 2

Completion

Variables:

- diploma2008 [dipl2008]
- diploma2011 [dipl_2011]
- dipl2008tris
- dipl2011tris
- diploma_beide
- opleidingbis_beide
- V45 [INKMN_2008]
- INKMN_2011
- INKOMEN [INKMN_2011_origi...]
- iden_gemeente_complexsam...
- tevreden_met_buurt_dichotoo...
- stratu_gewicht_vergelijking
- tevredenheid_stad
- definitiefgewicht_herschaald
- gemeente_2011 = 13 (FILTE...
- definitiefgewicht_herschaald_...
- stratum_gewichten

Units: **Population Sizes**

☐ Value:

☐ Unequal values for strata:

Define...

☒ Read values from variable:

populatietotaal

< Back Next > Finish Cancel Help

Analysis Preparation Wizard

Stage 1: Plan Summary

This panel summarizes the plan so far. You can add another stage to the plan.

If you choose not to add a stage the next panel is the Completion panel.

- Welcome
- Stage 1
 - Design Variables
 - Estimation Method
 - Size
 - Summary
- Add Stage 2
- Completion

Stage	Label	Strata	Clusters	Weights	Size	Method
1	(None)	iden_weegst		definitiefgew icht	(Read from populatielot	Equal WOR

File: C:\Users\Jaan\Documents\complex_...\complex_samples_plan_stadsmonitor_2011.csaplan

Do you want to add stage 2?

☐ Yes, add stage 2 now
Choose this option if the sample contains another stage.

☒ No, do not add another stage now
Choose this option if this is the last stage of the sample.

Analysis Preparation Wizard

Completing the Analysis Wizard

You have provided all of the information needed to create a plan.

You can use the plan file in any Complex Samples analysis procedure when you are ready to analyze the data.

- Welcome
- Stage 1
 - Design Variables
 - Estimation Method
 - Size
 - Summary
- Add Stage 2
- Completion

What do you want to do?

☒ Save your specifications to a plan file
☐ Paste the syntax generated by the Wizard into a syntax window

To close this wizard, click Finish.

