



Masterproef Politieke Communicatie

**Responsstijlen in Surveyonderzoek: Het Modelleren van de
Acquiescence Responsstijl doormiddel van *Structural Equation
Modelling* en de RIRSMACS-methode.**

Kaat Bots



**Universiteit
Antwerpen**

**Promotor: Prof. Dr. Peter Thijssen
Copromoter: Dr. Tanja Perko
Verslaggever: Prof. Dr. Jan Beyers**

Master Politieke Communicatie (www.politiekecommunicatie.be)
Rolnummer studente: s0155850
Faculteit Politieke en Sociale Wetenschappen
Academiejaar 2015-2016

**Responsstijlen in Surveyonderzoek: Het modelleren van de Acquiescence Responsstijl
door middel van *Structural Equation Modelling* en de RIRSMACS-methode.**

Kaat Bots

Woord vooraf

Deze thesis zou niet mogelijk zijn geweest zonder de hulp van enkele mensen. Via deze weg wil ik dan ook vooral mijn promotor Professor Dr. Peter Thijssen bedanken voor de vele hulp en de immer interessante inzichten. Ook zou ik graag Danny Rouckhout bedanken om te helpen met het technische aspect van het onderzoek.

De data die gebruikt worden in dit onderzoek zijn afkomstig van het Studiecentrum voor Kernenergie in Mol. Graag zou ik hen daarvoor bedanken, en dan vooral mijn copromotor Dr. Tanja Perko. Zonder hun medewerking was deze thesis niet mogelijk geweest.

Abstract

Bij veel surveyonderzoek wordt er gebruik gemaakt van zogenaamde Likert schalen om attitudes te meten. Deze Likert schalen laten de respondent de verschillende surveyitems beoordelen aan de hand van een vijf-punten schaal die begint bij ‘helemaal niet akkoord’ en eindigt bij ‘helemaal akkoord’. Dit soort schalen zijn erg gevoelig aan responsstijlen. Responsstijlen wijzen op de neiging van de respondent om de surveyitems systematisch te beoordelen, los van de inhoud. In dit onderzoek ligt de focus op de *acquiescence* responsstijl (ARS), die wijst op het ja-knikken van de respondent. Door middel van twee statistische methodes, namelijk *Structural Equation Modelling* en de RIRSMACS-methode zal deze responsstijl (ARS) in kaart worden gebracht. Verschillende hypotheses over de oorzaken van ARS zullen getest worden aan de hand van deze methodes. Op het einde van het onderzoek kan er zo een discussie gevoerd worden over de oorzaken van ARS, en over de voor- en nadelen die beide methodes hebben bij het modelleren ervan.

Lijst met afkortingen:

ARS:	Acquiescence Response Style
CFA:	Confirmatory Factor Analysis
DARS:	Disacquiescence Response Style
ERS:	Extreme Response Style
MLRS:	Mild Response Style
MRS:	Middlepoint Response Style
NARS:	Net Acquiescence Response Style
NRS:	Noncontingent Response Style
RIRSMACS:	Representative Indicators Response Style Means and Covariance Structure
RMSEA:	Root Mean Square Error of Approximation
RS:	Response Style
SCK-CEN:	Studiecentrum Kernenergie – Centre d'études de l'énergie Nucléaire
SEM:	Structural Equation Modelling

Inhoudstafel

Woord vooraf	3
Abstract	4
Lijst met afkortingen:	5
HOOFDSTUK 1 – Introductie	8
1.1 Inleiding	8
1.2 Waarom is onderzoek naar responsstijlen nodig?	10
1.2.1 Maatschappelijke relevantie	10
1.2.2 Wetenschappelijke relevantie	11
1.2.3 Relevantie voor Politieke Communicatie	13
HOOFDSTUK 2 – De Theorie achter Responsstijlen	14
2.1 Wat is een Responsstijl?	14
2.2 Welke soorten responsstijlen bestaan er?	15
2.3 Wat is ARS?	16
2.4 Oorzaken van ARS	18
2.4.1 Het niveau van de stimuli	18
2.4.2 Het niveau van de respondent	19
HOOFDSTUK 3 – <i>Structural Equation Modelling</i> (SEM)	23
3.1 Het modelleren van ARS	23
3.2 <i>Structural Equation Modelling</i> (SEM)	24
3.3 Het onderzoek van Billiet en McClendon (2000)	25
3.4 Voorbeeld van een ongepubliceerd onderzoek naar ARS met SEM	28
3.5 Analyse van de SCK-Barometer 2015	35
HOOFDSTUK 4 – De Resultaten	43
4.1 Hoe rapporteer je over <i>Structural Equation Modelling</i> ?	43
4.2 Resultaten van de hypotheses testen	43
4.2.1 Hypothese 2: Dubbelzinnige items	43
4.2.2 Hypothese 3: Gender	45
4.2.3 Hypothese 4: Leeftijd	47
4.2.4 Hypothese 5a en 5b: Opleidingsniveau en Loon	47
4.2.5 Hypothese 6a en 6b: Verschil tussen Waalse en Vlaamse respondenten	49
HOOFDSTUK 5 – De RIRSMACS-methode	53
5.1 RIRSMACS-methode geïllustreerd	53
5.2 RIRSMACS-methode toegepast op SCK-Barometer 2015	56
5.3 Combinatie van RIRSMACS-methode en SEM	58
HOOFDSTUK 6 – Discussie	61

Bibliografie	64
Bijlagen	67
1. Tabellen.....	67
2. Figuren	71

HOOFDSTUK 1 – Introductie

1.1 Inleiding

Onderzoek naar het meten van de ‘publieke opinie’ is niet weg te denken uit de politieke wetenschappen. Voor politici is kennis over de publieke opinie van groot belang om hun legitimiteit te kunnen versterken. Door middel van responsiviteit kan men reageren op wat er leeft bij de massa. Naast het versterken van hun legitimiteit, verhoogt het ook hun kansen om herkozen te worden. Het belang van dergelijk onderzoek is dus enkel gestegen doorheen de laatste jaren. Naast een belangrijke rol in politiek onderzoek, is de publieke opinie tegenwoordig ook van belang in de commerciële wereld. Als kennis macht is, dan geldt dit ook voor kennis over de publieke opinie. Wie weet wat er leeft bij het volk kan hierop in spelen en er politiek voordeel uit halen. Wie weet wat het publiek denkt over zijn/haar product, kan hierop in spelen, en er zo commercieel voordeel uit halen. Met het groeien van alle commerciële markten, werd kennis over de preferenties van de massa steeds waardevoller. Het is dus de vraag naar kennis over de publieke opinie die erg is toegenomen. Zoals dat gaat, is met deze vraag ook het aanbod gegroeid. Er zijn ondertussen al meerdere methodes om de publieke opinie te meten. Elk met zijn voor- en nadelen. Over het algemeen worden er vier methodes om publieke opinie te meten gebruikt, namelijk: surveyonderzoek, focusgroepen, experimenteel onderzoek en de inhoudsanalyse van massamedia (Glynn, Herbst, Lindeman, O’Keefe, Shapiro, 2016).

In dit onderzoek zal de focus liggen op surveyonderzoek. Surveyonderzoek wijst op het ondervragen van grote groepen respondenten door middel van vragenlijsten, ook wel enquêtes genoemd. Het gebruik van deze methode is zeer populair sinds het begin van de 20^{ste} eeuw, omdat toen de interesse in de massaopinie groeide. Surveyonderzoek is immers de methode bij uitstek waar grote groepen respondenten mee bereikt kunnen worden (Weijters et al., 2008). Hoewel het niet persé de goedkoopste methode is, kan er vaak wel een groot sample mee verzameld worden. Dit is van belang voor de validiteit en betrouwbaarheid van het onderzoek. Ondanks de voordelen en de grote populariteit van surveyonderzoek, heeft de methode ook verschillende nadelen. Het grootste nadeel is de mogelijke aanwezigheid van ‘bias’. Deze bias, die slaat op vooringenomenheid of partijdigheid, kan zowel aanwezig zijn bij de respondenten als bij de surveyitems en interviewers. Zo kunnen de uitkomsten van surveyonderzoek significant beïnvloed worden door de vragen op een bepaalde manier te stellen, of enkel een bepaald type respondenten te ondervragen. De aanwezigheid van bias zou de betrouwbaarheid en de validiteit van surveyonderzoek stevig aantasten. (Glynn, Herbst, Lindeman, O’Keefe, Shapiro, 2016).

Bij veel van de surveyonderzoeken die tegenwoordig plaatsvinden wordt er gretig gebruik gemaakt van Likert-schalen. Dit betekent dat het surveyitem als een statement wordt weergegeven waarop de

respondent dan moet aangeven in welke mate hij het al dan niet eens is met dat statement (De Vellis, 2003). Het vaakst wordt er gebruik gemaakt van de vijf-punten Likert schaal. Dit betekent dat de respondent elk statement kan beoordelen met vijf mogelijke antwoorden, die gaan van uitermate akkoord tot helemaal niet akkoord¹. Deze antwoorden worden dan gelinkt aan een cijfer om later de statistische analyse mee uit te voeren (De Vellis, 2003). Deze Likert schalen worden bijzonder vaak gebruikt in surveyonderzoek omdat het duidelijk is voor de respondent hoe de items te beoordelen, en gebruiksvriendelijk voor de onderzoeker om de data te verwerken. Dit type van surveyonderzoek brengt echter wel een probleem met zich mee, namelijk ‘responsstijlen’ (RS) (Welkenhuysen-Gybels et al., 2003). De Engelse definitie van dergelijke RS luidt als volgt: *RS are the respondent’s systematic tendency to respond to a range of survey items on a different basis from what the items are designed to measure*” (Paulhus, 1991 in Van Vaerenbergh & Thomas, 2012).

Dit probleem (RS) bevindt zich op het niveau van de respondenten. De vraag bij surveyonderzoek blijft dus of de onderzoekers er wel degelijk vanuit kunnen gaan dat de antwoorden van de respondenten hun daadwerkelijke mening over het bedoelde thema weergeven. Doorheen de jaren is er steeds meer onderzoek verschenen dat dit in twijfel trekt. Steeds meer is het geloof gegroeid dat er bepaalde RS zijn die de resultaten van onderzoeken verstoren. Een RS wordt dus gedefinieerd als de systematische neiging om surveyitems op een andere manier te beoordelen dan hoe ze bedoeld zijn om iets te meten (Paulhus, 1991 in Van Vaerenbergh, Thomas, 2012). De aanwezigheid van dergelijke RS in onderzoeken die gebruik maken van Likert schalen, kan dus een sterk effect hebben op de empirische resultaten van de onderzoeken. Toch wordt er in recentelijk onderzoek zelden rekening gehouden met de aanwezigheid van RS.

Van Vaerenbergh en Thomas (2012) definiëren in hun artikel acht verschillende soorten RS. Deze staan allemaal voor een andere systematische neiging bij het beantwoorden van surveyvragen met Likert schalen. In dit onderzoek zullen deze zeven RS besproken worden, maar de focus zal liggen op de *acquiescence response style* (ARS). Deze staat voor de neiging om het eens te zijn met alle surveyitems, los van de inhoud. Het is dus de neiging om ‘ja’ te knikken. ARS is één van de RS waar reeds het meeste onderzoek naar voltrokken is en die verondersteld wordt om erg vaak voor te komen (Johnson et al., 2005).

De brede puzzel is dus dat het meten van de publieke opinie nog steeds belangrijk is, en dat surveyonderzoek als methode in de lift zit, maar dat er toch zelden rekening wordt gehouden met RS die de resultaten van surveyonderzoek significant kunnen beïnvloeden. Desalniettemin heeft bestaand onderzoek reeds aangetoond dat er bepaalde statistische technieken zijn om RS in rekening te brengen of zelfs te corrigeren. In deze paper zullen twee van deze methodes uitgebreid besproken worden voordat ze gedemonstreerd worden. Ten eerste is er *Structural Equation Modelling* (SEM) en ten

¹ Bijvoorbeeld: *Helemaal Akkoord, Akkoord, Niet Akkoord maar ook niet Niet Akkoord, Niet Akkoord, Helemaal Niet Akkoord.*

tweede is er de RIRSMACS-methode. Deze zullen uitgebreid worden toegelicht na de theoretische uitleg over RS. Verder zal er antwoord komen op andere interessante vragen over de ARS. Bij welke soort respondenten is deze het duidelijkst terug te vinden? Aan de hand van de literatuur zullen we resultaten van vorige onderzoeken bekijken en zelf RS modelleren aan de hand van de uitgelegde statistische methodes. Hierna zullen we een duidelijker beeld scheppen over de ARS.

De dataset die gebruikt zal worden is de Barometer van het SCK-CEN (Studiecentrum voor Kernenergie – Centre d'études de l'énergie Nucléaire). De Barometer onderzoekt het veiligheidsgevoel tegenover nucleaire energie bij de Belgische bevolking. De studie werd voor het eerst uitgevoerd in 2002 en werd sindsdien periodiek herhaald. In deze paper zal er gewerkt worden met de Barometer van 2015². Deze vormt een ideale dataset voor onderzoek naar RS. Niet alleen bevat het een groot aantal respondenten, het bevat ook de nodige surveyitems voor de methodes die gehanteerd zullen worden. Alvorens deze methodes worden toegelicht volgt er uiteraard eerst een theoretische uiteenzetting die de lezer duidelijkheid moet brengen over RS, meer bepaald over de ARS. Niet alleen over hun wetenschappelijke achtergrond en indicators, maar ook over hun empirische implicaties.

1.2 Waarom is onderzoek naar responsstijlen nodig?

1.2.1 Maatschappelijke relevantie

Het belang van het meten van publieke opinie is zoals eerder aangehaald reeds vaak aangetoond als zijnde zeer belangrijk. Niet alleen in de politieke wereld, maar ook in de commerciële wereld wordt er veel belang gehecht aan accurate informatie over de publieke opinie. De interesse in de massaopinie is blijven groeien. Om hierover echter uitspraken te kunnen doen is een waardige representatieve sample essentieel. Kortom, surveyonderzoek is op dit moment één van de meest toegankelijke manieren om de publieke opinie te onderzoeken. De grootte van de sample is aantrekkelijk in politiek onderzoek, waar de mening van de massa vaak centraal staat. Een kwaliteitsvol surveyonderzoek mag niet onderschat worden qua kostprijs. Deze kostprijzen lopen vaak snel op als de survey bijvoorbeeld dient afgenomen te worden door interviewers, of als de lijst met surveyitems verschillende keren vertaald dient te worden.

Het feit dat surveyonderzoek zo populair is, roept natuurlijk vragen op over de betrouwbaarheid en de validiteit van de methode. Er is bijvoorbeeld het onderzoek van Schuman en Scott (1987) waarin zij de nadelen van surveyonderzoek onder de loep nemen. In de meeste boeken komen deze nadelen ook vaak aan bod en worden ze zeker benoemd. In de wetenschappelijke onderzoeken naar

² Website van het SCK: <https://www.sckcen.be/nl>

surveyonderzoek worden de nadelen dus wel benoemd en niet genegeerd. RS vormen één van die nadelen. En het klopt dat er inderdaad al behoorlijk wat onderzoek is verricht naar RS door mensen die bezig zijn met het theoretische concept van publieke opinie en het meten ervan (Schuman, Scott, 1987).

Het probleem is echter dat dit nadeel vaak genegeerd wordt in onderzoeken waar het surveyonderzoek als methode wordt gebruikt. Er wordt dan geen rekening gehouden met de mogelijke implicaties die deze RS zouden kunnen hebben op de resultaten van dat onderzoek. Dit zou aan meerdere factoren kunnen liggen. Aan de ene kant omdat de methodes om RS op te sporen te moeilijk zijn om toe te passen en aan de andere kant omdat er nog niet ‘genoeg’ empirisch bewijs is verschenen voor het bewijs van RS. Hoewel het doorheen deze paper duidelijk zal worden dat de nodige aandacht reeds besteed werd aan RS en hun implicaties

Deze paper zal dus naast de bespreking van RS en meer specifiek de ARS, ook een theoretische uiteenzetting voorzien van twee van de meest belangrijke statistische methodes om deze ARS op te sporen. In verschillende onderzoeken werden immers vele soorten methodes gebruikt. Van Vaerenbergh en Thomas (2012) geven een overzicht van alle verschillende methodes met hun voor- en hun nadelen. In deze paper bespreken we twee van deze methodes uitgebreid alvorens ze toe te passen op onze dataset, en empirische bewijzen van ARS te leveren. De bedoeling is om aan het einde van het onderzoek een discussie te kunnen voeren over de oorzaken van ARS en hoe deze RS opgespoord kan worden. Gezien de rijzende populariteit van surveyonderzoek en het gestegen aantal empirische bewijzen voor de aanwezigheid van ARS, is een onderzoek als dit zeker van toepassing. Het ‘niet rekening houden’ met ARS is een maatschappelijk probleem aangezien het de resultaten van zowel wetenschappelijk als commercieel onderzoek sterk kan beïnvloeden.

1.2.2 Wetenschappelijke relevantie

Hoewel het onderzoek naar RS nog niet bijzonder populair of bekend is, verschenen er reeds verschillende onderzoeken. Het onderzoek van Van Vaerenbergh en Thomas (2012) bijvoorbeeld geeft een overzicht van de zeven voornaamste RS en bespreekt deze alsook de manieren om ze op te sporen en te controleren. Ze wijzen in hun onderzoek op de mogelijke implicaties op empirische resultaten. Dit vinden ze een argument om het opsporen van RS extra op de kaart te zetten. Omdat dit vaak nog niet gebeurt door onderzoekers, willen ze duidelijkheid scheppen door (1) verschillende soorten RS te definiëren, (2) eerdere onderzoeken over RS te bespreken, (3) een overzicht voor te stellen van verschillende statistische bewerkingen die een oplossing kunnen bieden voor de effecten van RS. Hun artikel zorgt voor de nodige achtergrondinformatie over RS, die volgens hen nog niet aanwezig was in de literatuur (Van Vaerenbergh, Thomas, 2012). De RS die zij bespreken zijn: ARS (*acquiescence*

response style), MRS (*middlepoint response style*), ERS (*extreme response style*), MLRS (*mild response style*), NRS (*noncontingent response style*), DRS (*disacquiescence response style*) en NARS (*net acquiescence response style*).

Hun artikel is dus het eerste dat een compleet overzicht geeft over eerder onderzoek naar RS. Het onderzoek naar RS is immers al langer bezig. De eerste onderzoeken naar RS, allicht niet onder diezelfde naam, zijn afkomstig uit de jaren '60. Er werd onderzoek gedaan naar de stabiliteit en consistentie van de RS, maar ook naar mogelijke oorzaken van RS. Dit werd gedaan op twee niveaus. Aan de ene kant op het niveau van de respondent en aan de andere kant op het niveau van de survey. Op het niveau van de respondent werd er gekeken naar demografische variabelen (e.g. leeftijd, geslacht, opleidingsniveau etc.) om de RS te kunnen verklaren. Op het niveau van de survey werd er gezocht naar eventuele bias in de vragen en de opstelling van de surveyschalen, die als verklaring konden dienen. Deze onderzoeken, die begonnen in de jaren '60, komen uit verschillende disciplines maar vormden desalniettemin een goede basis voor latere onderzoeken (e.g. Das & Dutta, 1969; Hamilton, 1968; Crandall, 1982). In deze oudere onderzoeken werd er vooral gebruik gemaakt van descriptieve statistiek en telmethodes aangezien veel statistische methodes en software nog niet bestonden.

Eén van de eerste comparatieve onderzoeken over RS was dat van Stenning & Everett in 1984. Dit was een vergelijkende studie waarin onderzocht werd welke soort RS vaker voorkwamen bij welke nationaliteiten. De focus lag in deze studie vooral op het verschil tussen de neiging om extreem, en de neiging om zeer gemiddeld te antwoorden. Er werden ook effecten vastgesteld van leeftijd en opleidingsniveau op de aanwezigheid van bepaalde RS (Stenning, Everett, 1984). Harzing (2006) besteedde ook aandacht aan de cross-nationale verschillen bij de aanwezigheid van de ARS. Comparatief onderzoek werd ook verricht over de verschillende soorten data of de zelfs de verschillende manieren waarop de data gecollecteerd werden (e.g. Weijters et al., 2008).

Verder is er een nog een hele reeks goede onderzoeken verschenen, die meestal focussen op één soort RS. Zo is er nog een onderzoek van Greenleaf uit 1992. In dit artikel schrijft hij over het meten van de extreme respons stijl (ERS). Greenleaf beschrijft een methode om ERS-metingen te creëren en te valideren. Hij geeft een overzicht van verschillende, en vaak tegenstrijdige, bevindingen over de ERS en past zijn eigen methode toe om deze te valideren. (Greenleaf, 1992b). Het werk van Baumgartner en Steenkamp (2001) is ook van grote waarde voor het onderzoek naar RS, waaronder de ERS en de ARS. Dit is een onderzoek naar vijf belangrijke RS in elf landen. Het onderzoek vindt plaats bij consumenten en bevestigt de aanwezigheid van RS, en de afwezigheid van aandacht voor dit fenomeen, hoewel het de resultaten vaak significant beïnvloedt. Volgens dit onderzoek moet er meer aandacht zijn voor RS bij het construeren en het interpreteren van de meetschalen (Baumgartner, Steenkamp, 2001).

Ook over de methodologie verschenen er al verschillende onderzoeken. In verband met *Structural Equation Modelling* (SEM) publiceerden Billiet en McClendon reeds enkele onderzoeken over hoe ze de methode toepaste om RS te onderzoeken³. Ook Cheung en Chan (2002) publiceerden in de *Journal of Structural Equation Modelling* over hoe ze de methode in kwestie gebruikte om RS te modelleren. De tweede methode die we zullen bespreken, de RIRSMACS-methode is een soort verderzetting van SEM. De onderzoeken die deze methode hanteren zijn echter zeldzaam en dateren uit recente tijden. Vooral Weijters et al. zijn bekend voor het hanteren van deze RIRSMACS-methode. In dit onderzoek zal hun werk dan ook als een basis gebruikt worden, om er zelf op verder te kunnen bouwen (Weijters et al., 2009; Thomas et al., 2014). Er zijn dus verschillende soorten onderzoeken rond dit thema. Verschillende combinaties van RS worden onderzocht en verschillende methodes worden toegepast. Zoals eerder opgemerkt zal de focus in dit onderzoek liggen op de ARS. Na een algemene kadrering van de theorie achter RS, zal de ARS verder worden toegelicht. Deze RS werd reeds onderzocht, gebruikt en vermeld in verschillende onderzoeken⁴. Deze paper zal een uiteenzetting geven van alle resultaten over de ARS. Waarna deze kunnen worden nagegaan in een empirisch onderzoek aan de hand van de twee statistische methodes om de ARS op te sporen die worden uitgelegd in deze paper. Dit is niet alleen een aanvulling voor de onderzoeken naar de ARS maar ook een aanvulling bij de methodologische onderzoeken naar manieren om RS op te sporen.

1.2.3 Relevantie voor Politieke Communicatie

Het onderzoeksveld ‘politieke communicatie’ draait om de relatie tussen de politiek, de media en de massa. Laat de massa nu net dat ene zijn waarvoor het onderzoek naar publieke opinie nodig is. Om uitspraken te kunnen doen over de relatie met de massa, moet er eerst geweten zijn wat er leeft bij de massa, welke attitudes er spelen. Bij het meten van attitudes in surveyonderzoek wordt er vaak gebruik gemaakt van Likert schalen, en deze schalen zijn nu net extra vatbaar voor RS. Net om deze reden is onderzoek naar dergelijke RS van groot belang voor politieke communicatie. Het modelleren van RS is positief voor de validiteit van het meten van attitudes. Deze paper zal het onderzoek naar RS overlopen, repliceren en evalueren. Vooral de methodologie is belangrijk voor politieke onderzoekers, wanneer deze in de toekomst RS in rekening zouden brengen. Daarom ook worden in dit onderzoek SEM en de RIRSMACS-methode toegepast op een niet-politieke dataset, die de nodige elementen bevat. Op deze manier kan er na de toepassing een suggestie gemaakt worden over welke methode het meest-gepast is voor politieke wetenschappers, om hun onderzoek te beschermen van niet-gedetectede bias in de vorm van RS, die de resultaten kan beïnvloeden.

³ Zie Billiet en McCledon 1998 en 2000

⁴ Zie Baumgartner & Steenkamp (2001), Greenleaf (1992b), Billiet & McClendon (2008)

HOOFDSTUK 2 – De Theorie achter Responsstijlen

2.1 Wat is een Responsstijl?

Alvorens de theorie over RS verder toegelicht kan worden is het nuttig om de definitie te herhalen: *“RS are the respondent’s systematic tendency to respond to a range of survey items on a different basis from what the items are designed to measure”* (Paulhus, 1991 in Van Vaerenbergh & Thomas, 2012). Dit betekent dus dat respondenten de survey niet invullen naargelang hun voorkeuren over de surveyvragen, maar dat er andere factoren meespelen. Meer specifiek is er sprake van een ‘stijlfactor’ die erop wijst dat respondenten alle vragen in dezelfde stijl beantwoorden. Hun antwoorden staan op die manier los van de inhoud, maar zijn toch systematisch. Op die manier kunnen antwoorden bijvoorbeeld zeer gematigd zijn of juist zeer extreem zijn. En zo bestaan er vele soorten RS die allemaal verder toegelicht zullen worden. Deze RS kunnen de uitkomst van een onderzoek beïnvloeden op verschillende manieren. Ten eerste kan de aanwezigheid van een RS de verdeling van een variabele beïnvloeden. Zowel het gemiddelde als de variantie (= spreiding) kan vertekend worden door de aanwezigheid van RS, waardoor de onderzoeker zijn resultaat foutief kan interpreteren. RS kunnen ook de correlatie tussen verschillende variabelen beïnvloeden. Correlaties kunnen zowel versterkt als verzwakt worden. Aangezien correlaties aan de basis liggen van de meeste statistische bewerkingen, kan dit dus grote gevolgen hebben (Van Vaerenbergh, Thomas, 2012). Doordat RS zulke gevolgen kunnen hebben voor wetenschappelijke resultaten, is het onderzoek ernaar belangrijker dan momenteel wordt ingeschat.

De gevolgen zijn significant en eisen dus de aandacht op van onderzoekers. Aan de andere kant blijft de vraag echter wat de oorzaken van deze RS zijn. Ook hierover is reeds een beduidende hoeveelheid onderzoek verschenen. Het onderzoek naar de oorzaken van RS wordt opgedeeld in twee soorten. Aan de ene kant zijn er de oorzaken op het niveau van de stimuli, de survey zelf. En aan de andere kant zijn er de oorzaken op het niveau van de respondent. Dit laatste zal uitgebreid besproken worden met betrekking tot ARS, en het onderzoek naar de oorzaken van ARS bij de karakteristieken van respondenten. Op het niveau van de survey zijn er echter ook oorzaken die van belang zijn. Zo is er bijvoorbeeld het formaat van de schaal. Er zijn verschillende onderzoeken die aangeven dat RS kunnen ontstaan door verschillen in schaal (Greenleaf, 1992a). Deze verschillen betreffen meestal de grootte. Zo kan het zijn dat een respondent eerder geneigd is om extreem te antwoorden op items met een kleinere schaal dan op items met een grotere schaal. Een andere mogelijke oorzaak, op surveyniveau, is de manier van datacollectie. Onderzoek heeft al aangetoond dat surveys die afgenomen werden via telefoon vaker zouden leiden tot ARS en ERS dan surveys die face-to-face of via papier werden afgenomen. Deze verschillen leiden tot de conclusie dat er rekening moet worden

gehouden met RS bij het gebruik van mixed data (Weijters, Geuens, Schillewaert, 2008). Een laatste relevante oorzaak is de mate van betrokkenheid bij het thema. Hoewel dit lijkt op een eventuele eigenschap van een respondent, verwijst dit eerder naar de relevantie van de vragen. Gezien er eerder een RS aanwezig zal zijn bij vragen die beschouwd worden als zijnde niet relevant. Onderzoek wees bijvoorbeeld uit dat het gebruik van ERS vaker voorkomt bij respondenten die sterk betrokken zijn bij het thema (Greenleaf, 1992b). Dit zijn dus maar enkele van de belangrijkste voorbeelden van de oorzaken van RS (Van Vaerenbergh, Thomas, 2012). Voordat het onderzoek zich volledig zal richten op ARS, is het belangrijk om de andere soorten RS te overlopen en toe te lichten. Er zijn doorheen de jaren namelijk wel enkele verschillende RS geïdentificeerd die allen slaan op een andere systematische manier waarop mensen (onbewust) surveyvragen beantwoorden. Er volgt een overzicht van de verschillende soorten RS.

2.2 Welke soorten responsstijlen bestaan er?

Middelpunt Respons Stijl (MRS): Deze RS slaat op de neiging om bij surveyitems met een ranking-schaal, steeds de middelste categorieën aan te duiden. De antwoorden van een persoon die volgens deze RS antwoordt zijn hierdoor allemaal zeer gematigd (Van Vaerenbergh, Thomas, 2012). Baumgartner en Steenkamp (2001) vonden bewijs voor MRS door middel van het tellen van gematigde antwoorden en de verhouding te berekenen tegenover alle antwoorden. Hun theoretische uitleg voor deze RS is dat ze vaak toegepast wordt door mensen die hun echte mening over de surveyitems liever niet willen weergeven. Het wijst dus op het ontwijken van vragen of het ontwijken van een mening te moeten hebben over bepaalde items. Ook in andere onderzoeken wordt deze RS vooral gelinkt aan ontwijkend gedrag (Weijters, Geuens, Schillewaert, 2008).

Extreme Respons Stijl (ERS): ERS wijst op de neiging om vaak de hoogste of de laagste antwoorden in een surveyonderzoek aan te duiden. Het staat recht tegenover de MRS (Van Vaerenbergh, Thomas, 2012). Greenleaf (1992b) schreef een artikel om consensus te zoeken over het meten van deze ERS, waar veel discussie over bestond. Om deze RS te kunnen meten is het namelijk wenselijk dat de extreme respons proporties gelijk zijn bij de verschillende vragen. Ook over deze RS is reeds veel onderzoek verschenen. Er zijn vier mogelijke theoretische verklaringen voor ERS. Ten eerste zou het een reflectie kunnen zijn van onverdraagzaamheid, stijfheid, dubbelzinnigheid en kortzichtigheid. ERS wordt ook geassocieerd met nerveusheid en afwijkend gedrag. De RS kan ook verklaard worden door karakteristieken van de respondent. Zo zou de respondent minder ontwikkelde cognitieve structuren bezitten wat zou kunnen leiden tot ERS. Een laatste mogelijke verklaring die wordt gegeven voor ERS is de aanwezigheid van uitzonderlijk belangrijke stimuli. Wanneer de inhoud van de survey van uitzonderlijk groot belang is voor de respondent zou dit kunnen leiden tot ERS (Greenleaf, 1992b).

Milde Respons Stijl (MLRS): Deze RS is niet te verwarren met MRS, en betekent dat het allerlaagste en het allerhoogste antwoord te allen tijde vermeden wordt door de respondent. MLRS is dus complementair met ERS, waar extreme antwoorden juist verkozen worden (Van Vaerenbergh, Thomas, 2012). Naar deze RS is er minder onderzoek gevoerd maar toch komt deze af en toe aan bod in onderzoek naar andere RS. De voorkeur voor onderzoek gaat echter duidelijk naar ERS, omdat de implicaties van deze RS duidelijker op te merken zijn.

Noncontingent Respons Stijl (NRS): Wanneer deze RS wordt vastgesteld betekent het dat de respondent de neiging had om de vragen/item willekeurig en zonder nadenken te beantwoorden (Van Vaerenbergh, Thomas, 2012). Het onderzoek naar deze NRS is beperkt, gezien de mogelijke gevolgen van deze theorie moeilijk in te schatten zijn. In zekere zin is het een non-RS, wat onderzoek ernaar ook minder interessant maakt. De effecten worden ook geacht minder significant te zijn. Willekeurige antwoorden zouden elkaar in zekere zin toch moeten opheffen in plaats van extra implicaties voor de resultaten te hebben. Deze RS vormt een interessante factor bij het meten van non-attitudes.

Disacquiescence Respons Stijl (DRS): Deze RS wijst op de neiging om het niet eens te zijn met surveyitems uit de vragenlijst. Los van de inhoud kiest de respondent ervoor om steeds zo min mogelijk akkoord te gaan met de voorgeschotelde survey items (Van Vaerenbergh, Thomas, 2012). DRS kan het gevolg zijn van een respondent met een introverte en bedachtzame persoonlijkheid die stimuli van buitenaf wil vermijden (Weijters, Geuens, Schillewaert, 2008). Onderzoek naar deze RS wordt vaak gecombineerd met onderzoek naar ARS. Dit onderzoek zal echter volledig focussen op ARS en niet op DRS.

Net Acquiescence Respons Stijl (NARS) is een samenvoeging van DRS en ARS waarin deze worden vergeleken en de assumptie wordt gemaakt of men al dan niet de neiging heeft om meer *acquiescence* te tonen dan *disacquiescence*.

2.3 Wat is ARS?

In de onderzoekswereld is er van alle RS al het meeste aandacht besteed aan ARS en zijn tegenhanger DRS (Harzing, 2006). Deze RS wijzen op het systematisch akkoord of niet akkoord gaan met de items uit een survey. De populariteit van ARS zorgt ervoor dat er veel onderzoekers het reeds opgespoord hebben, en getracht hebben om claims te maken over de oorzaken. Er zijn dus veel verschillende bevindingen over ARS en dat zorgt ervoor dat het een interessante focus is voor verder onderzoek. Dit onderzoek zal de claims over ARS uitéénzetten en uitleggen. Waarna de hypothesen die gevormd

kunnen worden uit deze bestaande literatuur, getest zullen worden. Het testen zal, in tegenstelling tot veel van de oudere onderzoeken, gebeuren door middel van geavanceerde statistische bewerkingen. Waar veel van de eerdere onderzoeken een tel- of verhoudingsmethode hanteerde, zal dit onderzoek gebruik maken van *Structural Equation Modelling* (SEM) en de RIRSMACS-methode. Maar eerst is er meer kennis over de ARS nodig. De A in ARS staat uiteraard voor *acquiescence*, wat letterlijk ‘instemmen’ betekent. ARS is dus de neiging om het eens te zijn met alle items/ vragen uit de survey (Van Vaerenbergh, Thomas, 2012).

Zoals reeds vermeld worden attitudes meestal gemeten door middel van items die beoordeeld moeten worden via een Likert schaal. Het probleem hiermee is dat deze schalen erg vatbaar zijn voor ARS, wat bias veroorzaakt. Het gebruik van dit soort surveys blijft echter zeer populair, vandaar de nood om methodes te verkennen die ARS kunnen opsporen en corrigeren. Men gebruikt dus erg vaak surveyonderzoeken met Likert schalen omdat deze bestaan uit meerdere indicatoren die verschillende attitudes kunnen aangeven (Billiet, McClendon, 1998, pg. 130). Toen het duidelijk werd dat dit soort van surveyonderzoeken waarschijnlijk gepaard gingen met RS, werd er gezocht naar oplossingen. Er waren verschillende onderzoekers die voorstelden om het gebruik van dergelijke schalen bij survey items te bannen. In de plaats hiervan zou er dan gebruik kunnen gemaakt worden van *forced-choice items*. Bij dit soort surveyitems moeten de respondenten items rangschikken in plaats van aangeven in hoeverre ze het eens zijn met een item (Schuman, Presser, 1981). De meningen hierover waren echter verdeeld. Ray (1990) vond dat deze methode verschillende problemen over het hoofd zag en stelde daarom voor om gebruik te maken van gebalanceerde schalen. Dit betekent dat een deel van de vragen positief, en een deel van de vragen negatief geformuleerd wordt in de survey. Op deze manier kan de *acquiescence* die aanwezig is bij de positieve schalen, de *acquiescence* die aanwezig is bij de negatieve schalen opheffen. Zo corrigeert het probleem eigenlijk zichzelf (Ray, 1979). Hier wordt er wel vanuit gegaan dat positief en negatief verwoorde items evenveel kans hebben op ARS. In de praktijk is echter bewezen dat dit niet noodzakelijk waar is (McClendon, 1992).

Niet iedereen is echter overtuigd van het bestaan van ARS. Rorer (1965) bracht een onderzoek uit waarin hij ARS een mythe noemde, die nooit bewezen was. Hij kwam tot deze conclusie omdat er volgens hem geen doorslaggevend bewijs was voor RS. Volgens Rorer correleren de ontdekkingen van RS niet over de verschillende onderzoeken heen. Rorer wijt dit aan het feit dat RS in alle testen anders geconceptualiseerd werden, waardoor er geen eenduidig concept van ARS bestaat (Billiet, McClendon, 1998). Dit nieuws kwam natuurlijk gelegen voor vele onderzoekers aangezien het de validiteit van surveyonderzoek in zekere zin herstelde. De steun voor Rorer's claim bleef echter beperkt. Het onderzoek naar RS ging door en al gauw kwamen er nieuwe resultaten aan het licht die het bestaan van ARS bevestigde (Rorer, 1965).

2.4 Oorzaken van ARS

Het onderzoek van Billiet en McClendon (2009) onderzoekt de aanwezigheid van ARS voor items die gebalanceerd zijn. Met ‘gebalanceerd’ wordt er bedoeld dat er zowel positief geformuleerde als negatief geformuleerde items aanwezig zijn in de dataset. Zo wordt er een *least likely case*⁵ gecreëerd, waarbij het alsnog aantonen van ARS duidelijk significant is. Er wordt gebruik gemaakt van *Structural Equation Modelling (SEM)* om ARS op te sporen (Billiet, McClendon, 2009). Ook met behulp van simpelere methodes werd er al ARS vastgesteld in verschillende datasets. In het comparatief onderzoek van Anne-Wil Harzing werd er gebruik gemaakt van een telmethode, waar de proportie van het aantal vragen waarmee akkoord gegaan werd, vergeleken werd met het totale aantal vragen. Dit comparatief onderzoek vergelijkt de aanwezigheid van ARS in verschillende landen. Waarna er geconcludeerd kan worden in welke landen er meer of minder ARS aanwezig is bij de bevolking (Harzing, 2006). Het onderzoek van Billiet en McClendon (2000) was echter niet comparatief, maar bevestigde de aanwezigheid van ARS in een heterogene algemene populatie. De dataset in dit onderzoek zal een sample zijn uit de Belgische bevolking. Daarom kunnen we de volgende hypothese opstellen:

H1: Acquiescence is vast te stellen in een dataset met losstaande theoretische concepten, die is verkregen door een bevraging van een sample van de heterogene Belgische bevolking.

2.4.1 Het niveau van de stimuli

Een deel van de onderzoekers die werken rond RS, claimt dat deze RS beter beschouwd kunnen worden als een persoonlijkheidskenmerk. Jackson en Messick (1962) stellen dat een RS een mogelijk te meten persoonlijkheidskenmerk is. Zij zien RS als systemen in het gedrag die gebruikt worden bij het beoordelen van allerlei soorten items die inhoudelijk los staan van elkaar (Billiet McClendon, 2009). Toch zijn er ook onderzoekers die ARS niet beschouwen als een persoonlijkheidskenmerk. Ray (1983) spreekt niet alleen het ongeloof van Rorer in ARS tegen, hij maakt ook een nieuwe claim over ARS. Volgens hem is de grootste reden voor de aanwezigheid van ARS, dubbelzinnigheid in de items van een survey. De claim is dus dat er meer gebruik van ARS zal zijn wanneer de survey-items dubbelzinnig of moeilijk geformuleerd zijn, dan wanneer de items duidelijk en helder geformuleerd worden (Ray, 1983; Billiet, McClendon, 2009). Er is nog één ander onderzoek dat dit bevestigt. McBride en Moran (1967) wezen in 1967 al op de gevolgen van vragen die op een dubbelzinnige manier geformuleerd worden. Onze tweede hypothese luidt dus als volgt:

H2: Er zal meer ARS aanwezig zijn bij sets van items die op een dubbelzinnige, dus niet-duidelijke manier, geformuleerd worden.

⁵ Een scenario waarin het weinig waarschijnlijk lijkt dat men ARS zal aantreffen. In dit geval omdat de vragen zodanig geformuleerd zijn dat het respondenten zou moeten stimuleren om niet met alles akkoord te gaan.

2.4.2 Het niveau van de respondent

Zoals reeds eerder vermeld kunnen de oorzaken van ARS gezocht worden op twee niveaus. De vorige hypothese gaat dan eerder over het niveau van de stimuli in de dataset. Hoewel ook dat niveau reeds wetenschappelijke aandacht ontving, is er al meer onderzoek gedaan naar het andere niveau. Namelijk het niveau van de respondent. De vraag is dus bij welke soort respondenten er meer ARS valt te detecteren en bij welke er minder ARS is vast te stellen. Verschillende onderzoeken deden hier reeds onderzoek naar. Niet alleen over ARS, maar ook over andere RS. Desalniettemin zijn er over de ARS al vele resultaten verschenen. Deze resultaten waren niet geheel geharmoniseerd, er zaten tegenstrijdige uitkomsten in. Deze kunnen we overlopen om ze vervolgens zelf te kunnen testen in onze dataset. We overlopen de meest voor de hand liggende demografische variabelen.

De eerste demografische variabele die we gaan bespreken is **gender**. Deze oude natuurlijke breuklijn vormt al sinds jaar en dag een punt van onderzoek. Ook wat betreft RS hebben onderzoekers zich dus reeds de vraag gesteld of RS misschien vaker voorkomen bij mannen dan bij vrouwen of omgekeerd. In het onderzoek van Austin et al. kwam via de *Rash methode*⁶ aan het licht dat er bij vrouwen meer ARS viel vast te stellen dan bij mannen. Deze resultaten waren in lijn met de resultaten van Eid en Rauber (2000) die dezelfde conclusie trokken over vrouwen (Austin et al., 2006). In ander onderzoek werd ook vastgesteld dat er bij vrouwelijke respondenten zowel meer ARS als DRS kon worden vastgesteld (Weijters et al., 2010c). Dit is dus een interessant gegeven om na te gaan in dit onderzoek.

H3: Er is meer ARS aanwezig bij vrouwelijke respondenten dan bij mannelijke respondenten.

Een andere demografische variabele die reeds onderzocht werd is de **leeftijd** van de respondenten. Het lijkt logisch te veronderstellen dat jongeren en ouderen op een andere manier vatbaar zijn voor ARS, gezien de verschillen in levenservaring. Waar iemand misschien zou denken dat jongere mensen vaker akkoord zouden gaan met surveyitems door minder levenservaring of zelfzekerheid, bleek juist het tegendeel. In enkele onderzoeken kwam duidelijk naar boven dat er een positieve relatie bestaat tussen iemands leeftijd en de aanwezigheid van ARS (Billiet, McClendon, 2000; Greenleaf, 1992a; Mirowsky, Ross, 1991). Dit betekent dat er meer ARS kan worden vastgesteld bij respondenten met een hogere leeftijd dan bij respondenten met een lagere leeftijd. Er is echter ook onderzoek waar er geen effect wordt gerapporteerd over leeftijd op de aanwezigheid van ARS (Eid, Rauber, 2000). Om na te gaan wat er klopt in onze dataset zullen we de relatie tussen ARS en leeftijd onderzoeken.

H4: Er is een significante relatie tussen ARS en leeftijd, in de zin dat ARS vaker voorkomt bij oudere respondenten dan bij jongere respondenten.

⁶ Een specifieke statische methode die wordt gebruikt om RS op te sporen.

Het gebruik van een RS wordt niet altijd gezien als een toevalligheid, daarom onderzoeken we ook de mogelijke eigenschappen van respondenten. Zo is er ook de demografische variabele **opleidingsniveau**. Tijdens de schoolcarrière van kinderen wordt hen aangeleerd om kritisch te zijn, en na te denken over dingen zonder er zo maar akkoord mee te gaan. Daaruit zou de veronderstelling kunnen komen dat hoe langer of beter de opleiding is die iemand genoten heeft, invloed zou kunnen hebben op die respondent zijn vatbaarheid voor responsstijlen. Verschillende onderzoeken bevestigen deze redenering en stellen dat deze negatieve relatie een wereldwijd fenomeen is (Greenleaf, 1992a; Meisenberg, Williams, 2008; Mirowsky, Ross, 1984). Een tegenstellend resultaat tot deze redenering is nog niet naar boven gekomen in onderzoek en zou dan ook een hele verrassing zijn. Het lijkt duidelijk dat we, ook in onze dataset, meer ARS zullen aantreffen bij de respondenten met lagere opleidingsniveau.

H5a: Er is een negatieve relatie tussen de aanwezigheid van ARS en het opleidingsniveau van de respondenten. Hoe hoger het opleidingsniveau hoe lager de aanwezigheid van ARS.

Een andere ontdekking die toch ook werd gedaan met betrekking tot deze demografische variabelen was dat er een verband is tussen de aanwezigheid van ARS en de duur van tewerkstelling van een respondent, los van opleidingsniveau (Johnson et al., 2005). Zonder dat dit H5a echt tegenspreekt, ondermijnt het deze wel. Daarom is het extra interessant om H5a te testen en na te gaan of dit wel degelijk duidelijk valt vast te stellen. In deze zelfde categorie, ligt het **inkomen** van de respondenten. Deze demografische variabele is, niet noodzakelijk, maar toch meestal verbonden aan het opleidingsniveau. Ook in eerder onderzoek kwam dit reeds naar boven. ARS is minder vaak aanwezig naar gelang het inkomen van de respondenten hoger is (Meisenberg, Williams, 2008; Ross, Mirowsky, 1984). Deze resultaten liggen dus in dezelfde lijn als de resultaten over opleidingsniveau. Wat uiteraard de verwachting was. Vandaar ook onze volgende hypothese:

H5b: Er is een negatieve relatie tussen de aanwezigheid van ARS en het inkomen van respondenten. Dit betekent dat hoe meer een respondent verdient, hoe kleiner de kans is dat er ARS is vast te stellen.

Dit zijn dus enkele van de mogelijke persoonlijkheidskenmerken die bijdragen tot het verklaren van de variantie van RS. Deze moeten echter in context gezien worden, en het moet gezegd worden dat er andere variabelen zijn die meer van die variantie kunnen verklaren. Zo wordt bijvoorbeeld beweerd dat cross-nationaal en cross-cultureel onderzoek, meer potentieel heeft om variantie van bijvoorbeeld ARS te verklaren. Onze dataset is echter een sample van de Belgische bevolking dus een cross-nationaal onderzoek is onmogelijk. Wat echter wel binnen de mogelijkheden van dit onderzoek ligt is een cross-cultureel onderzoek. De Belgische bevolking bestaat uiteraard uit enkele subculturen, waarvan de bekendste de breuklijn tussen Walen en Vlamingen is. Er is geen precedent voor deze stelling, dus geen literatuur met directe informatie hierover om een hypothese uit te vormen. Op het vlak van cross-nationaal onderzoek is er echter wel reeds een hoop literatuur verschenen.

Om echter toch een idee te krijgen over een mogelijk verschil is het interessant om naar gelijkaardige groepen te kijken. Er wordt vaak gezegd dat Vlamingen meer neigen naar Nederlanders en Walen naar de Fransen. Deze stelling is afkomstig uit de geschiedenis, aangezien België ook zo verdeeld was voor de eenmaking in 1830. In het cross-nationaal onderzoek van Anne-Wil Harzing, werd de aanwezigheid van ARS gemeten bij respondenten uit verschillende landen. Zo kon er een cross-nationale vergelijking gemaakt worden⁷. Frankrijk scoorde een percentage van 54% op de aanwezigheid van de ARS, wat aan de hoge kant was. Dit percentage stond voor het aantal vragen waar akkoord mee gegaan werd door de respondent. Nederland daarentegen scoorde 48%, en deed dus ‘beter’ in de zin dat er minder ARS werd vastgesteld. Om deze cijfers in context te kunnen plaatsen is het nuttig om te weten dat bij de Deense respondenten het minste ARS werd vastgesteld (44%) en bij de Maleisische respondenten het meeste ARS, maar liefst 61% (Harzing, 2006). Dit doet ons vermoeden dat er meer ARS aanwezig zal zijn bij de Waalse respondenten dan bij de Vlaamse respondenten. Een andere reden om dit te vermoeden is het onderzoek naar het verschil tussen de aanwezigheid van RS in rurale gebieden en in stedelijke gebieden (Thomas et al., 2014). Met 212,32 inwoners/km² is Wallonië aanzienlijk minder dichtbevolkt dan Vlaanderen met zijn 554, 42 inwoners/km². Dit zou geïnterpreteerd kunnen worden op de manier dat er in Wallonië een meer rurale sfeer heerst dan in Vlaanderen. En het is reeds aangetoond dat er in stedelijk gebied minder ARS voorkomt (Thomas et al., 2014). Daarom stellen we dat:

H6a: Er meer ARS aanwezig zal zijn bij Waalse respondenten dan bij Vlaamse respondenten.

Er is echter ook een andere redenering die gevolgd kan worden omtrent deze hypothese. In Vlaanderen heerst er immers al jaren een Christelijke traditie. Ondanks de trends van de laatste jaren blijft de Christelijke-Democratische en Vlaamse partij (CD&V) stevast vertegenwoordigd in het Vlaamse politieke landschap. Deze christelijke traditie gaat stereotypisch gezien gepaard met een zekere volgzaamheid, die zou kunnen wijzen op de aanwezigheid van ARS. In Wallonië heerst deze christelijke traditie veel minder. Daar heerst er eerder een socialistische traditie met de Parti Socialiste (PS) die nog steeds de grootste is. Deze traditie is afkomstig uit het feit dat het grootste deel van de Waalse bevolking arbeiders zijn/waren of afstammen uit dergelijke families. Deze redenering zorgt er dus voor dat we een tegengestelde hypothese kunnen vormen voor H6a, we kunnen namelijk ook vermoeden dat:

H6b: Er meer ARS aanwezig zal zijn bij Vlaamse respondenten dan bij Waalse respondenten.

Voor de hypothesen daadwerkelijk getest kunnen worden, zullen de verschillende methodes om ARS te meten toegelicht worden. Het moet dan ook gezegd worden dat niet alle hypothesen getest kunnen worden met dezelfde methodes. Zo maakt de RIRSMACS-methode bijvoorbeeld gebruik van niet-

⁷ Het aantal respondenten per nationaliteit was echter beperkt, dus de conclusies moeten genuanceerd worden.

gecorrleerde items terwijl men bij SEM gecorreleerde items gebruikt bij het modelleren van ARS. Daarom worden de methodes nu verder toegelicht.

Tabel 2.4.2.a - Overzicht hypotheses:

H1	Er is ARS vast te stellen in een sample van de heterogene Belgische bevolking.
H2	Er valt meer ARS vast te stellen bij dubbelzinnig-geformuleerde items dan bij duidelijk geformuleerde items.
H3	Er is meer ARS aanwezig bij vrouwelijke respondenten dan bij mannelijke respondenten.
H4	Er is een positieve relatie tussen de aanwezigheid van ARS en de leeftijd van een respondent.
H5a	Er is een negatieve relatie tussen de aanwezigheid van ARS en het opleidingsniveau van de respondent.
H5b	Er is een negatieve relatie tussen de aanwezigheid van ARS en het inkomen van een respondent.
H6a	Er is meer ARS aanwezig bij Waalse respondenten dan bij Vlaamse respondenten.
H6b	Er is meer ARS aanwezig bij Vlaamse respondenten dan bij Waalse respondenten.

HOOFDSTUK 3 – *Structural Equation Modelling* (SEM)

3.1 Het modelleren van ARS

Theoretisch gezien is er dus al heel wat onderzoek verschenen over de verschillende RS. Natuurlijk moet dit fenomeen, om het bestaan ervan te bewijzen, ook in de praktijk reeds zijn aangetoond. Er worden verschillende methodes gehanteerd om ARS aan te tonen of zelfs om het te corrigeren in bepaalde onderzoeken. Van Vaerenbergh en Thomas (2012) geven een overzicht van deze verschillende methodes, en overlopen de voor- en nadelen ervan. Doorheen de tijd zijn de technieken uiteraard een heel stuk geëvolueerd. Waar in het begin eerder basis teltechnieken gebruikt werden, wordt ARS nu meestal gemodelleerd en wordt er gebruik gemaakt van *confirmatory statistics*. Dit is ook te wijten aan de evolutie in software, die nodig is om bepaalde methodes toe te passen.

De eerste methodes die werden toegepast waren meestal telmethodes. Dit houdt in dat, voor het geval van ARS, het aantal keer dat een respondent het eens was met de survey items geteld werd. Op die manier werd er dan een beeld gevormd over de hoeveelheid ARS. In het onderzoek van Harzing (2016) wordt de getelde ARS in verhouding gezet tot het totaal aantal items. Zo kan over, in dat geval verschillende nationaliteiten, een beeld gevormd worden over de hoeveelheid ARS. Deze gemakkelijke methode kan ook gebruikt worden bij surveyonderzoeken met gebalanceerde schalen⁸ (Johnson et al., 2005). De methode is gebruiksvriendelijk aangezien er geen extra tools of software nodig zijn. Het nadeel is dat wanneer men de telmethode wil gebruiken, er geen rekening mee gehouden kan worden dat het ‘akkoord gaan’ van de respondent niet gewoon zijn/haar daadwerkelijke mening demonstreert. En als men de methode wil toepassen op gebalanceerde items, is er de uitdaging om de vragen, in tegengestelde richting, duidelijk te kunnen formuleren. Het gebruikmaken van gebalanceerde schalen werd gezien als een methode om ARS te vermijden. Zoals reeds gezegd wordt er op die manier een *least likely case* gecreëerd waardoor de aanwezigheid van ARS niet waarschijnlijk lijkt. Door de moeilijkheden met het ‘niet-dubbelzinnig’ formuleren van dergelijke vragen is er vaak een grote kans dat er ‘interpretatie-factoren’ meespelen bij het beoordelen van de surveyitems (Van Vaerenbergh, Thomas, 2012).

Statistisch gezien betekent de aanwezigheid van RS, dat bepaalde variabelen, die inhoudelijk los staan van elkaar, variantie delen (Paulhus, 1991). De variantie staat voor de spreiding van de antwoorden op een bepaald surveyitem. Als bepaalde variabelen variantie gemeen hebben, terwijl ze inhoudelijk los staan van elkaar betekent dit dat de variabelen (surveyitems) toch systematisch beantwoord zijn. En aangezien er inhoudelijk geen verband is, slaat dat systeem dan op een stijlfactor. Deze stijlfactor kan

⁸ Evenveel negatief als positief geformuleerde items in de survey

verschillende stijlen inhouden, zoals bijvoorbeeld RS. Bij ARS zou de gedeelde variantie dan slaan op het akkoord gaan met de surveyitems. Een deel van de statistische uitdaging is om deze gedeelde variantie te kunnen verklaren. Bij het ‘verklaren van gedeelde variantie’ denken we automatisch aan een factoranalyse, omdat deze deductiemethode één factor haalt uit verschillende variabelen die allemaal variantie delen. Bij *Exploratory Factor Analysis* is het de bedoeling om een onderliggende structuur te vinden in de variabelen, een latent construct waarvan men niet wist dat het aanwezig was (Hair et al., 2010). Wanneer er getracht wordt om RS te vinden, is er dus weet van een onderliggende structuur en is het doel om deze te kunnen bewijzen. In dat geval wordt er gebruikgemaakt van *Confirmatory Factor Analysis* (CFA). Zoals er bij een verkennende factoranalyse gezocht wordt naar een latent construct en dus een onderliggende structuur, wordt hier bij CFA vanuit gegaan. Om ARS te vinden wordt er dus uit gegaan van een voorbeeldpatroon (akkoord gaan met elke vraag) dat dan toegepast kan worden op de data. Zo kan gezien worden of het model al dan niet past, en of de RS dus van toepassing is (Hair et al., 2010). Om te zien of het model ‘past’, wordt er dikwijls gekeken naar de *Root Mean Square Error of Approximation* (RMSEA). Deze indicator voor de ‘fit’ van een model vermijdt problemen met de sample door deze te vergelijken met populatie matrix. Een andere indicator die gebruikt wordt om de fit van een model aan te geven is de Chi-kwadraattoets. Deze wordt gebruikt om verdelingen met elkaar te vergelijken. De verdeling die verwacht wordt met het model, wordt vergeleken met de verdeling die effectief wordt waargenomen. Door de kwadraten van de verschillen tussen deze waarnemingen te toetsen wordt er dan gekeken hoe goed een verwachting op de realiteit past. In dit geval dus hoe goed het model past. In dit onderzoek worden de zelfgecreëerde modellen echter beoordeeld volgen de RMSEA en niet volgens de Chi-kwadraattoets. Deze lijkt geschikter aangezien de Chi-kwadraattoets niet geschikt is voor grote samples. Deze zal plots veel groter zijn en zo vaker onterecht een slechte fit aangeven (Hair et al., 2010).

In dit onderzoek zal er echter verder worden gegaan dan testen of ARS al dan niet aanwezig is. Er zal onderzocht worden wat de relatie van ARS tot bepaalde demografische variabelen inhoudt. De latente constructen die gecreëerd worden zullen dus in relatie tot elkaar en tot andere variabelen gebracht worden. Om dit te kunnen uitvoeren wordt er gebruik gemaakt van *Structural Equation Modelling* (SEM) waarvoor CFA eigenlijk de basis vormt. Dit zal dan ook de eerste methode zijn, die we uitgebreid zullen bespreken en later zullen toepassen op de data.

3.2 *Structural Equation Modelling* (SEM)

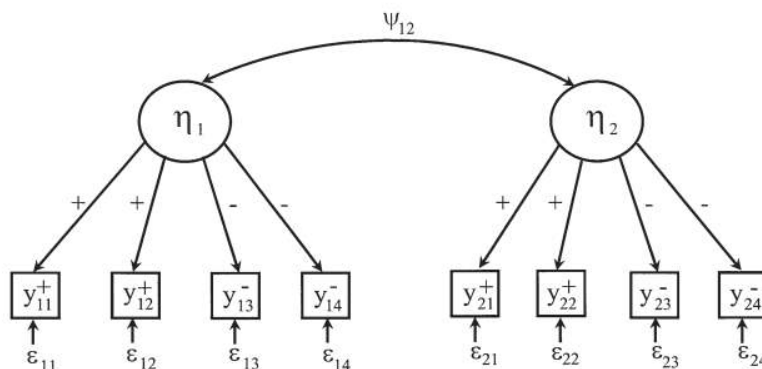
SEM wordt gebruikt om bestaande theorieën te testen op data. Het behoort dus tot *confirmatory statistics* en combineert twee veelgebruikte statistische bewerkingen. Aan de ene kant factoranalyse, om latente constructen te creëren. En aan de andere kant een regressieanalyse, die gebruikt wordt om

de relaties tussen de latente variabelen te onderzoeken (Hair et al., 2010). SEM zorgt ervoor dat deze bewerkingen niet afzonderlijk uitgevoerd moeten worden, maar gecombineerd kunnen worden. Het gebruik van SEM heeft bepaalde voordelen. Zo kunnen er meerdere afhankelijke variabelen betrokken worden in een onderzoek, daar ze samen gemodelleerd kunnen worden. Er wordt vertrokken uit een theorie en worden verschillende variabelen en latente constructen in een structuur geplaatst zodat het model klopt en de RMSEA laag genoeg is. Het is algemeen aanvaard dat een $RMSEA < 0.05$ goed is, en een $RMSEA < 0.08$ acceptabel is als fit (McDonald, Ringo Ho, 2002). Het gebruik van latente constructen heeft als voordeel dat het makkelijker is om een theorie toe te passen maar ook omdat het de zogenaamde *measurement error* in het model betreft. Bij andere methodes wordt deze vaak genegeerd, maar het moet gezegd worden dat geen enkel concept perfect gemeten kan worden en dat er altijd sprake is van een zekere foutmarge bij het meten van de concepten (Hair et al., 2010). Bij telmethodes bijvoorbeeld worden deze foutmarges niet in rekening gebracht. Er wordt in dat geval geen rekening mee gehouden of die neiging om akkoord te gaan met de surveyitems al dan niet iets te maken heeft met de inhoud. SEM is dus een geavanceerdere methode van de laatste jaren. Onderzoekers zagen ook het nut in van deze voordelen en begonnen SEM te gebruiken om RS vast te stellen. Dit geldt ook voor ARS.

3.3 Het onderzoek van Billiet en McClendon (2000)

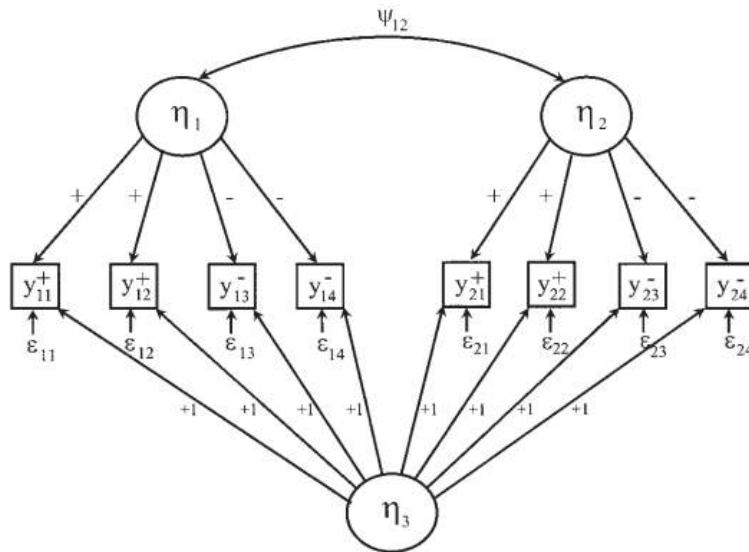
Het onderzoek van Billiet en McClendon (2000) is een goed voorbeeld hoe SEM kan dienen om RS in kaart te brengen. Zoals dat gaat met *confirmatory statistics* ontwikkelden zij eerst theoretische modellen die ze later zullen testen op de dataset. Billiet en McClendon ontwikkelen drie hypothetische modellen die allemaal op de dataset zouden kunnen passen. Het eerste model (figuur 3.3.a) is een simpel model waarin er twee inhoudsfactoren aanwezig zijn (η_1 en η_2). In het model wordt er geen stijlfactor gemodelleerd naast de inhoudelijke concepten. Als dit model het beste zou passen op de data betekent dit dat er geen stijlfactor op te merken valt in de data.

Figuur 3.3.a: Twee inhoudsfactoren en geen stijlfactor (bron: Billiet McClendon, 2000)



Het tweede model dat Billiet en McClendon voorstellen is gebaseerd op een model van Mirowsky en Ross (1992). In dit model (figuur 3.3.b) worden twee inhoudsfactoren afgebeeld (η_1 en η_2) en één stijlfactor (η_3). Alle factorladingen⁹ van de variabelen op de stijlfactor worden standaard op +1 gezet. Dit omdat er verwacht wordt dat alle surveyitems op een gelijkaardige manier beïnvloed worden door de stijlfactor. Een andere belangrijke eigenschap van dit model is, dat de covarianties tussen de inhoudsfactoren en de stijlfactor op 0 worden ingesteld. Dit doen ze omdat ze verwachten dat er geen correlaties zijn tussen de stijl- en inhoudsfactoren (Billiet, McClendon, 2000).

Figuur 3.3.b: twee inhoudsfactoren en één stijlfactor (Bron: Billiet, McClendon, 2000)

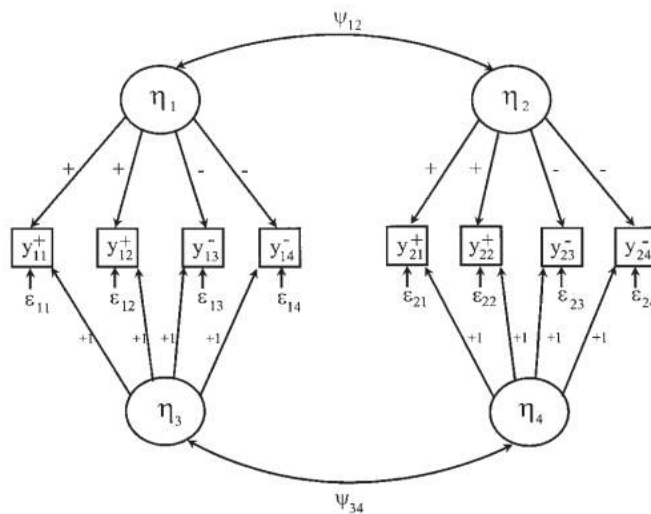


Hetgeen wat dit tweede model (figuur 3.3.b) praktisch inhoud is dat er verschillende variabelen zijn die allemaal staan voor bepaalde surveyitems. Onderliggend zijn er twee latente variabelen, die wijzen op een concept, waar telkens vier variabelen bij horen. Deze variabelen laden positief of negatief (hangt af van de richting van de vraag) op die concepten. Hoe beter een variabele laadt op een bepaalde factor, hoe hoger de correlatie dus tussen de variabele en de factor (of het concept). Als we dan een concept hebben met vier surveyitems die allemaal zeer goed laden op dezelfde factor (die staat voor het concept) weten we dat die vier surveyitems geschikt zijn om een respondent zijn mening ten opzichte van dat concept te verklaren. Naast de twee inhoudelijke factoren, heeft dit model (figuur 3.3.b) dus ook een stijlfactor. Deze stijlfactor is onderliggend aan alle acht surveyitems, ongeacht de inhoud. Dit is te wijten aan het systematisch beantwoorden van surveyitems op een manier die los staat van hun inhoud (zie definitie Paulhus, 1991).

Het derde hypothetische model (figuur 3.3.c) dat ze voorstellen bestaat uit twee inhoudsfactoren (η_1 en η_2) én twee stijlfactoren (η_3 en η_4). De twee verschillende stijlfactoren zouden dus wijzen op twee verschillende responsstijlen. Aan de ene kant achter de positief-geformuleerde items, en aan de andere kant achter de negatief-geformuleerde surveyitems.

⁹ Een factorlading geeft het correlatiecoëfficiënt weer tussen een variabele en de factor.

Figuur 3.3.c : twee inhoudsfactoren en twee stijlfactoren (Bron: Billiet, McClendon, 2000)



Dit bovenstaande model (figuur 3.3.c) kan op twee manieren bestaan. Het is mogelijk dat de covariantie (= gedeelde variantie) tussen de twee stijlfactoren 0 is, wat zou betekenen dat we met twee totaal verschillende RS te maken hebben op de twee verschillende sets van items. De tweede optie is dat de covariantie niet beperkt wordt tot 0. De verwachting hier is dat er een positieve en substantiële correlatie zou zijn tussen beide stijlfactoren (Billiet, McClendon, 2000). Deze relatie zou vervolgens wijzen op het feit dat het gaat om dezelfde RS, die in een andere mate aanwezig is in de twee verschillende sets van items. Wat er dus op zou wijzen dat respondenten, afhankelijk van het concept, gevoelig zijn voor RS. Dat het bij bepaalde surveyitems makkelijker is om RS vast te stellen dan bij andere.

Het doel van het onderzoek van Billiet en McClendon was om na te gaan welk van deze modellen het beste op de data paste. Zo kan er bevestigd worden met welke soort RS we te maken hebben. Ze passen de verschillende modellen toe en kijken welke het beste past op de data. Om dit na te gaan kijken ze naar RMSEA. Uit hun bevindingen concluderen ze dat het tweede model het beste past op de data (Billiet, McClendon, 2000). Ze concluderen dus dat er wel degelijk ARS valt vast te stellen bij de respondenten van hun dataset.

Deze conclusie komt voort uit het feit dat ze de stijlfactor wel degelijk kunnen identificeren uit de dataset, aangezien het model waar deze in zit, een goede RMSEA heeft. Beter dan het model waarin er geen stijlfactor gemodelleerd wordt. Om te weten of deze geïdentificeerde RS wijst op ARS, wordt de correlatie tussen de stijlfactor en de ARS-variabele¹⁰ berekend. Deze correlatie bedraagt 0.90, waardoor ze kunnen uitgaan van ARS. Een zwakte in hun onderzoek is echter dat er voor het berekenen van deze ARS-variabele dezelfde items gebruikten die ze in hun model betrokken. De extreem hoge correlatie tussen deze variabelen en de stijlfactor is daardoor gedeeltelijk veroorzaakt

¹⁰ Een variabele die voor de respondenten weergeeft of ze al dan niet akkoord gingen met de surveyitems. Een telvariabele, waar de scores gewogen worden.

door inhoudelijk correcte antwoorden. Hier wordt geen rekening mee gehouden door Billiet en McClendon (2000). Het had echter vermeden kunnen worden door de ARS-variabelen te berekenen aan de hand van compleet andere surveyitems. De verkregen relatie zou dan waarschijnlijk veel minder extreem zijn, en meer legitiem. Nu is het hoge regressiecoëfficiënt vooral te wijten aan een gedeelde variantie die bijna volledig inhoudelijk te verklaren is. Het andere argument waardoor ze zeggen dat de RS die ze vinden zeker ARS is, blijkt uit de zeer lage correlatie die ze vinden met de MRS-variabele. Deze variabele telt het aantal keer dat de respondenten de surveyitems beoordeelden met de middelste optie van de schaal. Door de goede fit van hun model, de correlatie van hun stijlfactor met de ARS-variabele en de niet-bestaande correlatie tussen hun stijlfactor en de MRS-variabele, lijkt het aan te nemen dat zij met hun methode wel degelijk ARS hebben gemodelleerd.

De methode die ze hanteren is een goed voorbeeld voor de methode die we in dit onderzoek zullen toepassen. Het positieve is dat er steeds variabelen kunnen worden toegevoegd en dat er dan onmiddellijk duidelijk wordt of deze al dan niet in het model passen (cfr. RMSEA). Dit is handig om de hypothesen te testen die gaan over demografische variabelen. De manier waarop Billiet en McClendon de methode hanteren is één manier, maar er zijn ook andere manieren. In hun onderzoek ontwierpen ze eerst drie complete modellen die dan vervolgens op de data werden toegepast. Zo werd er dan naar de ‘fit’ gekeken om te bevestigen of ze al dan niet gepast waren. Een andere methode is om het model stap per stap op te bouwen, en telkens te kijken naar de evolutie in de fit. Als voorbeeld van dergelijke methode maken we gebruik van een ongepubliceerd onderzoek.

3.4 Voorbeeld van een ongepubliceerd onderzoek naar ARS met SEM

Dit ongepubliceerd onderzoek is het werk van Thijssen en Rouckhout van de Universiteit Antwerpen. Deze gebruikten SEM om een RS op te sporen in een oudere SCK-Barometer. Dit vormt een perfect voorbeeld omdat de dataset die wij gaan gebruiken, een nieuwere versie is van deze oudere dataset. De dataset die ze gebruiken is dus ook ontwikkeld door het Studiecentrum voor Kernenergie om bij de Belgische bevolking te polsen naar hun mening over nucleaire energie. De dataset bevat vele variabelen en dus is het belangrijk om de juiste eruit te kiezen om in het model te betrekken.

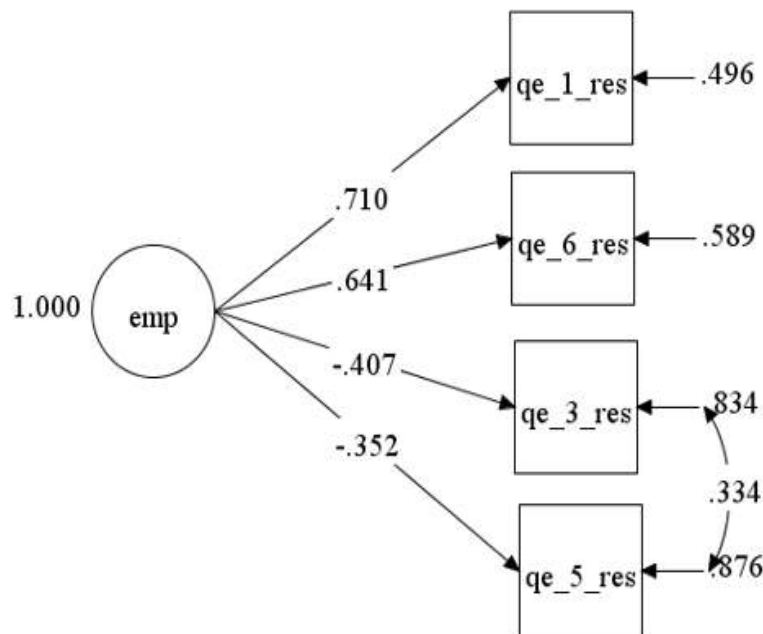
Net als Billiet en McClendon (2000) maken Thijssen en Rouckhout gebruik van twee inhoudelijke latente constructen waarachter ze op zoek gaan naar een stijlfactor. In plaats van de modellen in één keer voor te stellen maken doen zij dit stap per stap. De allereerste stap is dus om het eerste latente construct voor te stellen. Uit de Barometer werd het concept ‘empathie’ gekozen. In de dataset werd het concept gedefinieerd door vier variabelen (zie tabel 3.4.a) die te maken hebben met het thema. De surveyitems moesten beoordeeld worden aan de hand van een vijf-punten Likert schaal.

Tabel 3.4.a: surveyitems ‘empathie’

variabele	Surveyitems ‘empathie’
Qe_1_res	Ik heb vaak tedere, bezorgde gevoelens voor mensen die minder geluk hebben dan ik. (P)
Qe_6_res	Ik ben vaak geraakt door dingen die ik zie gebeuren. (P)
Qe_3_res	Soms heb ik niet veel medelijden met andere mensen wanneer ze problemen hebben. (N)
Qe_5_res	Andermans ongeluk stoort me meestal niet erg. (N)

Wanneer ze dit construct modelleren ziet het er als volgt uit:

Figuur 3.4.a – empathie als latente variabele (on gepubliceerd onderzoek Thijssen en Rouckhout)



Op de pijlen die van de factor (empathie) naar de surveyitems gaan, staan de factorladingen. Deze geven de correlatie tussen de items en de factor weer. De RMSEA van dit model bedraagt 0.063. Een acceptabele fit. Net zoals Billiet en McClendon (2000) is het de bedoeling om een stijlfactor te vinden achter verschillende inhoudelijke latente constructen. Daarom moeten er meerdere concepten zijn die worden weergegeven in het model. Behalve het concept ‘empathie’ voegen ze dus ook nog andere concepten toe. ‘Empathie’ wordt vervoegd met ‘persoonlijkheid’ dat door drie variabelen gedefinieerd wordt (zie tabel 3.4.b) die meer op het inlevingsvermogen van de respondent focussen. Samen definiëren ze de algemene factor ‘empathie’.

Tabel 3.4.b – surveyitems ‘persoonlijkheid’

variabele	Surveyitems ‘persoonlijkheid’
Qe_4_res	Ik probeer bij een meningsverschil ieders standpunt te bekijken alvorens ik een beslissing neem. (P)
Qe_7_res	Ik geloof dat er twee kanten zijn aan elke vraag en probeer naar beide te kijken. (P)
Qe_8_res	Alvorens iemand te bekritisieren probeer ik mij voor te stellen hoe ik mij zou voelen mocht ik in zijn/haar plaats zijn. (P)

Als tweede concept wordt er dus gebruik gemaakt van ‘risicoperceptie’. Dit concept wordt gedefinieerd door risicoperceptie van stralingen (zie tabel 3.4.c) en het gebruiksrisico van bepaalde objecten (zie tabel 3.4.d). De surveyitems die over risico gaan worden iets anders gemeten. Er wordt de respondent gevraagd hoe ze bepaalde risico’s inschatten voor een doorsnee Belg. Er worden allerlei mogelijke risico’s gegeven die de respondenten dan moeten beoordelen met een zes-punten schaal die gaat van ‘geen risico’ tot een ‘zeer hoog’ risico. De respondenten hebben ook de optie om aan te geven dat ze het niet weten.

Tabel 3.4.c – surveyitems ‘stralingsrisico’

variabele	Surveyitems ‘stralingsrisico’
Qrp_2_resp	Het gezondheidsrisico voor een doorsnee Belg - Radioactief aval
Qrp_5_resp	Het gezondheidsrisico voor een doorsnee Belg - Een ongeluk in een nucleaire installatie
Qrp_9_resp	Het gezondheidsrisico voor een doorsnee Belg - Een terroristische aanslag met een radioactieve bron

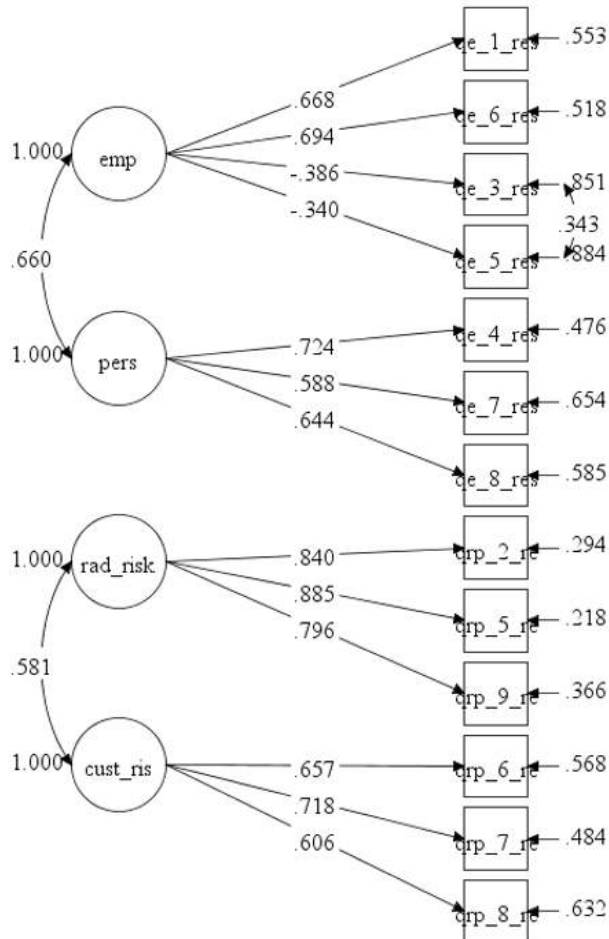
Tabel 3.4.d – surveyitems ‘gebruiksrisico’

variabele	Surveyitems ‘gebruiksrisico’
Qrp_6_resp	Het gezondheidsrisico voor een doorsnee Belg - Straling van GSM toestellen
Qrp_7_resp	Het gezondheidsrisico voor een doorsnee Belg - Natuurlijke straling
Qrp_8_resp	Het gezondheidsrisico voor een doorsnee Belg - Röntgenfoto’s in de geneeskunde

Al deze concepten worden in het model gezet met de variabelen die hun definiëren. Belangrijk om te weten is dat de correlaties tussen de concepten onderling op 0 worden gezet. De correlatie tussen empathie en risicoperceptie kan dus niet bestaan in het model. Deze stap wordt genomen omdat er geen perfect gebalanceerde schalen zijn. Bij gebalanceerde schalen is een correlatie tussen de twee inhoudelijke constructen niet dramatisch. Bij deze twee schalen zou dit echter wel een probleem zijn,

omdat de gevonden stijlfactor eventueel verklaard zou kunnen worden door inhoudelijke aspecten, wat niet de bedoeling is. Wanneer we al deze latente constructen modelleren krijgen we dit resultaat:

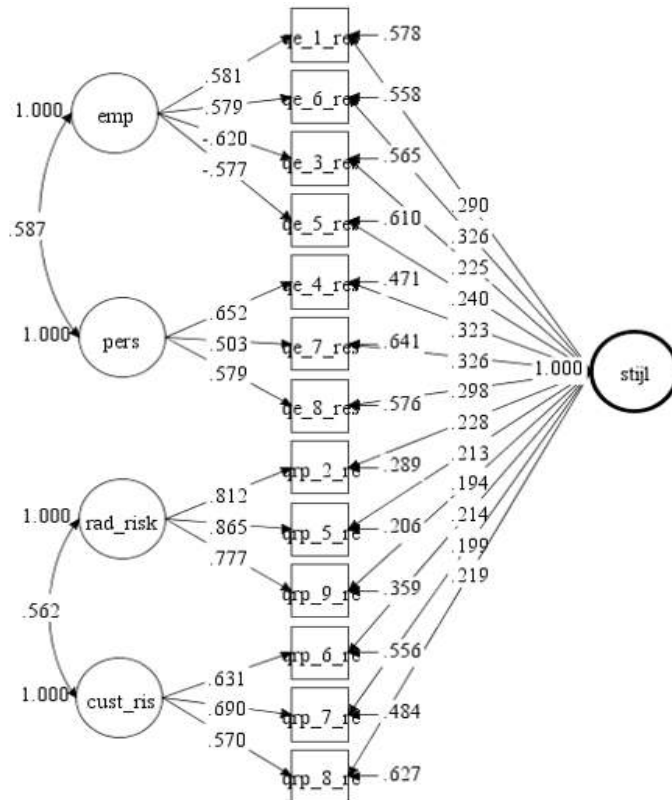
Figuur 3.4.b – Alle inhoudelijke constructen gemodelleerd (ongepubliceerd onderzoek Thijssen en Rouckhout)



We zien dat de factorladingen die de variabelen hebben op hun factor goed zijn. Dit is positief want het betekent dat de variabelen inderdaad bij de concepten passen. De RMSEA is ondertussen gezakt tot 0.046, wat al een goede fit aangeeft voor het model. Het model is dus reeds toegenomen in validiteit. De onderlinge relaties tussen de twee factoren per concept zijn ook significant, wat ook gewenst is aangezien zij bij dezelfde factor horen.

De inhoudelijke latente constructen zijn nu gemodelleerd. Het belangrijkste werk is echter om de stijlfactor toe te voegen aan het model. De stijlfactor is een factor die iets gemeen heeft met alle variabelen, los van de inhoud. Deze moeten we toevoegen aan het model. Belangrijk is dat de stijlfactor alle aanwezige variabelen in het model beslaat, en dat de lading op één wordt vastgezet op alle variabelen, omdat de stijlfactor alle variabelen evenveel beïnvloed. Ook is het belangrijk om de correlaties tussen de stijlfactor en de inhoudelijke factoren op 0 te zetten. We verwachten namelijk geen relatie. Wanneer we de stijlfactor toevoegen krijgen we het volgende resultaat:

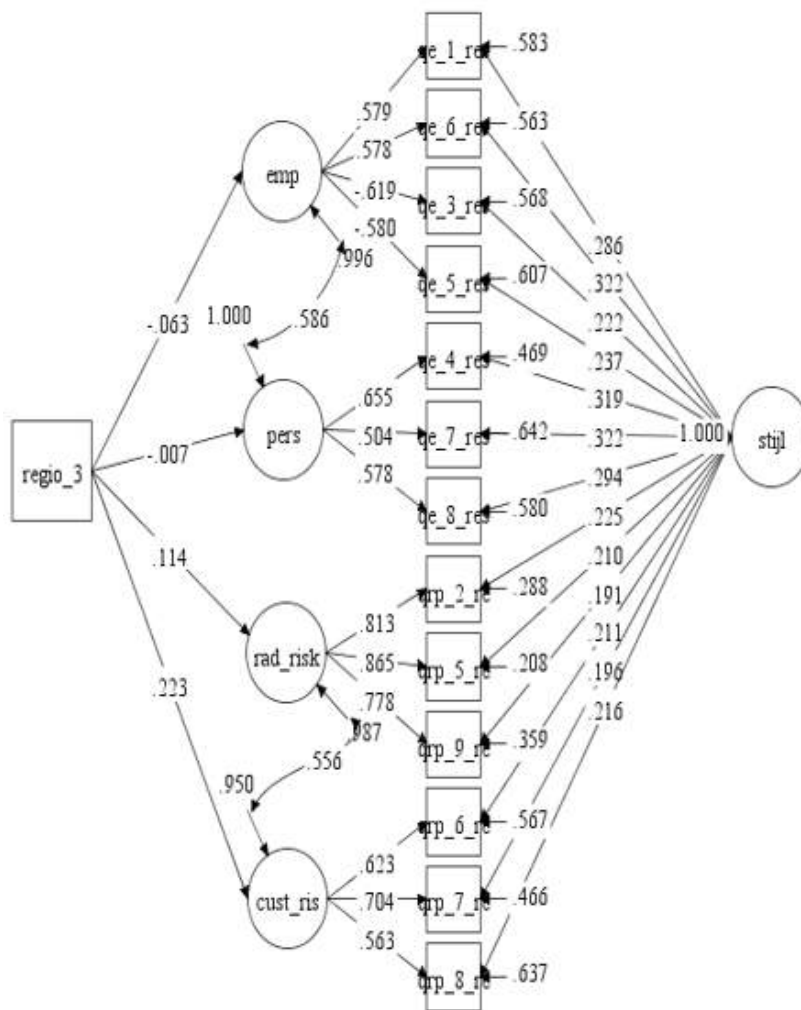
Figuur 3.4.c – stijlfactor toegevoegd aan inhoudelijke constructen (ongepubliceerd onderzoek Thijssen en Rouckhout)



Het is nu belangrijk om naar de RMSEA te kijken. Deze bedraagt ondertussen 0.045, een verbetering tegenover het vorige model. Het is dus gewenst om de stijlfactor te modelleren. De onderliggende factor is aanwezig, en verklaart een deel van de gedeelde variantie van de variabelen, die inhoudelijk losstaan van elkaar. Daarom wijst het op een stijlfactor. Inhoudelijk gezien is er geen verklaring voor deze gedeelde variantie omdat de concepten ‘empathie’ en ‘risicoperceptie’ losstaan van elkaar. Om te zien om welke RS het eventueel zou kunnen gaan, moeten we de stijlfactor laten correleren met de variabelen die als indicator voor RS dienen. Dit wordt echter pas later toegepast in de methode van Thijssen en Rouckhout. Het eerste wat ze doen is nagaan wat de invloed van de regio is op de stijlfactor.

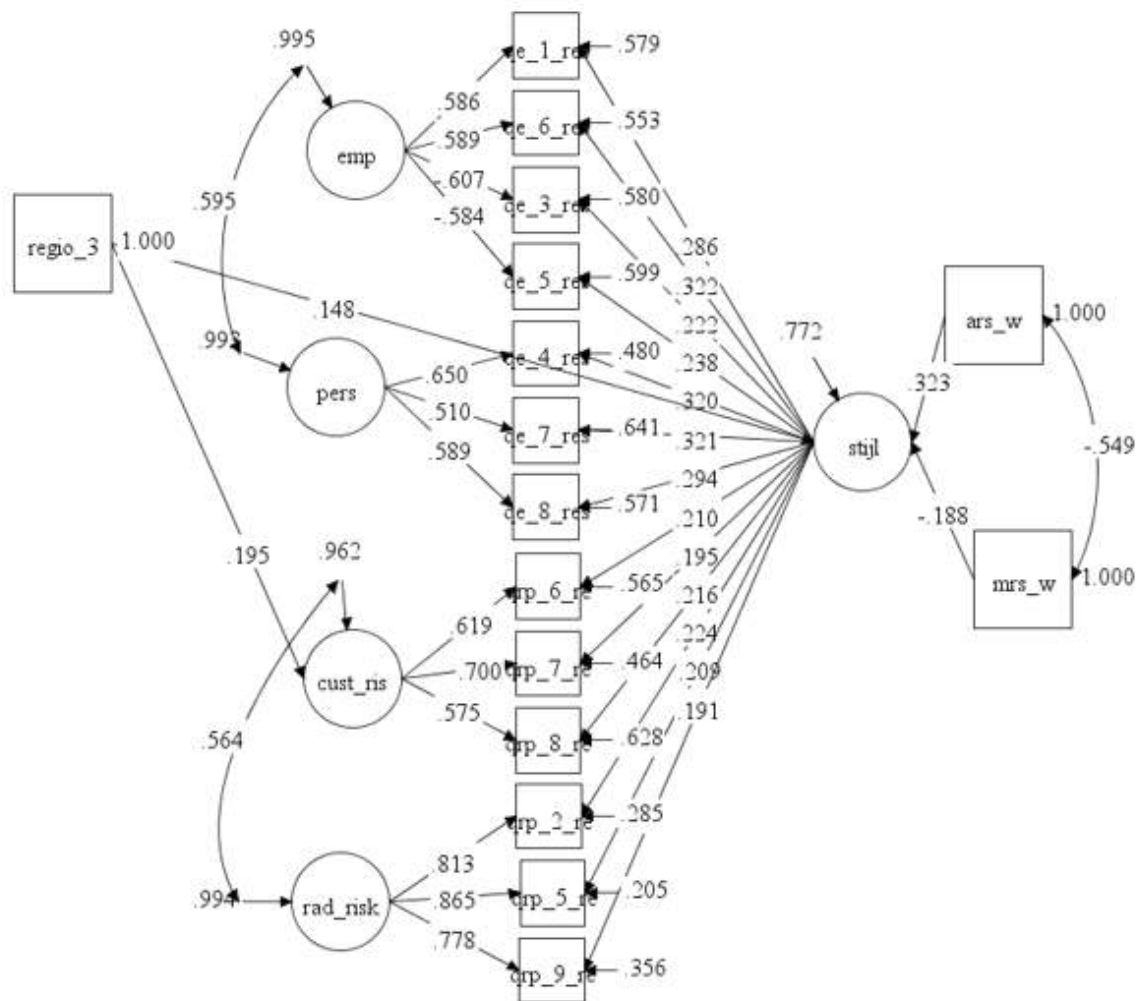
De variabele ‘regio’ wordt toegevoegd aan het model. Deze variabele slaat op het feit of de respondenten uit Vlaanderen of Wallonië komen. De variabele ‘regio’ kan zowel toegevoegd worden in het model zonder stijlfactor als in het model met de stijlfactor. Omdat we uiteindelijk uitspraken willen doen over regio in relatie tot de stijlfactoren, als responsstijlen bijvoorbeeld, voeren we regressieanalyses uit tussen regio en de inhoudelijke latente constructen. De regio-variabele ‘regio_3’ is een dummy waarbij ‘1’ staat voor een Waalse respondent, en ‘0’ staat voor een Vlaamse respondent. We verkrijgen het volgende model:

Figuur 3.4.d – regio-variabele toegevoegd aan het model (ongepubliceerd onderzoek Thijssen en Rouckhout)



Nog steeds is de fit goed met een RMSEA van 0.049. De bedoeling blijft echter om een RS te kunnen vaststellen. Nu de aanwezigheid van een stijlfactor bevestigd is aan de hand van het model, is de laatste stap om te ontdekken om welke RS het zou kunnen gaan. Daarom wordt de stijlfactor in een regressieanalyse gestoken met de variabele die staat voor ARS, en de variabele die staat voor MRS. Deze variabele telt het aantal keer dat de respondenten akkoord gingen met de surveyitems. In tegenstelling tot Billiet en McClendon (2000) berekenden Thijssen en Rouckhout deze variabelen aan de hand van surveyitems die niet in hun model voorkwamen. Op deze manier werd vermeden dat het verband zou slaan op inhoudelijke overeenkomsten. Een relatie tussen deze ARS-variabele en de stijlfactor zou wijzen op de aanwezigheid van ARS. Ondertussen wordt de regio-variabele, waarin 1 staat voor een Waalse respondent en 0 voor een Vlaamse respondent, ook toegepast op de stijlfactor. Op deze manier kan er worden nagegaan of er meer ARS is bij Vlaamse of bij Waalse respondenten. Zo krijgen we het volgende model:

Figuur 3.4.e – Volledig eindmodel van het ongepubliceerd onderzoek van Thijssen en Rouckhout



Wat meteen duidelijk wordt uit het model is dat er een positief verband is tussen de stijlfactor en de ARS-variabelen. De relatie is veel beter dan met de MRS-variabele waardoor we toch met enige zekerheid kunnen vaststellen dat de stijlfactor slaat op ARS. Het coëfficiënt bedraagt 0.323 dus is er een duidelijk verband vast te stellen. Het feit dat deze een stuk lager is dan de 0.9 van Billiet en McClendon komt voort uit het feit dat de ARS-variabelen berekend zijn aan de hand van andere surveyitems, om inhoudelijke aspecten uit te sluiten. Ook de regio-variabele heeft een positieve relatie met de stijlfactor, die significant is. Hieruit kunnen we concluderen dat er meer ARS aanwezig is bij Waalse respondenten dan bij Vlaamse respondenten. De fit van het model blijft ook bewezen door de goede RMSEA van 0.047.

Dit model vormde een voorbeeld van de methode die wij in dit onderzoek ook zullen toepassen op de nieuwe SCK-Barometer van 2015 om de hypothesen te kunnen testen. De data uit de SCK-Barometer die in dit onderzoek gebruikt zullen worden zijn echter beter omdat ze voorzien zijn van enkele goed-gebalanceerde sets met items. Deze zijn van groot belang om een goed model te kunnen ontwerpen

3.5 Analyse van de SCK-Barometer 2015

Het model dat we tot nu toe presenteerden was afkomstig uit een oudere versie van de Barometer. Het is nu de bedoeling om SEM toe te passen op de Barometer uit 2015 om eigen modellen te ontwikkelen. Deze Barometer werd ontwikkeld door het Studiecentrum voor nucleaire energie en bevat meer dan 200 surveyitems die polsen naar respondenten hun mening over nucleaire energie en alles wat daarbij komt kijken. De survey werd afgenomen bij meer dan 1000 respondenten. Hierdoor is de dataset geschikt om SEM op toe te passen aangezien het algemeen vereist wordt dat er voor SEM minstens 200 respondenten in de dataset moeten zitten. Dit is een minimumvereiste voor de validiteit van het model (Hair et al., 2010). Voordat we kunnen proberen om de aanwezigheid van RS te bewijzen, moeten er beslist worden welke variabelen we zullen opnemen in ons model. De dataset bevat meer dan 200 surveyitems dus een goede keuze is van groot belang. Net zoals in de andere onderzoeken die we reeds besproken hebben¹¹, hebben we twee – inhoudelijk losstaande – constructen nodig waarachter we de RS kunnen ontdekken. Er zijn twee inhoudelijke ‘concepten’ nodig, die ieder gedefinieerd worden door verschillende variabelen. De bedoeling is dan weer om achter al deze variabelen een achterliggende (=latente) factor (=concept) te vinden, die niet afkomstig is uit de inhoud en dus een stijlfactor wordt genoemd.

Het eerste inhoudelijke latente construct dat we gaan gebruiken bevat vier surveyitems. Deze moeten door de respondenten beoordeeld worden aan de hand van Likert schalen, dit betekent dat de respondent de statements kon beoordelen aan de hand van vijf mogelijkheden¹². De zesde optie was om de vraag open te laten, wat dan resulteerde in een *missing*. Missende resultaten worden niet betrokken in de analyse. De vier surveyitems polsen allemaal naar de declassering van Belgische kerncentrales¹³. Samen definiëren ze dus hoe de respondent staat tegenover de declassering van de Belgische kerncentrales. Het zijn gebalanceerde items, wat betekent dat de vragen in verschillende richtingen geformuleerd zijn. Om dit te illustreren staat er achter de positief-geformuleerde surveyitems een (P) en achter de negatief-geformuleerde surveyitems een (N). Hieronder een overzicht van de verschillende items die samen het latente construct vormen:

¹¹ Billiet en McClendon, 2010; Ongpubliceerd onderzoek van Thijssen en Rouckhout

¹² 1= Helemaal niet akkoord, 2= Niet akkoord, 3= Noch akkoord, noch niet akkoord, 4= Akkoord, 5= Helemaal akkoord

¹³ Het minder belangrijk maken van de Belgische kerncentrales. Ze naar de achtergrond laten verdwijnen.

Tabel 3.5.a – surveyitems ‘declassering’

Variabele	Surveystatements ‘declassering’
DE5	De nucleaire industrie beschikt over de vereiste technologie om kerncentrales succesvol te declasseren. (P)
DE6	De nucleaire industrie heeft niet de vereiste expertise om kerncentrales succesvol te kunnen declasseren. (N)
DE7	De eigenaar van de kerncentrales heeft de benodigde financiële middelen om kerncentrales te declasseren. (P)
DE8	Ik vertrouw het toezicht van de Belgische autoriteiten op de declasseringswerkzaamheden van de nucleaire industrie. (P)

Om zeker te weten of de surveyitems inderdaad geschikt zijn om te dienen als latent construct, voeren we een factoranalyse uit. Via deze analyse kunnen we nagaan of de variabelen inderdaad bij de factor ‘declassering van kerncentrales in België’ passen. Door naar de factorladingen van de verschillende items te kijken kunnen we de correlaties tussen de items en de factor nagaan. Zoals bij een normale correlatie wordt er bij de factorladingen verwacht dat ze minstens 0.5 bedragen. In de onderstaande tabel een overzicht van de factorladingen:

Tabel 3.5.b – factorladingen ‘declassering’

variabele	Surveyitems ‘wetenschap en technologie’	factorlading
DE5	De nucleaire industrie beschikt over de vereiste technologie om kerncentrales succesvol te declasseren. (P)	,821
DE6	De nucleaire industrie heeft niet de vereiste expertise om kerncentrales succesvol te kunnen declasseren. (N)	,830
DE7	De eigenaar van de kerncentrales heeft de benodigde financiële middelen om kerncentrales te declasseren. (P)	,762
DE8	Ik vertrouw het toezicht van de Belgische autoriteiten op de declasseringswerkzaamheden van de nucleaire industrie. (P)	,768

De factorladingen zijn allemaal hoger dan 0.7, wat wijst op een goede samenhang tussen de variabelen en de factor. Wanneer we de vragen zelf bestuderen merken we echter wel enig verschil op tussen de vragen. Zo gaat DE8 bijvoorbeeld als enige vraag over de Belgische autoriteiten, waar de andere items eerder focussen op de eigenaars van de kerncentrales. Ook is het Vlaamse woord ‘declassering’ waarschijnlijk minder bekend bij de respondenten dan het Franse woord “déclasser”, bij de Waalse respondenten. We gaan er echter vanuit dat de interviewers de respondenten hebben toegelicht mits zij vroegen waar het woord op wees. Respondenten hadden immers ook de optie om de vraag niet te beantwoorden omdat ze het niet wisten. Wanneer we echter naar de data kijken is het duidelijk dat er weinig *missings* zijn, wat ons doet vermoeden dat de respondenten zijn ingelicht. Bij het modelleren

zal er duidelijk worden of er inderdaad significante verschillen zijn die resulteren uit deze taal-technische opmerkingen, wanneer dit het geval is komen we hierop terug.

Het tweede latente construct moet inhoudelijk losstaan van ‘declassering van Belgische kerncentrales’. De hele barometer bevat echter items over nucleaire energie, dus volledig losstaan van elkaar is haast onmogelijk. Toch zijn er genoeg verschillende concepten, waardoor dit geen probleem voor het model mag vormen. De andere surveyitems die we in ons model gaan gebruiken zijn items die bij de respondenten polsen naar hun mening over de evolutie van wetenschap en technologie. Er zijn vier items over wetenschap en technologie die passen onder dit concept. Weer hebben we het geluk dat de items gebalanceerd zijn en dus zowel positieve als negatieve items bevat. Hieronder een overzicht van de items:

Tabel 3.5.c – surveyitems ‘wetenschap en technologie’

Variabele	Surveystatements ‘wete’
AX2	Wetenschap en technologie zullen zorgen voor hogere levenskwaliteit voor de toekomstige generaties. (P)
AX3	Wetenschap en technologie maken ons leven gemakkelijker. (P)
AX6	Wetenschap en technologie hebben de meeste mensen ongelukkiger gemaakt. (N)
AX7	Wetenschappelijke en technologische ontwikkeling hebben onvoorziene kwalijke effecten voor de gezondheid van mens en milieu. (N)

Weer gaan we door middel van een factoranalyse na of de items allemaal gepast zijn om het concept te definiëren. De factorladingen van de items zien er als volgt uit:

Tabel 3.5.d – factorladingen ‘wetenschap en technologie’

Variabele	Surveystatement ‘wete’	Factorlading
AX2	Wetenschap en technologie zullen zorgen voor hogere levenskwaliteit voor de toekomstige generaties. (P)	,752
AX3	Wetenschap en technologie maken ons leven gemakkelijker. (P)	,754
AX6	Wetenschap en technologie hebben de meeste mensen ongelukkiger gemaakt. (N)	,796
AX7	Wetenschappelijke en technologische ontwikkeling hebben onvoorziene kwalijke effecten voor de gezondheid van mens en milieu. (N)	,730

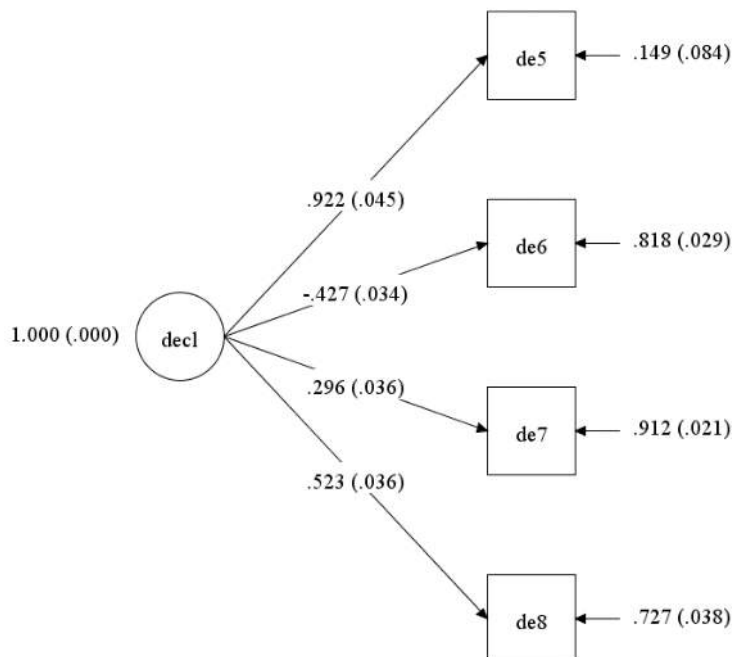
Net als de vorige surveyitems lijken deze items ook goed te passen bij hun factor. De factor die in dit geval slaat op de mening over de evolutie van wetenschap en technologie. We hebben dus twee latente constructen geselecteerd die we gaan modelleren. De bedoeling blijft om achterliggend aan al deze items een stijlfactor te vinden.

Om SEM te kunnen toepassen maken we gebruik van Mplus. De dataset wordt eerst in SPSS herwerkt tot enkel de noodzakelijke variabelen overblijven. Wanneer de dataset klaar is, gebruiken we Mplus om de variabelen op de juiste manier te kunnen modelleren. Daar kijken we naar de RMSEA om te kijken of de modellen al dan niet werken. Omdat SEM een ingewikkelde methode is, zullen we stap per stap toelichten hoe de modellen veranderd worden. Op deze manier kunnen we door middel van de RMSEA zien of het model goed op de data past en welke significante resultaten we kunnen vaststellen.

Stap 1

Voor we een achterliggende factor kunnen vinden moeten we eerst onze twee latente constructen modelleren. Het is belangrijk om te vertrekken vanuit een goed model met een goede fit. We beginnen met het eerste construct, waarin het concept ‘declassering’ gedefinieerd wordt door vier surveyitems. Uit deze vier variabelen halen we dus een factor die samenvat wat deze vier variabelen meten, in dit geval de mening van de respondenten over de ‘declassering’ van kerncentrales. We maken reeds gebruik van Mplus omdat we zo in één keer verder kunnen werken aan ons uiteindelijke model. Wanneer we deze commando’s geven in Mplus krijgen we het volgende model (3.5.a) waarbij ‘decl’ staat voor de factor die de respondenten hun mening over de declassering van kerncentrales weergeeft, en de surveyitems worden weergegeven met hun variabele-naam uit de dataset:

Figuur 3.5.a – latente variabele ‘declassering’

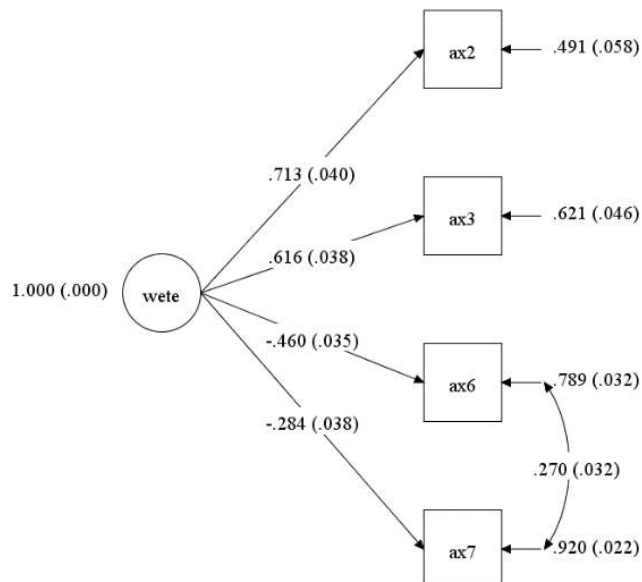


Zoals reeds vermeld kijken we naar de *Root Mean Square Error of Approximation* (RMSEA) om de fit van het model te beoordelen. In dit eerste model waar slechts één latent construct wordt weergegeven krijgen we een RMSEA van 0.033, wat meteen erg goed is.

Stap 2

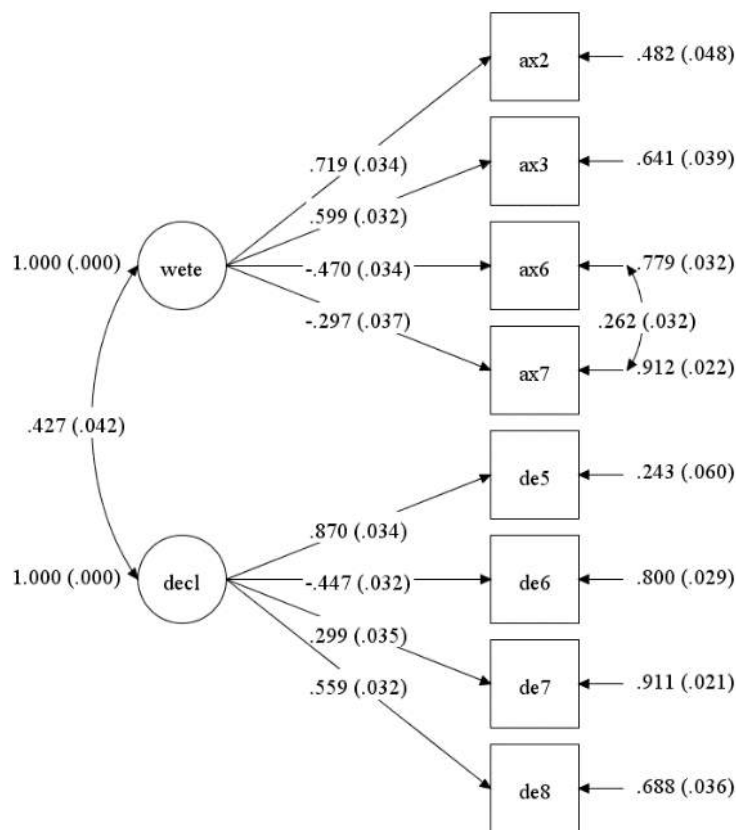
We moeten echter twee verschillende latente constructen hebben. Anders kunnen we geen achterliggende stijlfactor ontdekken. Het tweede construct dat we gaan gebruiken is een factor die de respondenten hun mening tegenover de evolutie van wetenschap en technologie weergeeft. Deze factor, die we ‘wete’ noemen wordt gedefinieerd door vier, zowel positief als negatief-geformuleerde, surveyitems. Wanneer we dit construct modelleren ziet het er als volgt uit:

Figuur 3.5.b – latente variabele ‘wetenschap en technologie’



Dit model, alleen geconstrueerd, heeft een iets mindere fit met een RMSEA van 0.079, wat nog steeds acceptabel. We verwachten dus een goede fit als we de twee constructen samen modelleren. Dit zal dan een goede basis vormen voor wanneer we het model verder ontwikkelen. Het model dat we verkrijgen wanneer we de twee constructen naast elkaar modelleren, ziet er zo uit:

Figuur 3.5.c – inhoudelijke constructen gemodelleerd

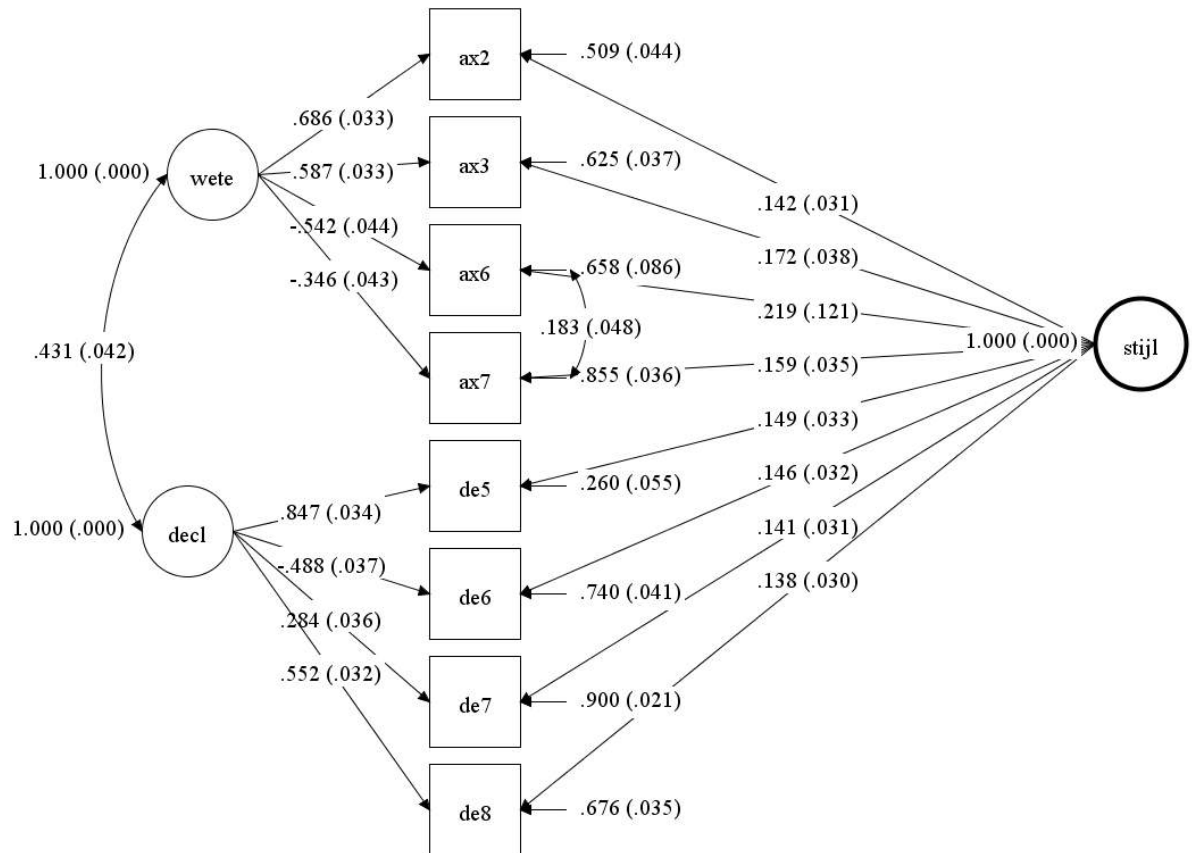


Zoals verwacht heeft dit model een erg goede fit. De RMSEA bedraagt slechts 0.046, dit geeft aan dat het model erg goed op de data past. Omdat het model goed past, vormt het een geschikt basismodel voor verdere analyse. Wat nog opvalt is dat er een behoorlijke correlatie te meten is tussen onze twee latente variabelen 'wete' en 'decl'. Dit is echter niet gewenst omdat een deel van de gelijke variantie dan verklaard zou kunnen worden door de inhoud. Omdat we echter twee gebalanceerde sets met items hebben vormt dit geen probleem voor het onderzoek. Door de gebalanceerde items wordt er immers een *least-likely case* gecreëerd. Deze gebalanceerde items worden visueel weergegeven door het feit dat de negatief-geformuleerde items negatief, en de positief-geformuleerde items positief laden op de factor.

Stap 3

Nu we de twee latente constructen naast elkaar gemodelleerd hebben is het mogelijk om de achterliggende stijlfactor te vinden. Deze factor is achterliggend dus komt in de vorm van een latente variabele. De bedoeling is dus om te zien dat alle surveyitems, van beide constructen, zowel negatief als positief-geformuleerd, laden op die achterliggende stijlfactor. Wanneer deze latente variabele, onder de naam 'stijl' toevoegen aan het model krijgen we het volgende resultaat:

Figuur 3.5.d – model met een stijlfactor



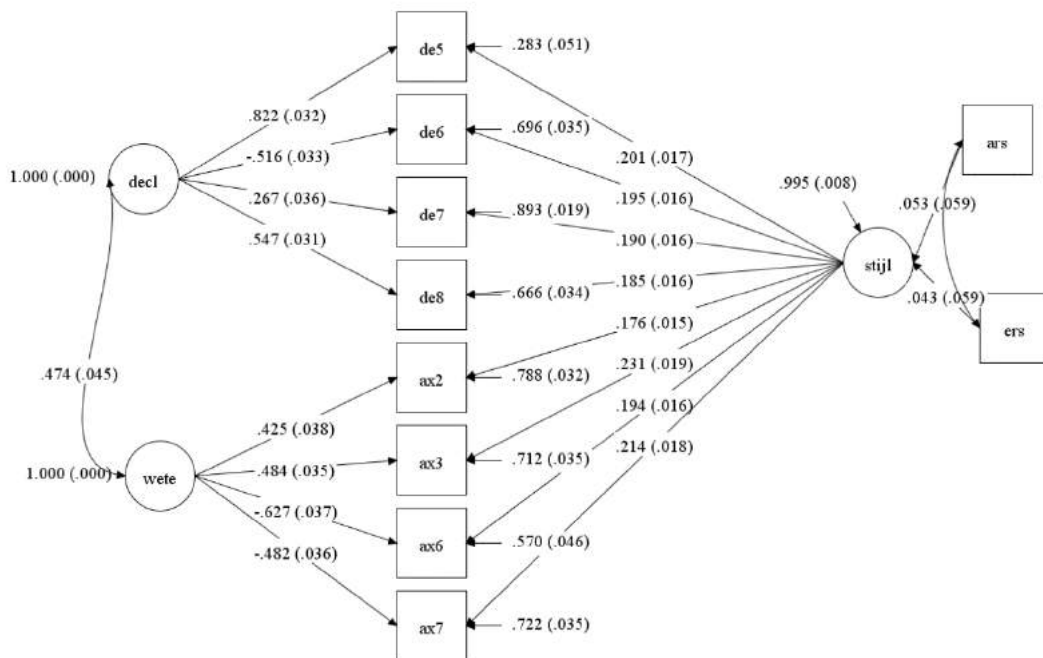
Het model lijkt goed op de data te passen aangezien de RMSEA gezakt is tot 0.043. De fit van het model is dus verbeterd nadat er een stijlfactor werd toegevoegd. De stijlfactor verklaart dus een gedeelde variantie die niet door de inhoud verklaard kan worden en het is dus gewenst om deze toe te voegen. Om te weten te komen of het RS zijn die verantwoordelijk zijn voor de stijlfactor, moeten we de relatie tussen de stijlfactor en variabelen, die staan voor de RS, bekijken.

Stap 4

Om de ARS variabele te berekenen, gebruiken we deze formule: $ARS = [f(4) + f(5) \times 2] / k$. We tellen dus hoeveel keer de respondent akkoord of helemaal akkoord was met een surveyitem en wegen hun responsen. Zo telt ‘helemaal akkoord gaan’ tweemaal zo veel mee als ‘akkoord gaan’. We berekenen deze ARS-variabele, aan de hand van variabelen die nog niet in ons model betrokken zijn. Zo vermijden we de fout die Billiet en McClendon (2000) maakten in hun onderzoek. Wanneer we namelijk variabelen zouden gebruiken die al in ons model voorkomen, zou de hoge correlatie tussen de stijlfactor en ARS-variabele voor een groot deel te wijten zijn aan de inhoud. We berekenen de

ARS-variabele aan de hand van vijftien niet-gecorrleerde surveyitems¹⁴. Iets wat eigen is aan de RIRSMACS-methode, die we later zullen bespreken. Door de variabele op deze manier te creëren voegt dit waarde toe aan de least-likely case die we ook al trachtten te creëren door de gebalanceerde items. Wanneer de ARS-variabele gemaakt is voegen we deze toe aan het model om te kunnen afleiden in hoeverre de stijlfactor inderdaad ARS inhoudt. Dit model ziet er zo uit:

Figuur 3.5.e – model met ARS en ERS-variabele



Net zoals in het ongepubliceerd onderzoek voeren we dus een regressieanalyse uit tussen de stijlfactor en de variabelen die staan voor een RS. Op deze manier hopen we te kunnen verklaren waar de stijlfactor voor staat. Ons verkregen model heeft een erg goede fit met een RMSEA van 0.042. De relatie tussen de stijlfactor en de ARS- en de ERS-variabelen zijn positief. Een deel van de stijlfactor houdt dus ARS in. De mindere relatie tussen de ARS-variabel en de stijlfactor dan in het ongepubliceerd onderzoek, kan verschillende oorzaken hebben. Het kan liggen aan de surveyitems, of aan het feit dat de *least-likely case* versterkt werd door het berekenen van de ARS-variabele aan de hand van willekeurige niet-gecorrleerde items. Hier komen we later nog op terug. Over het algemeen kunnen we wel voorzichtig aannemen **dat het nodig is om een stijlfactor te modelleren in een sample van de heterogene Belgische bevolking, waarvan een deel verklaard wordt door de aanwezigheid van ARS (H1).**

¹⁴ Zie tabel 5.2.a in hoofdstuk 5 of in de bijlagen.

HOOFDSTUK 4 – De Resultaten

4.1 Hoe rapporteer je over *Structural Equation Modelling*?

Op het einde van het vorige hoofdstuk hebben we reeds enkele resultaten van SEM weergegeven, voor we alle resultaten bespreken is het gewenst om uit te leggen hoe we deze gaan rapporteren. Er wordt gebruik gemaakt van het programma Mplus. In dit programma moet de gebruiker zelf een input schrijven in de vorm van een syntax. Alle geschreven syntaxen die gebruikt werden als input voor de modellen in deze paper zijn terug te vinden in de bijlagen. Wanneer je de input runt, geeft Mplus een output weer waar veel informatie in valt terug te vinden. Naast gemiddeldes, varianties, covarianties, regressiecoëfficiënt en correlaties, geeft deze output ook de fit van het model weer. Uit deze gegevens wordt voor deze paper de RMSEA gebruikt om de fit aan te geven. De uitkomsten van de modellen worden op verschillende manieren weergegeven. Er zijn verschillende opties om een gestandaardiseerde versie te verkrijgen die makkelijker te interpreteren valt. De gebruiker moet zelf kiezen welke versie er gebruikt wordt aan de hand van de soort van de variabelen die gebruikt worden. Een erg handige functie van Mplus is dat alle uitkomsten visueel kunnen worden weergegeven in een diagram. Dit is nodig om een overzichtelijk beeld te kunnen scheppen van het model. We opteren er dan ook voor om alle modellen visueel (als een diagram) weer te geven en enkel noodzakelijke elementen uit de output in de uitleg te verwerken. Het is ook nuttig om te weten dat latente variabelen steeds worden weergegeven als cirkels in plaats van vierkanten.

4.2 Resultaten van de hypothesen testen

4.2.1 Hypothese 2: Dubbelzinnige items

Onze tweede hypothese stelt dat er meer ARS zal voorkomen bij dubbelzinnig of moeilijk-geformuleerde items dan bij duidelijk en helder-geformuleerde items. Voor het vorige model gebruikten we het concept ‘declassering’ en het concept ‘evolutie van wetenschap en technologie’, meer bepaald de respondenten hun mening over deze concepten. De surveyitems die horen bij het concept ‘wetenschap en technologie’ zijn duidelijk en helder geformuleerd (zie tabel 3.5.c). De surveyitems die horen bij het concept ‘declassering’ zijn daarentegen moeilijker te snappen. Het woord ‘declassering’ alleen al is voor veel mensen geen evidentie. Om de hypothese te testen zullen we een stijlfactor zoeken achter het duidelijk-geformuleerde concept ‘wetenschap en technologie’ en een ander concept met duidelijk-geformuleerde items, en dus niet declassering. Op deze manier kunnen we de het model met en zonder ‘declassering’ met elkaar vergelijken.

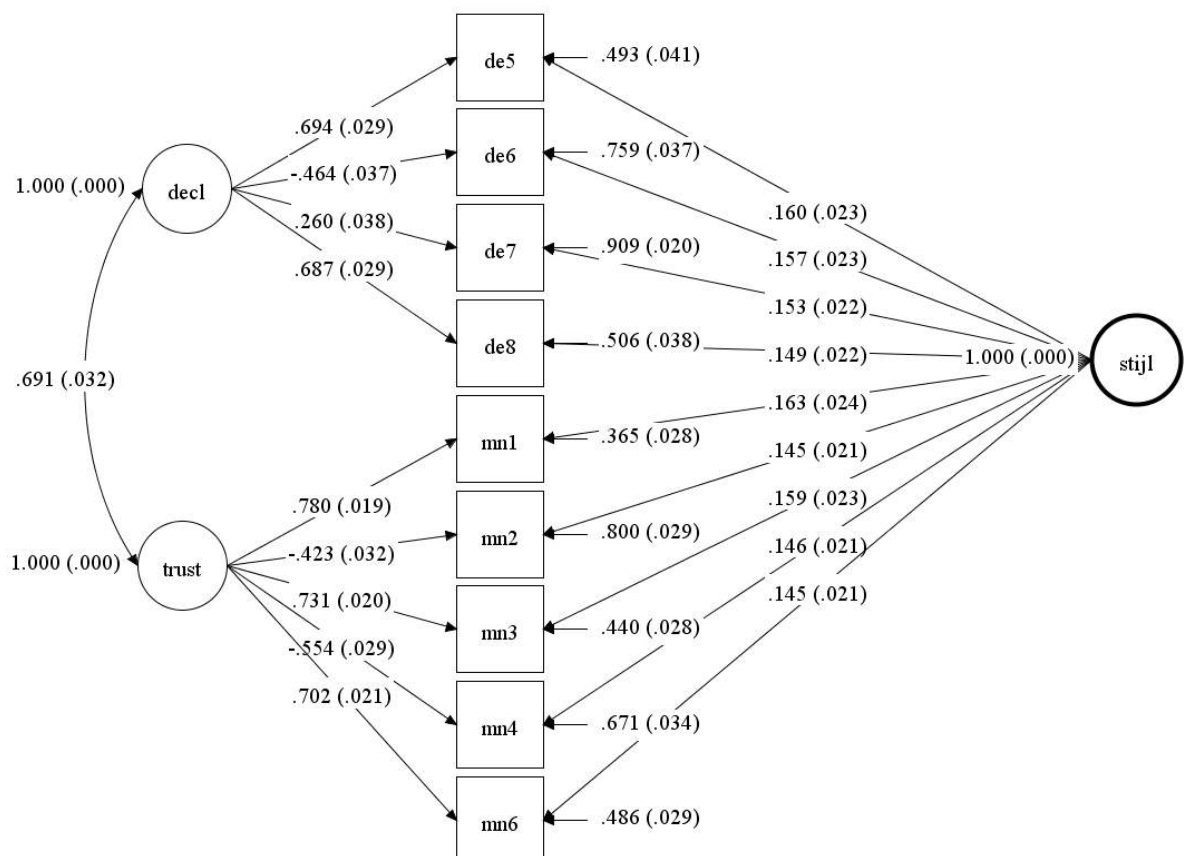
Als nieuw concept gebruiken we een concept dat we ‘trust’ noemen en dat gedefinieerd wordt door vijf surveyitems. Deze surveyitems gaan allemaal over het vertrouwen dat er is in de veiligheid van nucleaire energie in België. Hieronder een overzicht van de vijf surveyitems:

Tabel 4.2.1a – surveyitems ‘trust’

variabele	Surveystatements ‘trust’
MN1	Nucleaire reactoren in België worden op een veilige manier bestuurd. (P)
MN2	De autoriteiten controleren de veiligheid van de nucleaire installaties in België onvoldoende. (N)
MN3	In België wordt er veilig omgegaan met radioactief afval. (P)
MN4	Het transport van nucleaire materialen is onveilig. (N)
MN5	Ik voel me goed beschermd tegenover nucleaire installaties.(P)

Wanneer we dit construct modelleren naast het ‘wete’ construct, en er een stijlfactor achter zoeken, krijgen we het volgende model:

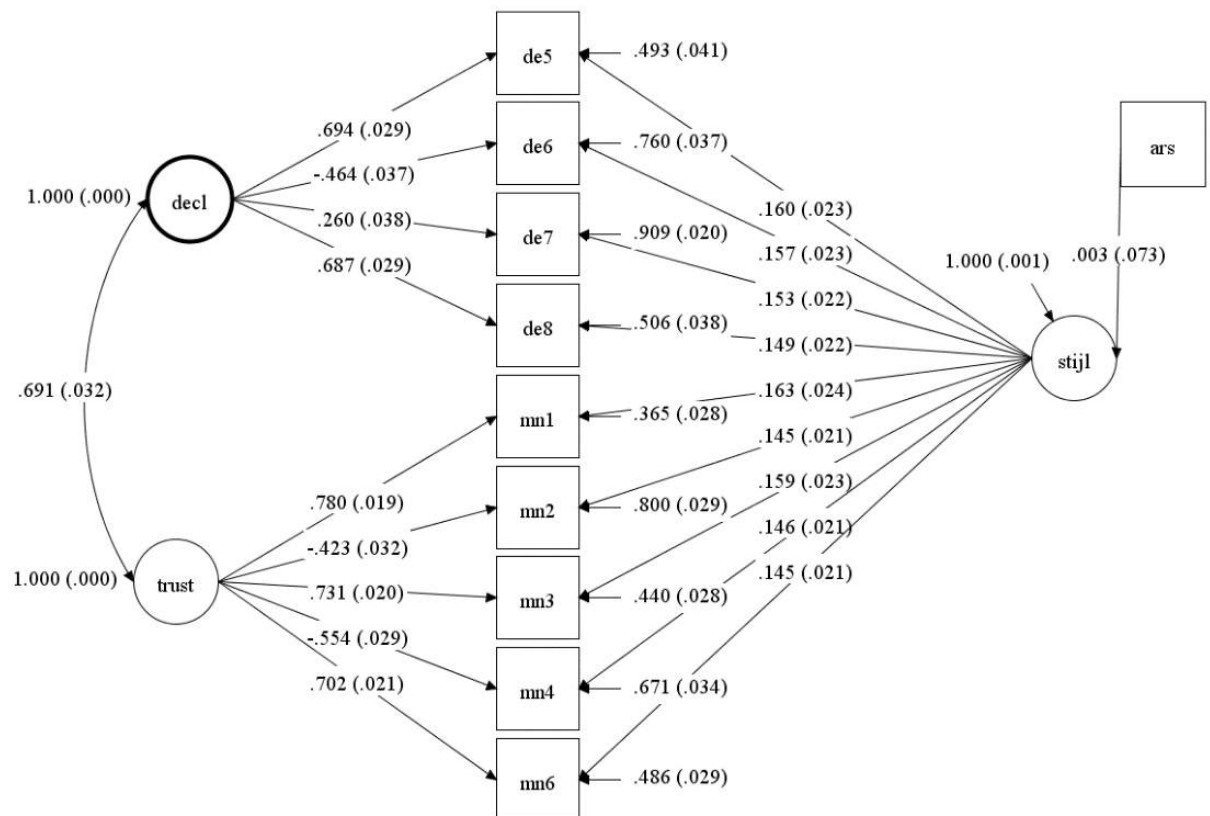
Figuur 4.2.1.a – model met ‘trust’ + stijlfactor



Het model heeft een acceptabele fit door de RMSEA die 0.063 bedraagt. Wat er echter wel opvalt aan dit model is de erg hoge correlatie tussen de twee latente constructen. De correlatie bedraagt 0.691 en is daarmee een stuk hoger dan de correlatie die we verkregen tussen ‘decl’ en ‘wete’. Vandaar dat er ook

gekozen wordt om de andere hypothese te testen terwijl we gebruik maken van de andere concepten. Om H2 te kunnen testen moeten we echter nog de ARS-variabele toevoegen aan dit model. Dit ziet er als volgt uit:

Figuur 4.2.1.b – model met ‘trust’ + ARS-variabele



De RMSEA is 0.053, het model past dus goed op de data. Wanneer we kijken naar de relatie tussen de ARS-variabele en de stijlfactor zien we dat er minder ARS is vast te stellen dan in het model dat het construct over ‘declassering’ bevatte. Daarom accepteren we de hypothese en stellen dat ***er meer ARS is vast te stellen in modellen met dubbelzinnig of moeilijk geformuleerde items (H2).***

4.2.2 Hypothese 3: Gender

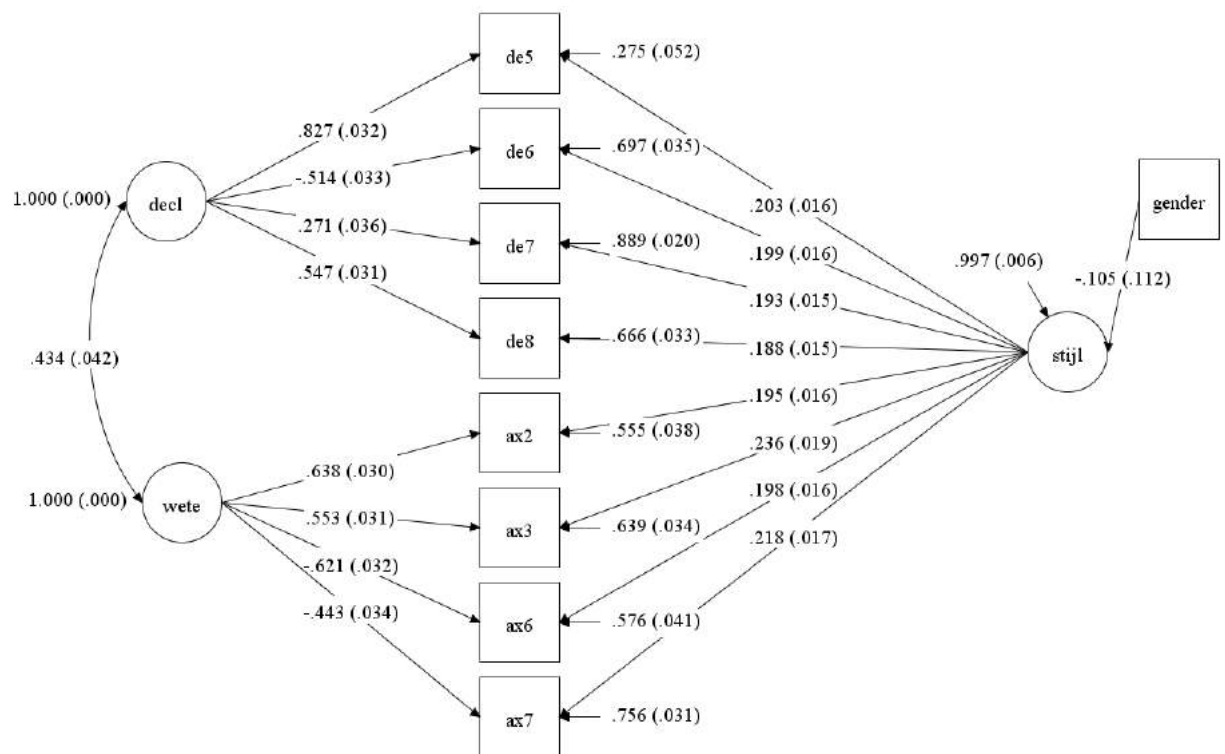
Onze derde hypothese luidt dat er een verband is tussen gender en ARS, namelijk in de zin dat er bij vrouwen vaker ARS voorkomt dan bij mannen. Deze hypothese komt voort uit de literatuur en is reeds vaker gebleken in andere onderzoeken. Het geslacht van de respondenten werd uiteraard opgenomen in de dataset als een demografische variabele. We zien dat de verhouding tussen mannen en vrouwen ongeveer in evenwicht is bij de respondenten van deze dataset:

Tabel 4.2.2.a – frequentietabel gender-variabele

	Aantal	Percentage
Mannen	513	49,9 %
Vrouwen	515	50,1 %
Totaal	1028	100 %

Aangezien de verhouding haast perfect gelijk is, lijkt het al onmogelijk om – al dan niet gewenste - resultaten hierop af te schuiven. Om de relatie tussen gender en ARS weer te geven moeten we de gender-variabele toevoegen aan het bestaande model. Dan krijgen we het volgende resultaat:

Figuur 4.2.2.a - model met de gender-variabele

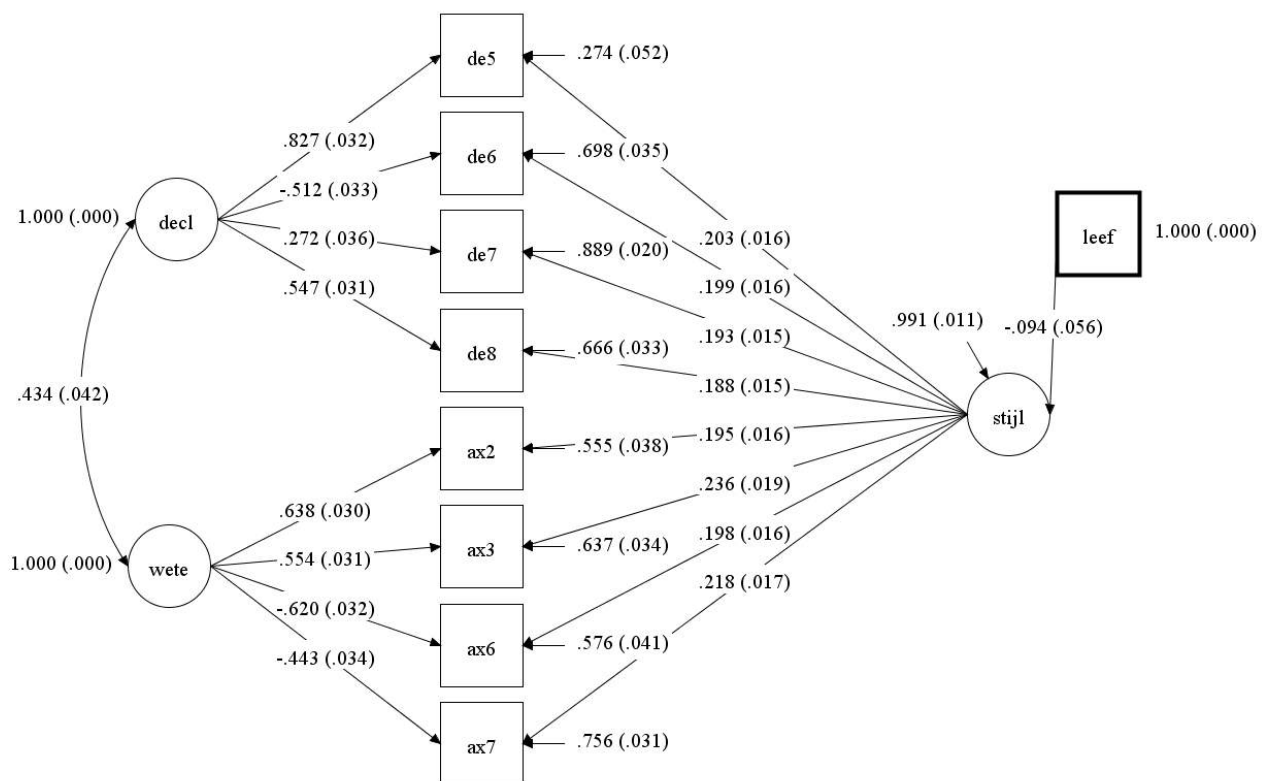


In de output van het model zien we dat de fit nog steeds goed is. De RMSEA bedraagt 0.044. In de output kunnen we ook kijken naar de resultaten van de regressieanalyse die ook in het model staat. We stellen een negatieve relatie van -0.105 vast. En aangezien de gender-variabele een dummy was waar '0' stond voor een vrouwelijke respondent en '1' voor een mannelijke, blijkt dat in deze dataset, er meer ARS valt vast te stellen bij vrouwen dan bij mannen. Wat op zich wel in lijn is met de onze hypothese en dus de eerder verschenen literatuur. Met een p-waarde van 0.260 weten we dat deze relatie die wel degelijk bestaat in deze dataset, maar moeilijk is te veralgemenen naar de Belgische bevolking. De uitkomst is dus niet significant waardoor we geen uitspraken kunnen doen over de gevormde hypothese. Wel zien we dat er in deze dataset over het algemeen meer ARS valt vast te stellen bij vrouwen dan bij mannen. Maar over het algemeen **kunnen we geen relatie vaststellen tussen gender en de stijlfactor (H3)**.

4.2.3 Hypothese 4: Leeftijd

Zoals bij de meeste onderzoeken werd er ook bij de Barometer uit 2015 naar de leeftijd van de respondenten gevraagd. Ze moesten hun geboortjaar kenbaar maken aan de enquêteur. We hebben deze variabele echter opnieuw gecodeerd tot een continue variabele die voor elke respondent zijn/haar leeftijd aangeeft door middel van het daadwerkelijke getal. De gemiddelde leeftijd van de respondenten bedraagt 49 jaar. De jongste respondent is 18 jaar oud en de oudste respondent heeft een leeftijd van maar liefst 92 jaar. De variabele heeft dus een brede spreiding. Wanneer we de leeftijd-variabele toevoegen aan ons model met de stijlfactor als onafhankelijke variabele, krijgen we het volgende model:

Figuur 4.2.3.a – model met de leeftijd-variabele



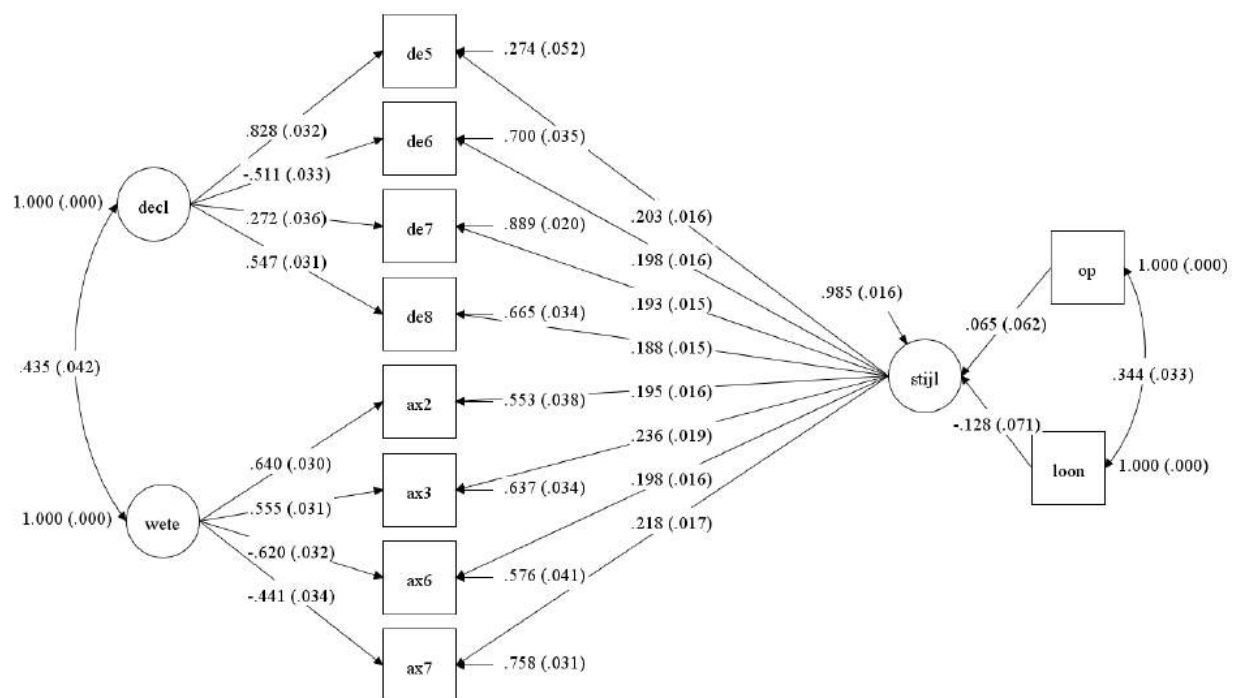
Het regressiecoëfficiënt dat we verkrijgen door leeftijd op de stijlfactor te zetten bedraagt helaas slechts - 0.094 en door te hoge p-waarde kunnen we geen significantie relatie vaststellen. Hierdoor kunnen we geen uitspraken doen over onze hypothese behalve dan ***dat er geen significante relatie is vast te stellen tussen de aanwezigheid van ARS en de leeftijd van de respondent (H4).***

4.2.4 Hypothese 5a en 5b: Opleidingsniveau en Loon

Het opleidingsniveau en het loon van respondenten is vaak, maar niet altijd, op een positieve manier met elkaar verbonden. Daarom lijkt het gepast om beide variabelen samen te modelleren. Om het opleidingsniveau van de respondenten te meten werd er hen gevraagd naar hun hoogst-behaalde diploma.

De variabele heeft dus waarden van 1 tot 9 waarbij 9 staat voor een universitaire opleiding en 1 voor het behalen van een diploma van het lager onderwijs. Er is in de dataset geen variabele aanwezig die het exacte loon van de respondenten weergeeft, omdat dit vaak informatie is die respondenten liever niet geven. Er werd de respondenten echter wel naar hun huidige beroep gevraagd. Deze variabele zal dan dienen als een indicator voor de lonen van respondenten. Omdat een deel van de respondenten al op zekere leeftijd is, werden de antwoorden ‘met pensioen’ of ‘met brugpensioen’ aangeduid als *missings*. We gaan er namelijk van uit dat de hypothese niet van toepassing is op mensen die met pensioen zijn, aangezien zij vroeger uiteraard wel een ander loon of bepaalde tewerkstelling hadden. Wanneer we deze variabelen modelleren krijgen we het volgende resultaat:

Figuur 4.2.4.a – model met loon en opleidingsvariabelen



We zien onmiddellijk dat positieve correlatie tussen loon en opleidingsniveau inderdaad bestaat. Dit is positief en betekent dat we een goede indicator voor loon hebben gebruikt. Er is namelijk een positieve relatie die 0.345 bedraagt. De afzonderlijke invloed van beide variabelen op de stijlfactor stelt echter teleur. We zien twee minimale relaties waarvan er geen enkele significant is. De kwaliteit en nauwkeurigheid van de variabelen kan hier gedeeltelijk aan de basis van liggen. De fit van het model blijft wel erg goed met een RMSEA van 0.055, wat een goede fit is. Over de hypothese kunnen we ons niet uitspreken aangezien **er geen relatie is vast te stellen tussen de stijlfactor en het opleidingsniveau en het loon van respondenten (H5a en H5b).**

4.2.5 Hypothese 6a en 6b: Verschil tussen Waalse en Vlaamse respondenten

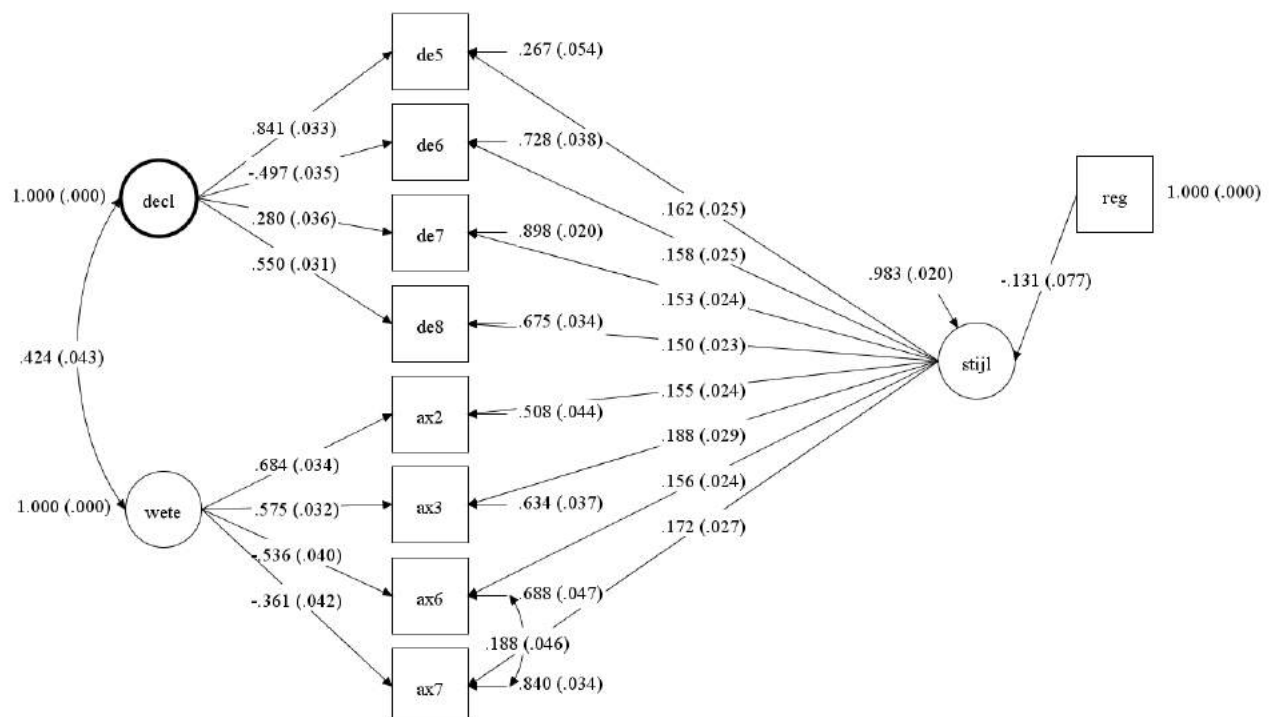
Deze survey is afgenomen bij meer dan 1000 respondenten over heel België. Dit wil zeggen dat er respondenten uit zowel Vlaanderen als Wallonië zijn terug te vinden in de dataset. Dit is van belang om onze laatste hypothese te kunnen testen. Voor we de demografische variabele gaan modelleren is het belangrijk om naar de verdeling van de variabele te kijken. Dan krijgen we het volgende resultaat:

Tabel 4.2.5.a – frequentietabel regio-variabele

	Aantal	Percentage
Vlaming	605	58.9
Waal	339	33.0
Brusselaar	84	8.2
Totaal	1028	100.0

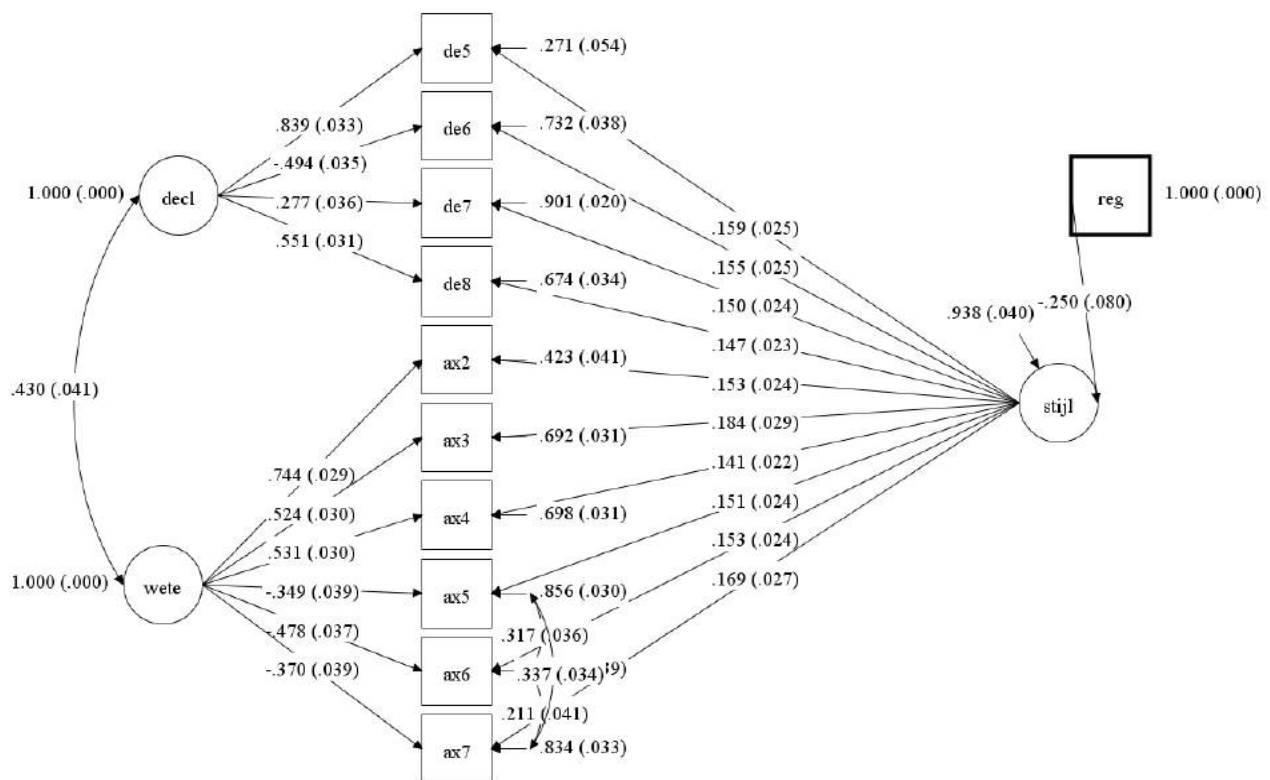
We zien dat er wel duidelijk meer Vlamingen vertegenwoordigd zijn in de dataset dan Walen. Ook werden er 84 Brusselaars ondervraagd, die voor de correctheid van de variabele tot de missings werden gerekend. De variabele die regio aangaft is namelijk een dummy waarbij ‘0’ staat voor Vlaming en ‘1’ voor een Waal, Brusselaar zijn dus ‘99’ en missings die niet in de analyse betrokken worden. Wanneer we de regio-variabele in het model implementeren om te kijken welke invloed het heeft op de aanwezigheid van de stijlfactor, krijgen we het volgende model:

Figuur 4.2.5.a – model met de regio-variabele toegevoegd



De fit van het model is acceptabel met een RMSEA van 0.079. Deze is wel minder dan de fit die we verkregen in onze vorige modellen. Naast de acceptabele fit vinden we ook een negatieve relatie die -0.131 bedraagt. Aangezien '0' stond voor Vlaming en '1' voor Waal, betekent deze negatieve relatie dat er meer ARS valt vast te stellen bij Vlamingen dan bij Walen. De p-waarde van deze relatie bedraagt 0.089. Over het algemeen worden p-waarden die kleiner zijn dan 0.05 gezien als significant. Onze p-waarde wordt eerder gezien als 'neigend naar significant'. Wat dit cijfer concreet betekent is dat in 8,9% van de gevallen onze nulhypothese niet verworpen kan worden. Wanneer we regio echter in een iets minder passend model zetten, vinden we echter ook een negatieve, en in dat geval een zeer significante, relatie tussen regio en ARS. Dit model dit er zo uit:

Figuur 4.2.5.b – model met de regio-variabele toegevoegd op groter model

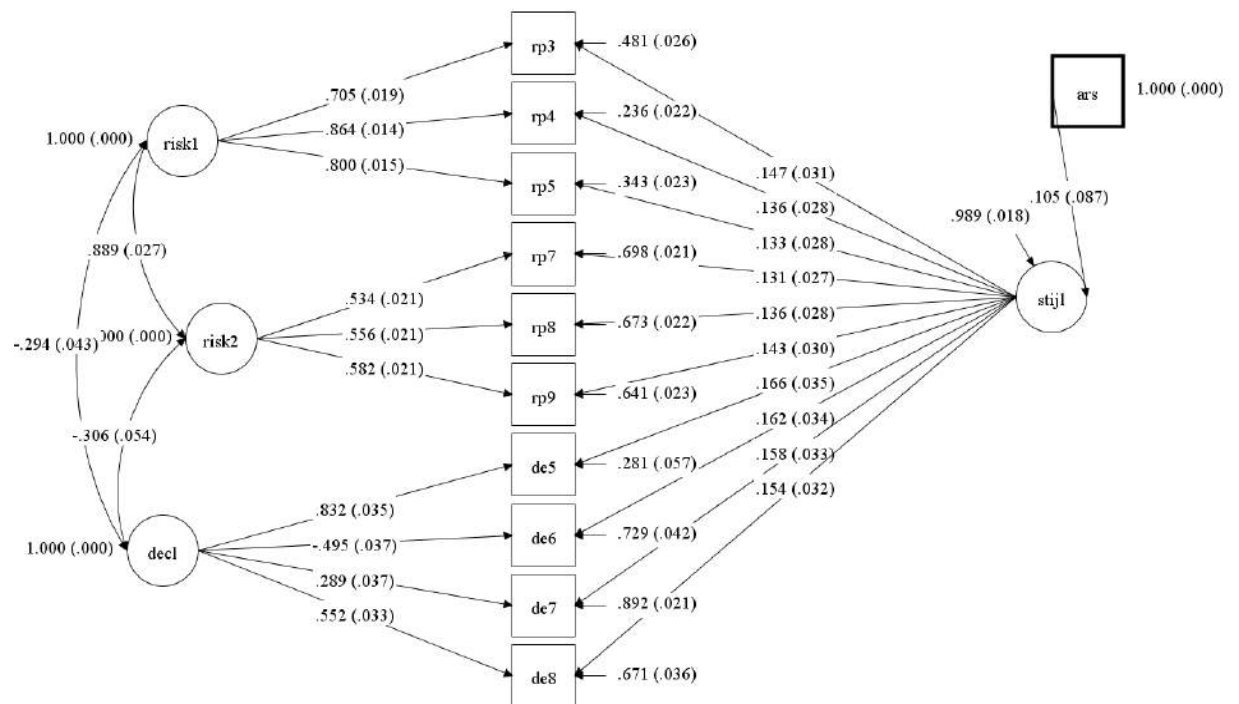


Daarom dat we toch voorzichtig onze hypothese 6b willen aannemen en stellen dat **er meer ARS is vast te stellen bij Vlamingen dan bij Walen (H6b)**.

Over het algemeen zijn er grote verschillen vast te stellen tussen de resultaten van dit onderzoek en de resultaten van het ongepubliceerd onderzoek van Thijssen en Rouckhout. Zo wordt er in ons model minder ARS vastgesteld. Ook vinden we in verband met de regio-variabele, een tegengesteld resultaat. Zo komt er uit dit onderzoek dat er meer ARS is vast te stellen bij Vlamingen dan bij Walen, en kwam er uit het ongepubliceerd onderzoek net het tegenovergestelde. Zij maakten gebruik van een andere dataset, een versie van de Barometer die eerder ontwikkeld werd en dus ook andere respondenten

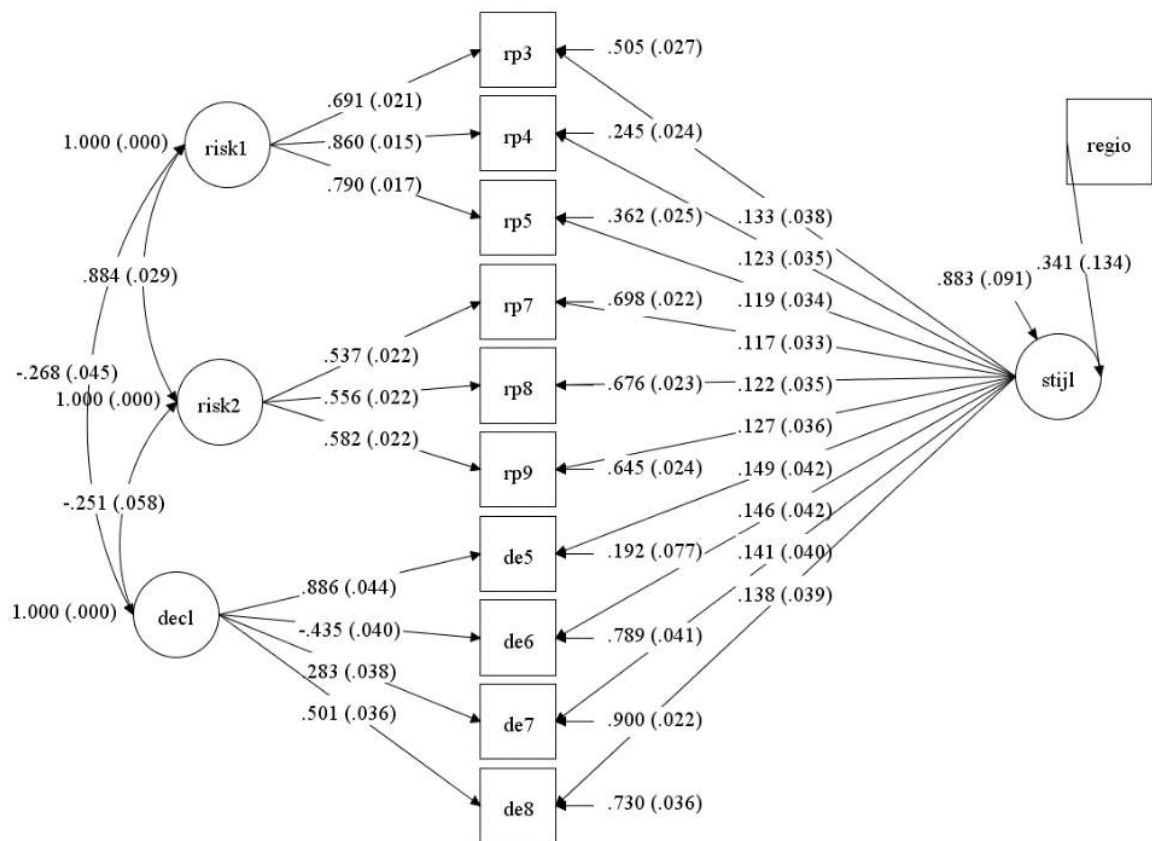
inhield. Technisch gezien is het dus mogelijk om deze tegenstrijdige uitkomsten te bekomen, maar theoretisch gezien is het wel opmerkelijk. Wat ons doet vermoeden dat de aanwezigheid van RS meer afhankelijk is van de inhoud van de surveyitems dan vooraf gedacht. In dit onderzoek werden namelijk andere surveyitems gebruikt. Dit zou dus aan de basis kunnen liggen van het verschil. Om dit vermoeden te bewijzen modelleren we ARS opnieuw, maar dan gebruikmakend van zoveel mogelijk gelijkaardige surveyitems als het ongepubliceerde onderzoek. We vervangen het construct ‘wetenschap en technologie’ door de items over risicoperceptie (zie tabel 3.4.c en 3.4.d) die ook in het ongepubliceerd onderzoek werden opgenomen. Op deze manier kunnen we kijken of het verschil in surveyitems aan de basis zou kunnen liggen. Het model ziet er dan als volgt uit:

Figuur 4.2.5.c – Model met risicoperceptie items



We zien dat er inderdaad meer ARS is vast te stellen in dit model. De relatie hier bedraagt 0.105 wat bijna gelijkaardig is aan het ongepubliceerd onderzoek. Het model is geschikt met een RMSEA van 0.059. Dit bevestigt ons vermoeden dat RS gevoelig zijn aan de inhoud. De vraag blijft echter of dit ook het geval is voor de verschillende resultaten in verband met regio. Wanneer we de onafhankelijke regio variabele toevoegen aan dit model krijgen we het volgende resultaat:

Figuur 4.2.5.d – Model met risicoperceptie surveyitems + regio-variabele



Wat we hier vinden is een positieve significante relatie tussen regio en de stijlfactor. Aangezien bij de regio-variabele '1' stond voor een Waalse respondenten en '0' voor Vlaamse respondenten, betekent dit dat er meer ARS is bij Waalse dan bij Vlaamse respondenten. Dit resultaat is tegengesteld aan het resultaat dat we vinden bij het model met de andere surveyitems. Het feit dat we de resultaten kunnen omdraaien door andere surveyitems te gebruiken bewijst dat de aanwezigheid van RS wel degelijk afhankelijk is van de inhoud. Een bevinding waarover verder onderzoek nog zeker gewenst is.

HOOFDSTUK 5 – De RIRSMACS-methode

5.1 RIRSMACS-methode geïllustreerd

Een deel van de tekortkomingen van SEM worden opgelost door de nieuwere RIRSMACS-methode. Dit staat voor *Representative Style Means and Covariance Structure* (RIRSMACS). Het allergrootste verschil aan deze methode is dat deze methode meerdere RS tegelijkertijd en gelijkaardig modelleert. Zoals reeds gezegd kan het zijn dat als je maar op zoek gaat naar één RS, je deze niet vindt omdat ze wordt gecompenseerd door een andere RS, die je niet in je model opneemt. Dit probleem kan vermeden worden met de RIRSMACS-methode omdat er meerdere RS tegelijkertijd gemodelleerd kunnen worden. Deze kunnen dan naast elkaar bestaan zonder elkaar op te heffen. Verder is de methode eerder een aanvulling op SEM aangezien er ook gebruik wordt gemaakt van latente constructen om de RS weer te geven.

RIRSMACS-methode is relatief nieuw. Dit betekent dat het nog niet bijzonder vaak is toegepast in onderzoeken. Het bekendste onderzoek dat gebruik maakte van de RIRSMACS-methode is het onderzoek van Weijters et al. (2008). Ook is er het onderzoek van Weijters, Geuens en Schillewaert (2009), waarin zij stap voor stap toelichten hoe de RIRSMACS-methode best kan worden toegepast. In hun geval gaat het om een comparatief onderzoek, waarin verschillende methode van datacollectie vergeleken worden, en gekeken wordt waar er de grootste aanwezigheid van RS valt vast te stellen. Ze gebruiken dus één lijst met surveyitems die aan de respondenten steeds op een andere manier werd voorgelegd. Door op al deze verschillende methodes de RIRSMACS-methode toe te passen, willen onderzoekers erachter komen welke soorten RS vaker voorkomen bij welke soorten manieren van dataverzameling.

De allereerste belangrijke stap – alvorens de RIRSMACS-methode te kunnen toepassen – is om na te gaan of de data geschikt zijn. Voor deze RIRSMACS-methode is het namelijk van groot belang dat er een aantal niet-gecorreleerde items aanwezig zijn in de dataset. Dit is nodig omdat de indicatoren voor de verschillende RS zullen worden afgeleid uit verschillende groepen niet-gecorreleerde surveyitems. Concreet betekent dit dat de survey, items moet bevatten die allemaal iets anders meten en niet samenhangen. Zoals er bij SEM op zoek wordt gegaan naar niet-samenhangende latente constructen, dus bijvoorbeeld twee keer vier surveyitems die twee concepten meten, zijn er voor de RIRSMACS-methode minimum zes niet-gecorreleerde items nodig. Zes is het minimum, maar het is toch gewenst dat een sample minstens vijftien dergelijke items bevat (Weijters et al., 2009). Moest dit een probleem vormen in de zin dat de survey te lang wordt, valt dit steeds op te lossen door de inhoudsfactoren te meten met behulp van minder variabelen. Zo worden er in de survey minder gelijkaardige variabelen opgenomen die toch hetzelfde meten (Weijters et al., 2009).

Deze niet-gecorrleerde items zijn dus nodig omdat de indicatoren voor de verschillende RS berekend zullen worden aan de hand van de groepen willekeurige surveyitems. De indicatoren die hieruit voortkomen zijn op hun beurt de variabelen die de latente stijlfactor, die staat voor een RS, aangeven. Het model is dus redelijk gelijkend aan een SEM-model, behalve dan het feit dat de RS gedefinieerd worden door indicatoren die bepaald worden door sets van willekeurige surveyitems, en worden weergegeven als een latente variabele. Het andere grote verschil is natuurlijk dat dit gedaan wordt voor meerdere RS tegelijkertijd. De indicatoren die berekend worden uit deze groepen met niet-gecorrleerde items worden volgens bepaalde formules berekend, namelijk:

$$ARS = [f(4) + f(5) \times 2] / k$$

$$DRS = [f(1) \times 2 + f(5)] / k$$

$$ERS = [f(1) + f(5)] / k$$

$$MRS = f(2) / k$$

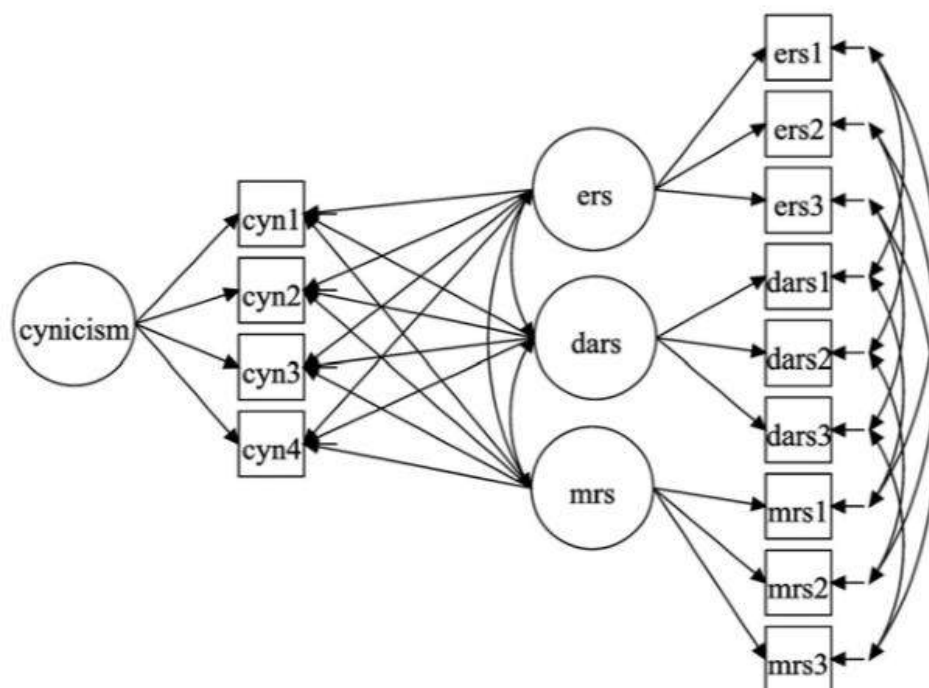
Deze formules zijn van toepassing in een dataset waar de surveyitems beoordeeld werden door middel van een Likert schaal, waarbij ‘5’ stond voor ‘helemaal akkoord’ en ‘1’ stond voor ‘helemaal niet akkoord’. De f die voor de getallen staat, wijst op de aantal¹⁵ keer dat de respondent die optie gebruikte om een surveyitem te beoordelen. De k slaat op het aantal variabelen waarop de formule wordt toegepast. Deze formules worden dus toegepast op de groepen niet-gecorrleerde items, uit elke groep wordt zo een indicator verkregen. Zo kan elk latent construct dat staat voor een RS gevormd worden uit de indicatoren van de verschillende groepen. Deze onderliggende factoren (RS) worden dan gemodelleerd op een ander inhoudelijk latent construct (zoals in SEM). Weer wordt er dikwijls gekeken naar de RMSEA om de ‘fit’ van het model te beoordelen. Dit klinkt natuurlijk allemaal erg abstract. Om de methode beter te illustreren zal een recent onderzoek van Thomas et al. (2014) uitgelegd worden.

In dit onderzoek wordt de RIRSMACS-methode gebruikt om het verschil aan te tonen in de aanwezigheid van RS in stedelijk en landelijk gebied. Op inhoudelijk vlak worden er vier attitude constructen onderzocht: politiek cynisme, gepercipieerde discriminatie, economische veiligheid en sociaal vertrouwen. Al deze concepten worden gemeten aan de hand van een aantal variabelen. Deze variabelen worden ook door de respondenten beoordeeld aan de hand van een vijf-punten Likert schaal. De RIRSMACS-methode wordt gebruikt om een confirmatorisch statistisch model op te zetten dat vier RS opspoort: ARS, ERS, DRS en MRS. Deze RS worden dus gemodelleerd als een latent construct, die ieder gebaseerd zijn op drie indicatoren. Deze indicatoren zijn afkomstig uit drie groepen die zijn samengesteld uit niet-gecorrleerde items (Thomas et al., 2014).

¹⁵ f = frequency

In de survey waren 45 niet-gecorrleerde items opgenomen. Uit deze set konden 35 items geselecteerd worden waarbij de inter-correlaties onder de .30 bleven. Uiteindelijk werden er 27 van deze items geselecteerd die werden opgedeeld in drie groepen van telkens 9 niet-gecorrleerde items. De bedoeling was dat uit elke groep van niet-gecorrleerde items, een indicator voor elke RS werd berekend. Op deze manier had elke RS 3 indicatoren voor ze in het model toegepast werden. Net als bij SEM werd er een inhoudelijk latent construct gebruikt om de achterliggende stijlfactor te vinden die kan wijzen op de aanwezigheid van een RS. Omdat de RIRSMACS-methode gebruik maakt van deze indicatoren is het genoeg om één inhoudelijk construct te hebben in het model, in plaats van twee, wat gebruikelijk is bij SEM. In het onderzoek van Thomas et al. (2014) worden de RS gemodelleerd op het concept ‘politiek cynisme’, dat in hun dataset wordt gemeten met behulp van vier surveyitems (Thomas et al., 2014). Dit ziet er als volgend uit:

Figuur 5.1.a: RIRSMACS-methode gemodelleerd (Thomas et al., 2014)



Net als bij SEM wordt er dus naar de RMSEA gekeken om de fit te evalueren. In het bovenstaande model bijvoorbeeld werd ARS uit het model gelaten omdat het een negatieve invloed had op de RMSEA. Het model waar ARS wel in betrokken was paste goed op de data uit het landelijk gebied maar had een slechtere RMSEA bij de data uit het stedelijk gebied. Hieruit konden de onderzoekers afleiden dat er minder ARS valt vast te stellen bij respondenten uit stedelijk gebied dan bij respondenten uit landelijk gebied (Thomas et al., 2014). Een andere conclusie die getrokken wordt uit hun onderzoek is dat de aanwezigheid van RS nogal afhankelijk is van de inhoud van de constructen. Niet onlogisch aangezien reeds bevestigd werd dat een bepaalde mate van betrokkenheid bij het thema een serieuze invloed kan hebben op de aanwezigheid van RS (Billiet, McClendon, 2000).

Het uitgelegde onderzoek vormt een goed voorbeeld hoe de RIRSMACS-methode succesvol valt toe te passen op data. Het doel is nu natuurlijk om ook deze methode toe te passen op de data van SCK-Barometer 2015, om ARS aan te tonen op een andere manier dan via SEM. Daarom moeten we eerst overwegen welke surveyitems we best kunnen betrekken bij het modelleren van RS aan de hand van de RIRSMACS-methode. De bedoeling is in grote lijnen om zo veel mogelijk dezelfde items te gebruiken die we reeds gebruikten bij SEM. Dit om de vergelijking tussen de modellen zo veel mogelijk gelijk te houden.

5.2 RIRSMACS-methode toegepast op SCK-Barometer 2015

De eerste belangrijke vereisten zijn dus die niet-gecorrleerde items die aanwezig moeten zijn in de datasets. Het is dus belangrijk om er genoeg te hebben, en ze in groepen te kunnen verdelen waaruit de indicatoren berekend kunnen worden. Doordat deze niet-gecorrleerde vaak niet aanwezig zijn in dataset (omdat ze vaak inhoudelijk geen nut hebben voor de opdrachtgever) is het moeilijk voor onderzoekers om de RIRSMACS-methode toe te passen. In de SCK-Barometer 2015 zijn echter wel speciaal vijftien niet-gecorrleerde items opgenomen. Deze items zijn statements over allerlei thema's die niets met elkaar te maken hebben. Er wordt dus niet van de respondenten verwacht dat ze deze items allemaal in dezelfde stijl beantwoorden. De surveyitems moeten door de respondenten beoordeeld worden ook aan de hand van vijf-punten Likert schalen. Hieronder een overzicht van deze vijftien niet-gecorrleerde items:

Tabel 5.2.a – niet-gecorrleerde surveyitems Barometer

Variabele	Surveystatements niet-gecorrleerde items
NEO1.	Over het algemeen ben ik geen piekeraar.
NEO2.	Ik heb een heel levendige fantasie.
NEO3.	Ik ben vaak sceptisch over de bedoelingen van anderen.
NEO4.	Ik sta bekend om mijn weloverwogenheid en gezond verstand.
NEO5.	Ik houd liever de mogelijkheden open dan alles op voorhand te plannen.
NEO6.	Zonder sterke emoties zou ik het leven niet interessant vinden.
NEO7.	Sommige mensen vinden mij egoïstisch.
NEO8.	Ik probeer alle aan mij opgedragen taken gewetensvol uit te voeren.
NEO9.	Mijn stijl in werk en spel is kalm en ongehaast.
NEO10.	Ik vind dat leerlingen alleen maar in verwarring worden gebracht door ze te laten luisteren naar sprekers met afwijkende ideeën.
NEO11.	Ik zou niet genieten van een vakantie in Las Vegas.
NEO12.	Ik vind dat we in morele vraagstukken beslissingen van onze politieke leiders mogen verwachten.

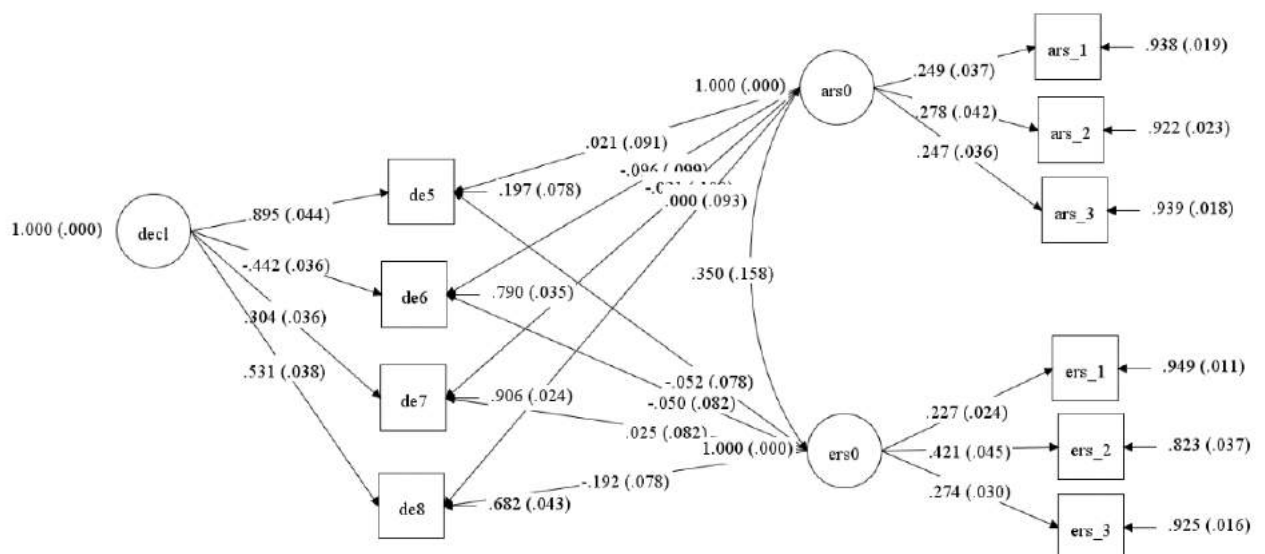
NEO13.	Met volslagen eerlijkheid kom je niet ver in het zaken doen.
NEO14.	Ik vermijd meestal films die schokkend of eng zijn.
NEO15.	Ik vind dat de wetgeving en sociaal beleid steeds moeten worden aangepast aan veranderingen in de wereld.

Hoewel het onderzoek zich focust op ARS, zullen we aan de hand van de RIRSMACS-methode twee RS modelleren. Namelijk ARS en ERS. Voor beiden RS zullen de indicatoren berekend worden uit de vijftien niet-gecorrleerde surveyitems die we opdelen in drie groepen van telkens vijf items. Deze indicatoren berekenen we als variabelen zodat elke respondent een score krijgt op elke indicator. Op deze manier kunnen we uit die variabelen, de latente factor halen die staat voor een ARS of ERS, afhankelijk van welke indicatoren natuurlijk. Deze modelleren we zoals reeds gezegd op een andere latent inhoudelijk construct dat voortkomt uit wel samenhangende variabelen

Stap 1:

Om zelf de RIRSMACS-methode toe te passen maken we gebruik van de niet-gecorrleerde items (zie tabel 5.3.a). We verdelen deze in drie blokken van telkens vijf items (NEO1 tem NEO5; NEO6 tem NEO10; NEO11 tem NEO15). Uit elke blok berekenen we een indicator (aan de hand van de eerder aangegeven formule) voor de RS die we gaan modelleren. In dit geval ARS en ERS. We hebben dus voor beide RS drie indicatoren berekend. Uit deze (2 maal) drie indicatoren halen we voor elke RS een latente factor om het te representeren. Deze twee latente factoren laten we vervolgens los op andere surveyitems van een inhoudelijk construct. In dit geval hergebruiken we het concept ‘declassering’. Wanneer we dit alles in model brengen door middel van de RIRSMACS-methode krijgen we het volgende model:

Figuur 5.2.a – RIRSMACS-model met ARS & ERS

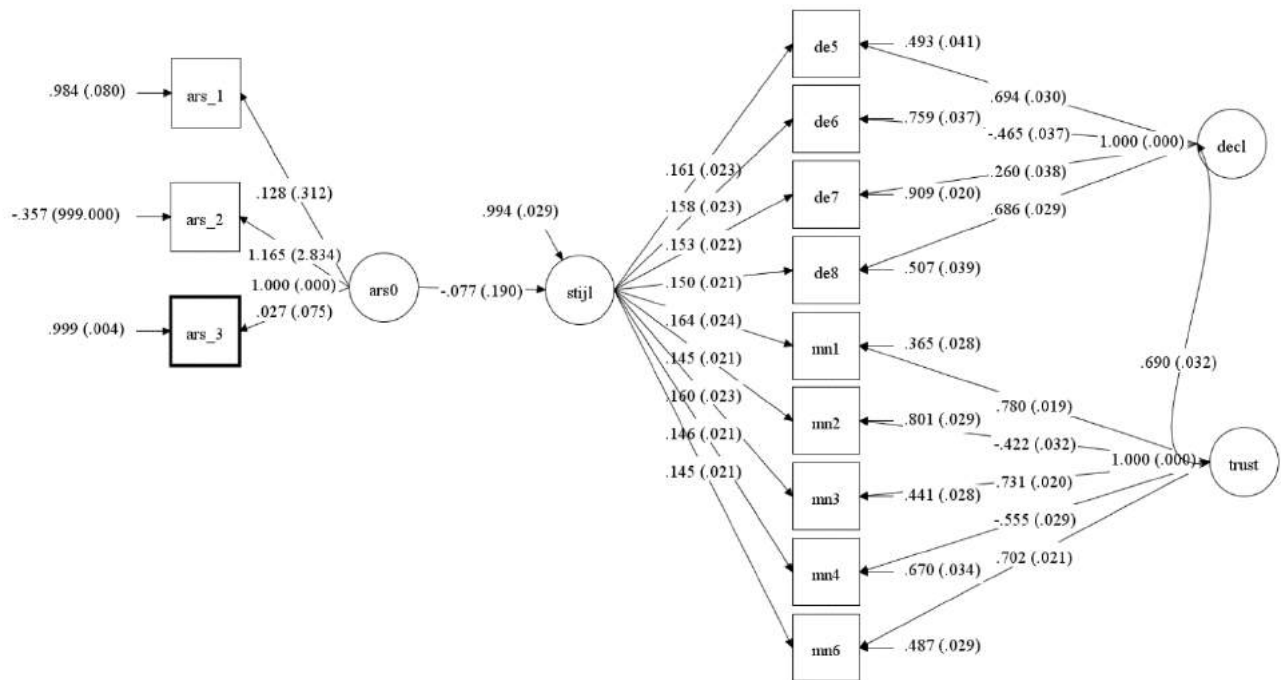


Dit is het model dat we verkrijgen na de uitvoering van de RIRSMACS-methode op onze dataset. Het model is complexer dan de modellen die we verkrijgen door SEM. Logisch aangezien we meerdere RS tegelijkertijd modelleren. Onder deze complexiteit lijdt de fit van ons model een beetje. De RMSEA van het model bedraagt 0.062, wat nog steeds een acceptabele fit is. Verder zien we tussen de inhoudelijk surveyitems over declassering en de latente ARS en ERS-variabelen, weinig significante resultaten. Net als bij het model dat we verkregen door SEM lijken de resultaten iets beter te zijn voor ARS. Wat ook opvalt is dat het ERS-construct wel een duidelijke relatie lijkt te hebben met één bepaalde variabele, namelijk DE8. Wanneer we deze variabele bekijken zien we dat het een surveyitem is waar we eerder reeds een opmerking bij plaatste. In het construct (zie tabel 3.5.a), is het namelijk het enige dat over de iets over de Belgische autoriteiten zegt. Deze bevinding bevestigt weer ons vermoeden dat de aanwezigheid van RS wel degelijk afhankelijk is van de inhoud van de surveyitems.

5.3 Combinatie van RIRSMACS-methode en SEM

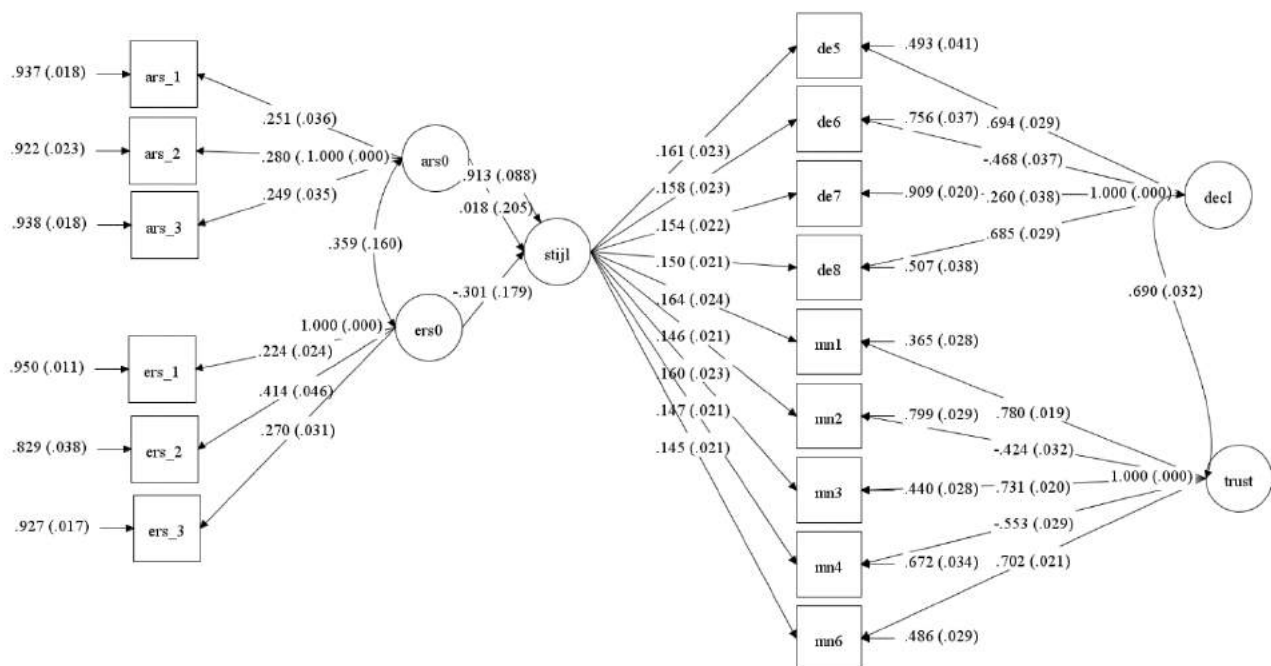
We hebben nu twee van de belangrijkste methodes om RS te modelleren uitgelegd in dit onderzoek. Hoewel de methodes hetzelfde doel hebben zijn de uitkomsten toch niet identiek. Zo lijken de uitkomsten bijvoorbeeld erg te verschillen naargelang van de surveyitems die gebruikt worden. Elke methode heeft zijn sterktes en zijn zwaktes. Daarom lijkt het een gepaste afsluiter om de twee methodes te combineren. Zo kunnen we de sterktes van beide methodes combineren en de zwaktes er trachten uit te laten. De sterktes van SEM lijken vooral te zitten in het model waar de RS in gezocht worden. In plaats van één concept te gebruiken, dat gedefinieerd wordt door enkele surveyitems, worden er meerdere gebruikt, waarachter een gemeenschappelijke stijlfactor gevonden kan worden. De items zijn gebalanceerd en de twee concepten staan inhoudelijk los van elkaar. In plaats van te kijken in hoeverre deze stijlfactor voorspeld wordt door de onafhankelijke ARS-variabele die gemaakt werd aan de hand van enkele andere surveyitems uit de vragenlijst, maken we van de ARS-variabele een latente variabele die gedefinieerd wordt door drie indicatoren. Dit gedeelte nemen we dus over van de RIRSMACS-methode, aangezien de drie indicatoren berekend werden aan de hand van de vijftien niet-gecorreleerde items. Wanneer we dit modelleren, ziet het er zo uit:

Figuur 5.3.a – combinatie van SEM en RIRSMACS-methode



We hebben nu dus een combinatie van de twee methodes. Het grote verschil met figuur 4.2.1.b (dat gecreëerd werd door middel van SEM) is dat de ARS-variaabele latent is, en dus anders berekend is. We vinden geen significante relatie tussen de twee latente constructen maar de relatie is wel groter dan met de variabele die gebruikt werd met SEM. Wanneer we de twee RS samen modelleren op deze wijze, ziet het er zo uit:

Figuur 5.3.b – Combinatie SEM en RIRSMACS-methode: ARS en ERS



De fit van model verbeterd bij het toevoegen van de latente variabele die staat voor ERS. De RMSEA bedraagt 0.049. Dit komt niet geheel verwacht maar is verklaarbaar door de duidelijk negatieve relatie tussen de stijlfactor en de ERS-variabele. Dit zou erop kunnen wijzen dat juist de tegenovergestelde stijlfactor (MLRS= het mijden van extreme antwoorden) aanwezig is in de stijlfactor. Dit verklaart in één keer de moeilijkheid om ARS te identificeren in de stijlfactor door middel van SEM. Het zou kunnen dat deze naar de achtergrond verdween door de aanwezigheid van andere RS. Dat is één van de dingen om rekening mee te houden bij het gebruik van een latente stijlfactor wanneer je SEM toepast op een dataset.

HOOFDSTUK 6 – Discussie

In dit onderzoek hebben we dus ARS opgespoord aan de hand van SEM. De methode die we gebruikt hebben is gebaseerd op de methode van Billiet en McClendon (2000) en op het ongepubliceerd onderzoek van Thijssen en Rouckhout. We hebben ARS gemodelleerd in de SCK-Barometer van 2015 die respondenten hun mening over nucleaire energie onderzoekt. Op het model zijn een verschillende demografische variabelen getest in de hoop om onze hypothesen te kunnen bevestigen of verwerpen. Hieronder een overzicht van de uitkomsten:

Tabel 6.a – Overzicht hypothesen

H1	Er is ARS vast te stellen in een sample van de heterogene Belgische bevolking.	BEVESTIGD
H2	Er valt meer ARS vast te stellen bij dubbelzinnig of moeilijk geformuleerde items dan bij duidelijk geformuleerde items.	BEVESTIGD
H3	Er is meer ARS aanwezig bij vrouwelijke respondenten dan bij mannelijke respondenten.	GEEN RELATIE
H4	Er is een positieve relatie tussen de aanwezigheid van ARS en de leeftijd van een respondent.	GEEN RELATIE
H5a	Er is een negatieve relatie tussen de aanwezigheid van ARS en het opleidingsniveau van de respondent.	GEEN RELATIE
H5b	Er is een negatieve relatie tussen de aanwezigheid van ARS en het inkomen van een respondent.	GEEN RELATIE
H6a	Er is meer ARS aanwezig bij Waalse respondenten dan bij Vlaamse respondenten.	VERWORPEN
H6b	Er is meer ARS aanwezig bij Vlaamse respondenten dan bij Waalse respondenten.	BEVESTIGD

Wat opvalt aan de resultaten is dat er met de individuele demografische variabelen weinig significante resultaten worden gevonden. Het meest verrassende resultaat komt er uit de vergelijking tussen regio's. Dit bewijst het belang aan van cross-cultureel onderzoek naar responsstijlen. Iets waar verschillende onderzoekers reeds voor geopteerd hebben (Harzing, 2006; Stenning & Everett, 1984; Thomas et al., 2014).

Het verrassende blijft echter wel dat de uitkomsten die we verkrijgen aan de hand van de regio-variabele, niet gelijk zijn aan de bevindingen van het ongepubliceerd onderzoek van Thijssen en Rouckhout. Waar zij besloten dat er meer ARS aanwezig was bij de Waalse respondenten, ondervinden we in dit onderzoek het tegenovergestelde. Hieruit rees het vermoeden dat RS afhankelijker zijn van de inhoud dan voorheen gedacht. Dit werd getest door de demografische variabele opnieuw te testen met dezelfde variabelen die gebruikt werden in het ongepubliceerd onderzoek van Thijssen en Rouckhout. Door dit te doen ondervonden we dat wanneer we de risicoperceptie-variabelen gebruikte om ARS te modelleren we inderdaad het tegenovergestelde

resultaat kregen. Niet allen ondervonden we dus dat er meer ARS was vast te stellen bij de Waalse respondenten, we stelden ook meer gewoon meer ARS vast. Hieruit besluiten we dat RS wel degelijk afhankelijk zijn van de inhoud van de surveyitems. Waarom we nu net meer ARS kunnen vaststellen bij de risicoperceptie surveyitems, laten we open. Deze puzzel laat ruimte voor een verder gedetailleerd onderzoek.

Een andere verschuiving in het onderzoek naar RS die we kunnen opmerken, is de aandacht voor het modelleren van RS. Waar de aandacht eerst lag op de theoretische kant van RS, zoals oorzaken ervan en eigenschappen van respondenten die er gevoelig aan zijn, verschuift de aandacht naar de praktijk. Er wordt meer aandacht besteed aan de methodologie. Waar men vroeger eerder RS in kaart bracht door simpele telmethodes, gebruikt men nu geavanceerdere methodes, waardoor er ook rekening wordt gehouden met de inhoudelijke covarianties. Deze verschuiving in het onderzoek toont aan dat het tijd is om RS ook in kaart te brengen bij datasets die gebruikt worden voor politiek onderzoek. Het is duidelijk in welke zin RS de resultaten vertekenen, waardoor het modelleren ervan van essentieel belang wordt.

In dit onderzoek werd er gebruikt gemaakt van twee methodes. Eerst was er *Structural Equation Modelling* (SEM). Deze methode bewees erg geschikt te zijn voor het modelleren voor RS. Een vereiste is wel dat er voor de inhoudelijke latente constructen gebalanceerde items gebruikt worden, omdat een gezamenlijke factor anders kan toegeschreven worden aan de inhoud. De andere grote moeilijkheid van SEM, is om te bepalen wat de stijlfactor werkelijk inhoudt. We gebruiken variabelen die als indicator dienen voor de verschillende RS. Het is van groot belang dat deze variabelen gecreëerd zijn aan de hand van surveyitems die nog niet in het model zijn opgenomen. Anders verkrijgen we een vertekend beeld zoals in het onderzoek van Billiet en McClendon waar de relatie tussen de ARS-variabel en de stijlfactor 0.9 bedraagt, maar waar niet erkend wordt dat dit een inhoudelijke oorzaak heeft aangezien de ARS-variabele aan de hand van diezelfde surveyitems berekend is. SEM is ook een goede methode om hypothesen te testen, aangezien er steeds variabelen aan het model kunnen worden toegevoegd en hun functie zo geëvalueerd kan worden.

Na het gebruiken van SEM, werd de RIRSMACS-methode toegepast. Deze recentere methode kan toegepast worden met dezelfde software maar is toch helemaal anders dan SEM. Het grote verschil is dat de RS niet worden afgebeeld als een latente factor (stijlfactor bij SEM). De variabelen die staan voor verschillende RS worden rechtstreeks toegepast op de surveyitems van een inhoudelijk construct. Deze variabelen worden berekend aan de hand van drie indicatoren die zelf berekend worden aan de hand van drie blokken met ieder vijf niet-gecorrleerde surveyitems in (zie tabel 5.2.a). Het feit dat de ARS-variabelen op deze manier gedefinieerd worden is één van de belangrijkste sterktes van de RIRSMACS-methode. Vandaar dat we dit stukje hebben overgenomen bij het modelleren van ARS door middel van SEM. Waar Billiet en McClendon (2000) hun ARS-variabele berekenden aan de hand

van dezelfde surveyitems die in het model gebruikt werden als inhoudelijk construct, en Thijssen en Rouckhout gebruik maakten van andere gelijkaardige surveyitems uit dezelfde dataset maar niet uit het model, gebruikt dit onderzoek de niet-gecorrleerde surveyitems om de ARS-variabele te creëren. Op deze manier wordt er eigenlijk een methode gebruikt die het beste van de twee methodes toepast.

Na dit onderzoek kunnen we wel concluderen dat het bijzonder nuttig kan zijn om RS te modelleren in allerlei soorten onderzoek, omdat ze resultaten kunnen vertekenen. Mits het naleven van enkele regels lijkt SEM het meest geschikt te zijn voor het modelleren van RS. Het is wel aan te raden om gebalanceerde surveyitems te gebruiken, en de ARS-variabele te creëren aan de hand van niet-gecorrleerde variabelen, om een *least-likely case* te hebben.

Met het oog op politieke communicatie, die meestal erg in de greep is van de publieke opinie, is het essentieel om RS op een correcte manier te kunnen meten, om ze zo in rekening te brengen. Zoals reeds gezegd is er een verdere focus op de relatie tussen RS en de inhoud van surveyitems gewenst in toekomstig onderzoek. Op deze manier kan er ook naar een gepaste manier worden gezocht om onderzoek te doen naar de publieke mening over politieke concepten, wat zo belangrijk is in het onderzoeksveld van de politieke communicatie.

Bibliografie

- Austin, E. J., et al. (2006). "Individual differences in response scale use: Mixed Rasch modelling of responses to NEO-FFI items." *Personality and Individual Differences* 40(6): 1235-1245.
- Baumgartner, H. and J.-B. E. Steenkamp (2001). "Response styles in marketing research: A cross-national investigation." *Journal of marketing research* 38(2): 143-156.
- Billiet, J. B. and M. McClendon (1998). "On the identification of acquiescence in balanced set items using structural models." *Advances in Methodology, Data Analysis, and Statistics. Metodološki zvezki* 14: 129-150.
- Billiet, J. B. and M. J. McClendon (2000). "Modeling acquiescence in measurement models for two balanced sets of items." *Structural Equation Modeling* 7(4): 608-628.
- Cheung, M. W.-L. and W. Chan (2002). "Reducing uniform response bias with ipsative measurement in multiple-group confirmatory factor analysis." *Structural Equation Modeling* 9(1): 55-77.
- Crandall, J. E. (1982). "Social interest, extreme response style, and implications for adjustment." *Journal of Research in Personality* 16(1): 82-89.
- Das, J. and T. Dutta (1969). "Some correlates of extreme response set." *Acta psychologica* 29: 85-92.
- De Beuckelaer, A., et al. (2010). "Using ad hoc measures for response styles: A cautionary note." *Quality & Quantity* 44(4): 761-775.
- De Vellis, R. F. and L. S. Dancer (1991). "Scale Development: Theory and Applications." *Journal of Educational Measurement* 31(1): 79-82.
- Glynn, C. J., et al. (2015). *Public opinion*, Westview Press.
- Greenleaf, E. A. (1992). "Improving rating scale measures by detecting and correcting bias components in some response styles." *Journal of Marketing Research* 29(2): 176.
- Greenleaf, E. A. (1992). "Measuring extreme response style." *Public Opinion Quarterly* 56(3): 328-351.

- Hair, J. F., et al. (2006). *Multivariate data analysis*, Pearson Prentice Hall Upper Saddle River, NJ.
- Hamilton, D. L. (1968). "Personality attributes associated with extreme response style." *Psychological bulletin* 69(3): 192.
- Harzing, A.-W. (2006). "Response Styles in Cross-national Survey Research A 26-country Study." *International Journal of Cross Cultural Management* 6(2): 243-266.
- Jackson, D. N. and S. Messick (1958). "Content and style in personality assessment." *Psychological bulletin* 55(4): 243.
- Johnson, T., et al. (2005). "The relation between culture and response styles evidence from 19 countries." *Journal of Cross-cultural psychology* 36(2): 264-277.
- McBride, L. and G. Moran (1967). "Double agreement as a function of item ambiguity and susceptibility to demand implications of the psychological situation." *Journal of personality and social psychology* 6(1): 115.
- McClendon, M. (1992). On the measurement and control of acquiescence in latent variable models. annual meeting of the American Sociological Association.
- McDonald, R. P. and M.-H. R. Ho (2002). "Principles and practice in reporting structural equation analyses." *Psychological methods* 7(1): 64.
- McDonald, R. P. and M.-H. R. Ho (2002). "Principles and practice in reporting structural equation analyses." *Psychological methods* 7(1): 64.
- Meisenberg, G. and A. Williams (2008). "Are acquiescent and extreme response styles related to low intelligence and education?" *Personality and Individual Differences* 44(7): 1539-1550.
- Mirowsky, J. and C. E. Ross (1991). "Eliminating defense and agreement bias from measures of the sense of control: A 2 x 2 index." *Social Psychology Quarterly*: 127-145.
- Moors, G. (2008). "Exploring the effect of a middle response category on response style in attitude measurement." *Quality & Quantity* 42(6): 779-794.

- Moors, G. (2010). "Ranking the ratings: A latent-class regression model to control for overall agreement in opinion research." *International Journal of Public Opinion Research* 22(1): 93-119.
- Ray, J. J. (1979). "Is the acquiescent response style problem not so mythical after all? Some results from a successful balanced F scale." *Journal of Personality Assessment* 43(6): 638-643.
- Ray, J. J. (1983). "Reviving the problem of acquiescent response bias." *The Journal of Social Psychology* 121(1): 81-96.
- Rorer, L. G. (1965). "The great response-style myth." *Psychological bulletin* 63(3): 129.
- Schuman, H. and S. Presser (1981). "Questions and answers." *Experiments in the*.
- Schuman, H. and J. Scott (1987). "Problems in the use of survey questions to measure public opinion." *Science* 236(4804): 957-959.
- Steiger, J. H. (2007). "Understanding the limitations of global fit assessment in structural equation modeling." *Personality and Individual Differences* 42(5): 893-898.
- Stening, B. W. and J. E. Everett (1984). "Response styles in a cross-cultural managerial study." *The Journal of Social Psychology* 122(2): 151-156.
- Thomas, T. D., et al. (2014). "Measurement Invariance, Response Styles, and Rural–Urban Measurement Comparability." *Journal of Cross-cultural psychology*: 0022022114532359.
- Tomarken, A. J. and N. G. Waller (2005). "Structural equation modeling: Strengths, limitations, and misconceptions." *Annu. Rev. Clin. Psychol.* 1: 31-65.
- Van Vaerenbergh, Y. and T. D. Thomas (2012). "Response styles in survey research: A literature review of antecedents, consequences, and remedies." *International Journal of Public Opinion Research*: eds021.
- Weijters, B., et al. (2009). "Response Styles and how to Correct them." *GfK Marketing Intelligence Review* 1(2): 44-53.

Bijlagen

1. Tabellen

HOOFDSTUK 1

/

HOOFDSTUK 2

Tabel 2.4.2.a - Overzicht hypothesen:

H1	Er is ARS vast te stellen in een sample van de heterogene Belgische bevolking.
H2	Er valt meer ARS vast te stellen bij dubbelzinnig-geformuleerde items dan bij duidelijk geformuleerde items.
H3	Er is meer ARS aanwezig bij vrouwelijke respondenten dan bij mannelijke respondenten.
H4	Er is een positieve relatie tussen de aanwezigheid van ARS en de leeftijd van een respondent.
H5a	Er is een negatieve relatie tussen de aanwezigheid van ARS en het opleidingsniveau van de respondent.
H5b	Er is een negatieve relatie tussen de aanwezigheid van ARS en het inkomen van een respondent.
H6a	Er is meer ARS aanwezig bij Vlaamse respondenten dan bij Waalse respondenten.
H6b	Er is meer ARS aanwezig bij Waalse respondenten dan bij Vlaamse respondenten.

HOOFDSTUK 3

Figuur 3.4.a - surveyitems 'empathie'

variabele	Surveyitems 'empathie'
Qe_1_res	Ik heb vaak tedere, bezorgde gevoelens voor mensen die minder geluk hebben dan ik. (P)
Qe_6_res	Ik ben vaak geraakt door dingen die ik zie gebeuren. (P)
Qe_3_res	Soms heb ik niet veel medelijden met andere mensen wanneer ze problemen hebben. (N)
Qe_5_res	Andermans ongeluk stoort me meestal niet erg. (N)

Tabel 3.4.b – surveyitems 'persoonlijkheid'

variabele	Surveyitems 'persoonlijkheid'
Qe_4_res	Ik probeer bij een meningsverschil ieders standpunt te bekijken alvorens ik een beslissing neem. (P)

Qe_7_res	Ik geloof dat er twee kanten zijn aan elke vraag en probeer naar beide te kijken. (P)
Qe_8_res	Alvorens iemand te bekritisieren probeer ik mij voor te stellen hoe ik mij zou voelen mocht ik in zijn/haar plaats zijn. (P)

Figuur 3.4.c - surveyitems 'stralingsrisico'

variabele	Surveyitems 'stralingsrisico'
Qrp_2_resp	Het gezondheidsrisico voor een doorsnee Belg - Radioactief aval
Qrp_5_resp	Het gezondheidsrisico voor een doorsnee Belg - Een ongeluk in een nucleaire installatie
Qrp_9_resp	Het gezondheidsrisico voor een doorsnee Belg - Een terroristische aanslag met een radioactieve bron

Tabel 3.4.d – surveyitems 'gebruiksrisico'

variabele	Surveyitems 'gebruiksrisico'
Qrp_6_resp	Het gezondheidsrisico voor een doorsnee Belg - Straling van GSM toestellen
Qrp_7_resp	Het gezondheidsrisico voor een doorsnee Belg - Natuurlijke straling
Qrp_8_resp	Het gezondheidsrisico voor een doorsnee Belg - Röntgenfoto's in de geneeskunde

Tabel 3.5.a – surveyitems 'declassering'

Variabele	Surveystatements 'declassering'
DE5	De nucleaire industrie beschikt over de vereiste technologie om kerncentrales succesvol te declasseren. (P)
DE6	De nucleaire industrie heeft niet de vereiste expertise om kerncentrales succesvol te kunnen declasseren. (N)
DE7	De eigenaar van de kerncentrales heeft de benodigde financiële middelen om kerncentrales te declasseren. (P)
DE8	Ik vertrouw het toezicht van de Belgische autoriteiten op de declasseringswerkzaamheden van de nucleaire industrie. (P)

Tabel 3.5.b – factorladingen 'declassering'

variabele	Surveyitems 'declassering'	factorlading
DE5	De nucleaire industrie beschikt over de vereiste technologie om kerncentrales succesvol te declasseren. (P)	,821
DE6	De nucleaire industrie heeft niet de vereiste expertise om kerncentrales succesvol te kunnen declasseren. (N)	,830

DE7	De eigenaar van de kerncentrales heeft de benodigde financiële middelen om kerncentrales te declasseren. (P)	,762
DE8	Ik vertrouw het toezicht van de Belgische autoriteiten op de declasseringswerkzaamheden van de nucleaire industrie. (P)	,768
DE9	Zelfs na declassering van een kerncentrale, blijft de site gevaarlijk voor meerdere decennia. (N)	,774

Tabel 3.5.c – surveyitems ‘wetenschap en technologie’

Variabele	Surveystatements ‘wete’
AX2	Wetenschap en technologie zullen zorgen voor hogere levenskwaliteit voor de toekomstige generaties. (P)
AX3	Wetenschap en technologie maken ons leven gemakkelijker. (P)
AX6	Wetenschap en technologie hebben de meeste mensen ongelukkiger gemaakt. (N)
AX7	Wetenschappelijke en technologische ontwikkeling hebben onvoorziene kwalijke effecten voor de gezondheid van mens en milieu. (N)

Tabel 3.5.d – factorladingen ‘wetenschap en technologie’

Variabele	Surveystatement ‘wete’	Factorlading
AX2	Wetenschap en technologie zullen zorgen voor hogere levenskwaliteit voor de toekomstige generaties. (P)	,752
AX3	Wetenschap en technologie maken ons leven gemakkelijker. (P)	,754
AX6	Wetenschap en technologie hebben de meeste mensen ongelukkiger gemaakt. (N)	,796
AX7	Wetenschappelijke en technologische ontwikkeling hebben onvoorziene kwalijke effecten voor de gezondheid van mens en milieu. (N)	,730

HOOFDSTUK 4

Tabel 4.2.1a– surveyitems ‘trust’

variabele	Surveystatements ‘trust’
MN1	Nucleaire reactoren in België worden op een veilige manier bestuurd. (P)
MN2	De autoriteiten controleren de veiligheid van de nucleaire installaties in België onvoldoende. (N)
MN3	In België wordt er veilig omgegaan met radioactief afval. (P)
MN4	Het transport van nucleaire materialen is onveilig. (N)
MN5	Ik voel me goed beschermd tegenover nucleaire installaties.(P)

Tabel 4.2.2.a – frequentietabel gender-variabele

	Aantal	Percentage
Mannen	513	49,9 %
Vrouwen	515	50,1 %
Totaal	1028	100 %

Tabel 4.2.5.a – frequentietabel regio-variabele

	Aantal	Percentage
Vlaming	605	58.9
Waal	339	33.0
Brusselaar	84	8.2
Totaal	1028	100.0

HOOFDSTUK 5

Tabel 5.3.a – niet-gecorrleerde surveyitems Barometer

Variabele	Surveystatements niet-gecorrleerde items
NEO1.	Over het algemeen ben ik geen piekeraar.
NEO2.	Ik heb een heel levendige fantasie.
NEO3.	Ik ben vaak sceptisch over de bedoelingen van anderen.
NEO4.	Ik sta bekend om mijn weloverwogenheid en gezond verstand.
NEO5.	Ik houd liever de mogelijkheden open dan alles op voorhand te plannen.
NEO6.	Zonder sterke emoties zou ik het leven niet interessant vinden.
NEO7.	Sommige mensen vinden mij egoïstisch.
NEO8.	Ik probeer alle aan mij opgedragen taken gewetensvol uit te voeren.
NEO9.	Mijn stijl in werk en spel is kalm en ongehaast.
NEO10.	Ik vind dat leerlingen alleen maar in verwarring worden gebracht door ze te laten luisteren naar sprekers met afwijkende ideeën.
NEO11.	Ik zou niet genieten van een vakantie in Las Vegas.
NEO12.	Ik vind dat we in morele vraagstukken beslissingen van onze politieke leiders mogen verwachten.
NEO13.	Met volslagen eerlijkheid kom je niet ver in het zaken doen.
NEO14.	Ik vermijd meestal films die schokkend of eng zijn.
NEO15.	Ik vind dat de wetgeving en sociaal beleid steeds moeten worden aangepast aan veranderingen in de wereld.

HOOFDSTUK 6

Tabel 6.a – Overzicht hypotheses

H1	Er is ARS vast te stellen in een sample van de heterogene Belgische bevolking.	BEVESTIGD
H2	Er valt meer ARS vast te stellen bij dubbelzinnig of moeilijk geformuleerde items dan bij duidelijk geformuleerde items.	BEVESTIGD
H3	Er is meer ARS aanwezig bij vrouwelijke respondenten dan bij mannelijke respondenten.	GEEN RELATIE
H4	Er is een positieve relatie tussen de aanwezigheid van ARS en de leeftijd van een respondent.	GEEN RELATIE
H5a	Er is een negatieve relatie tussen de aanwezigheid van ARS en het opleidingsniveau van de respondent.	GEEN RELATIE
H5b	Er is een negatieve relatie tussen de aanwezigheid van ARS en het inkomen van een respondent.	GEEN RELATIE
H6a	Er is meer ARS aanwezig bij Vlaamse respondenten dan bij Waalse respondenten.	BEVESTIGD
H6b	Er is meer ARS aanwezig bij Waalse respondenten dan bij Vlaamse respondenten.	VERWORPEN

2. Figuren

HOOFDSTUK 1

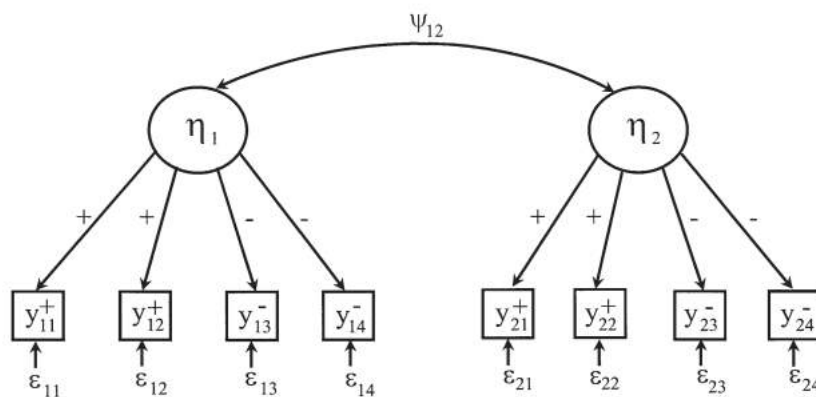
/

HOOFDSTUK 2

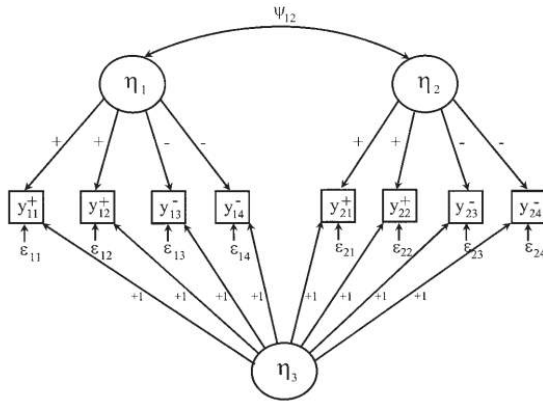
/

HOOFDSTUK 3

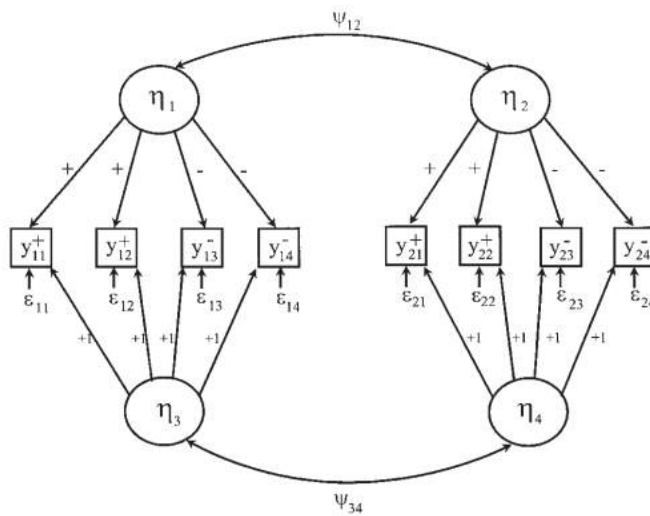
Figuur 3.3.a: Twee inhoudsfactoren en geen stijlfactor (bron: Billiet McClendon, 2000)



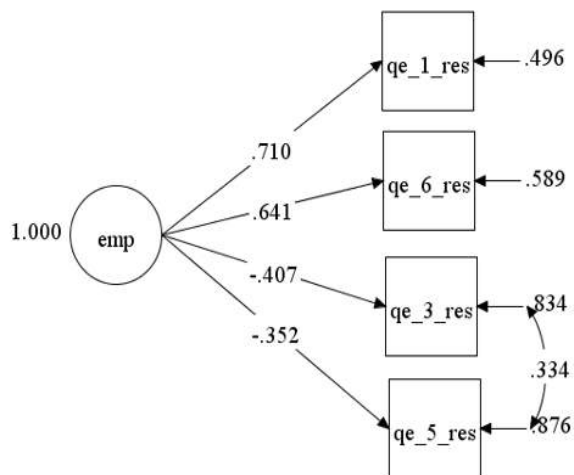
Figuur 3.3.b: twee inhoudsfactoren en één stijlfactor (Bron: Billiet, McClendon, 2000)



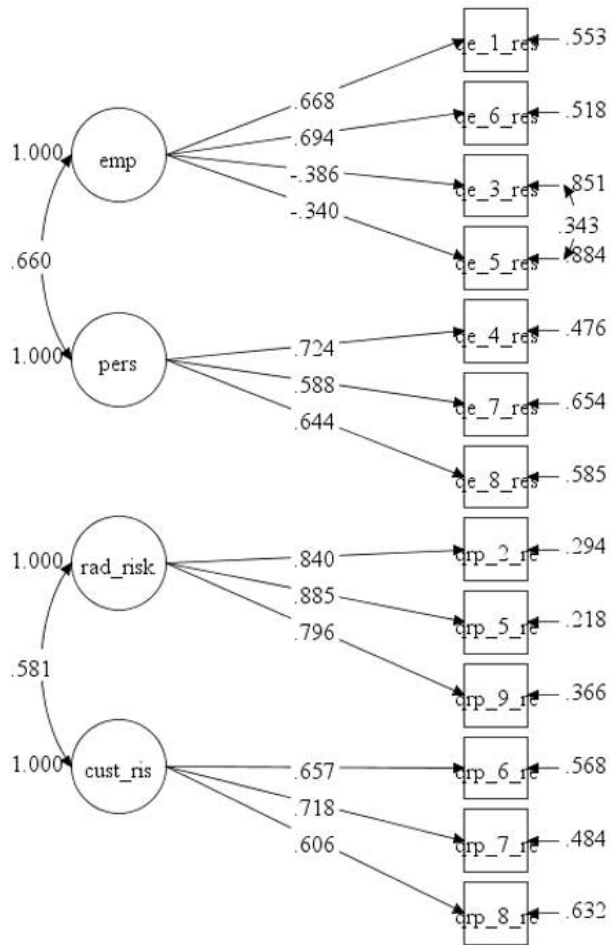
Figuur 3.3.c : twee inhoudsfactoren en twee stijlfactoren (Bron: Billiet, McClendon, 2000)



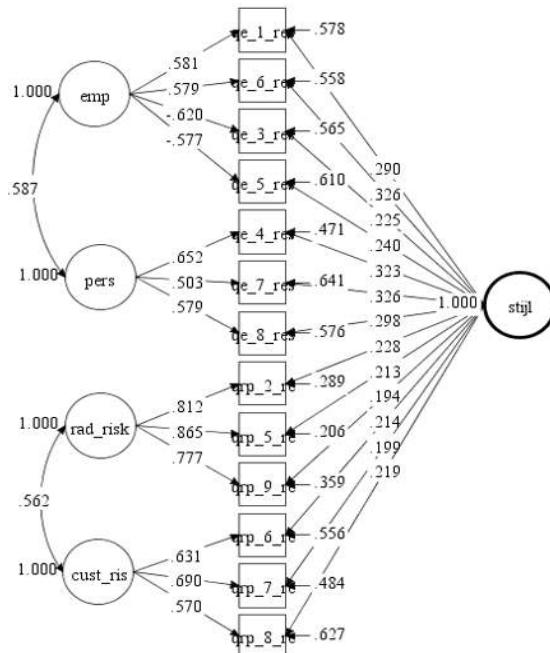
Figuur 3.4.a – empathie als latente variabele (ongepubliceerd onderzoek Thijssen en Rouckhout)



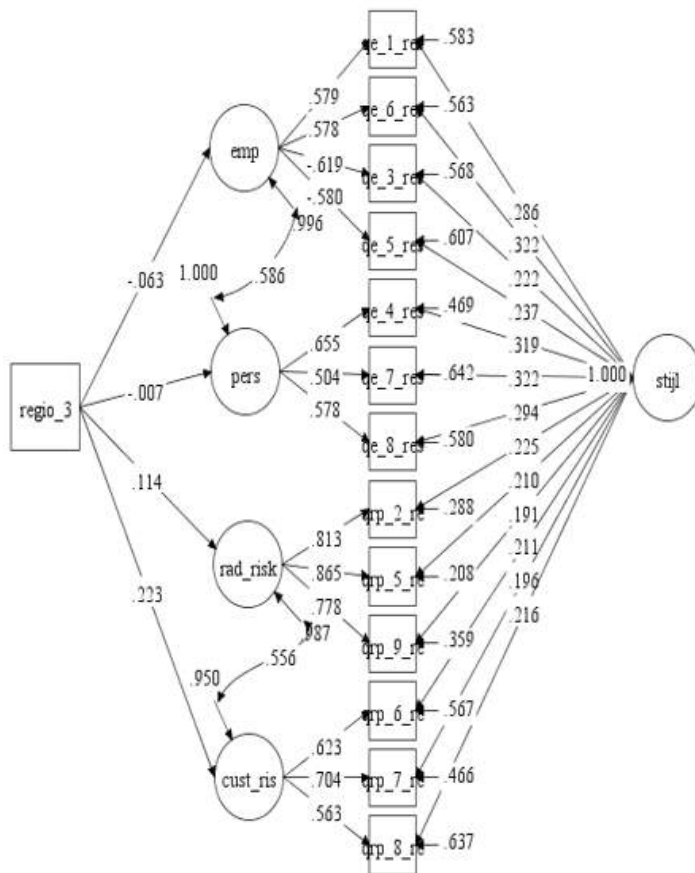
Figuur 3.4.b – Alle inhoudelijke constructen gemodelleerd (ongepubliceerd onderzoek Thijssen en Rouckhout)



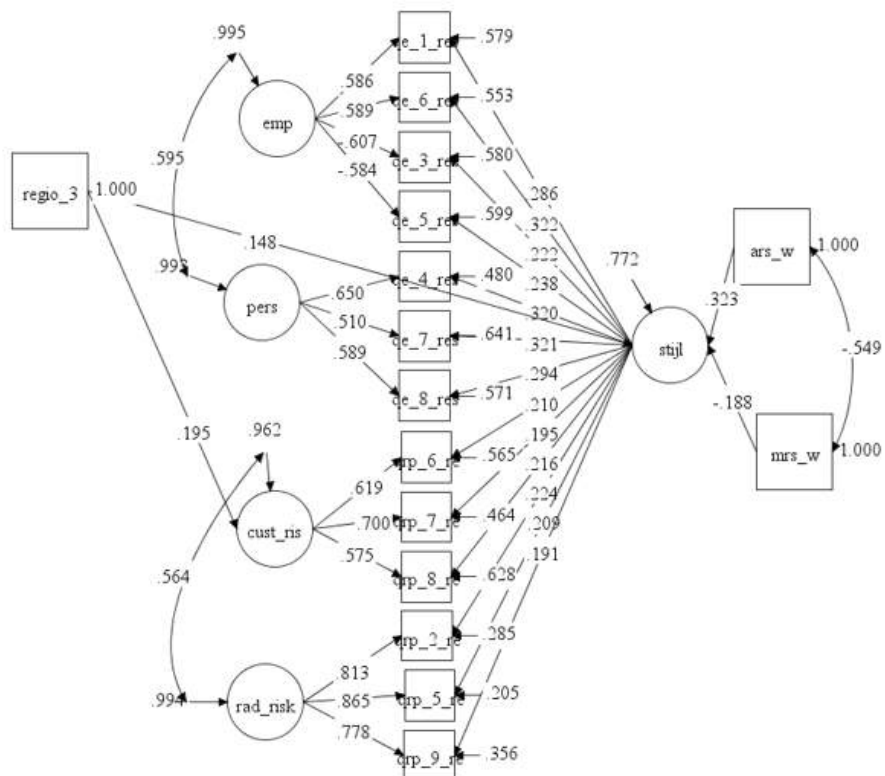
Figuur 3.4.c – stijlfactor toegevoegd aan inhoudelijke constructen (ongepubliceerd onderzoek Thijssen en Rouckhout)



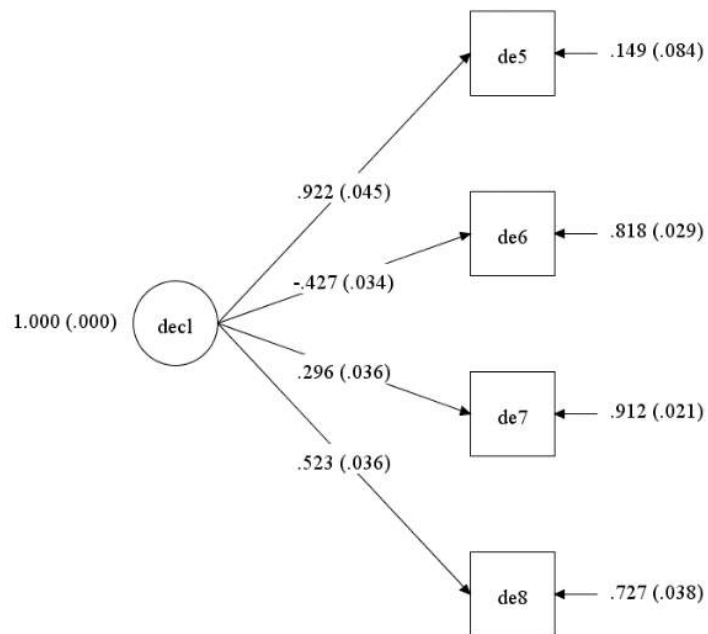
Figuur 3.4.d – regio-variabele toegevoegd aan het model (ongepubliceerd onderzoek Thijssen en Rouckhout)



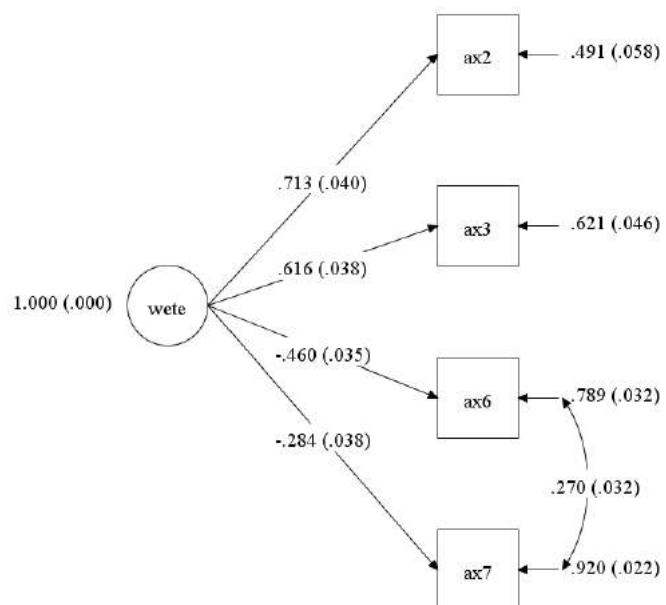
Figuur 3.4.e – Volledig eindmodel van het ongepubliceerd onderzoek van Thijssen en Rouckhout



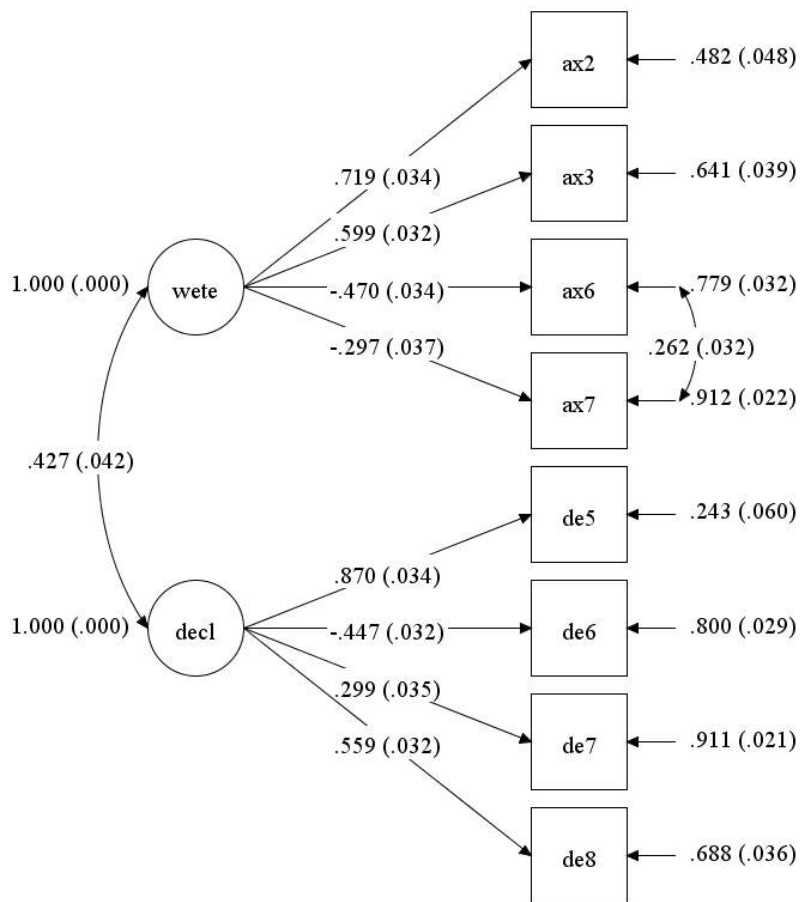
Figuur 3.5.a – latente variabele ‘declassering’



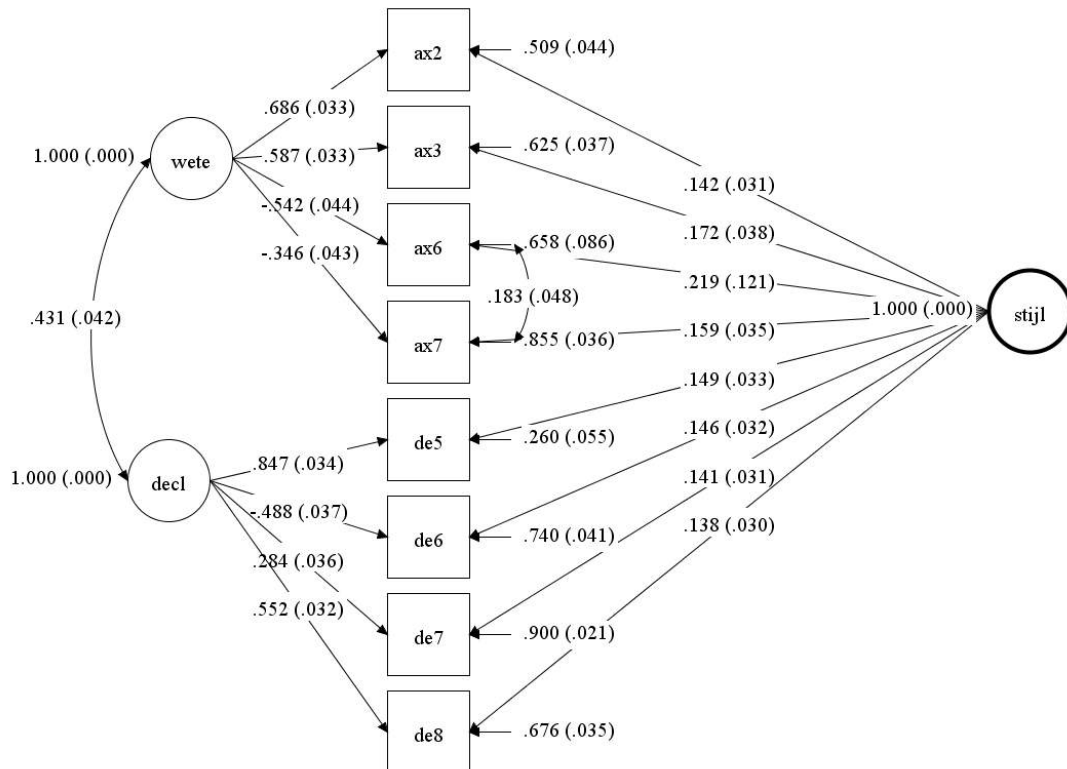
Figuur 3.5.b – latente variabele ‘wetenschap en technologie’



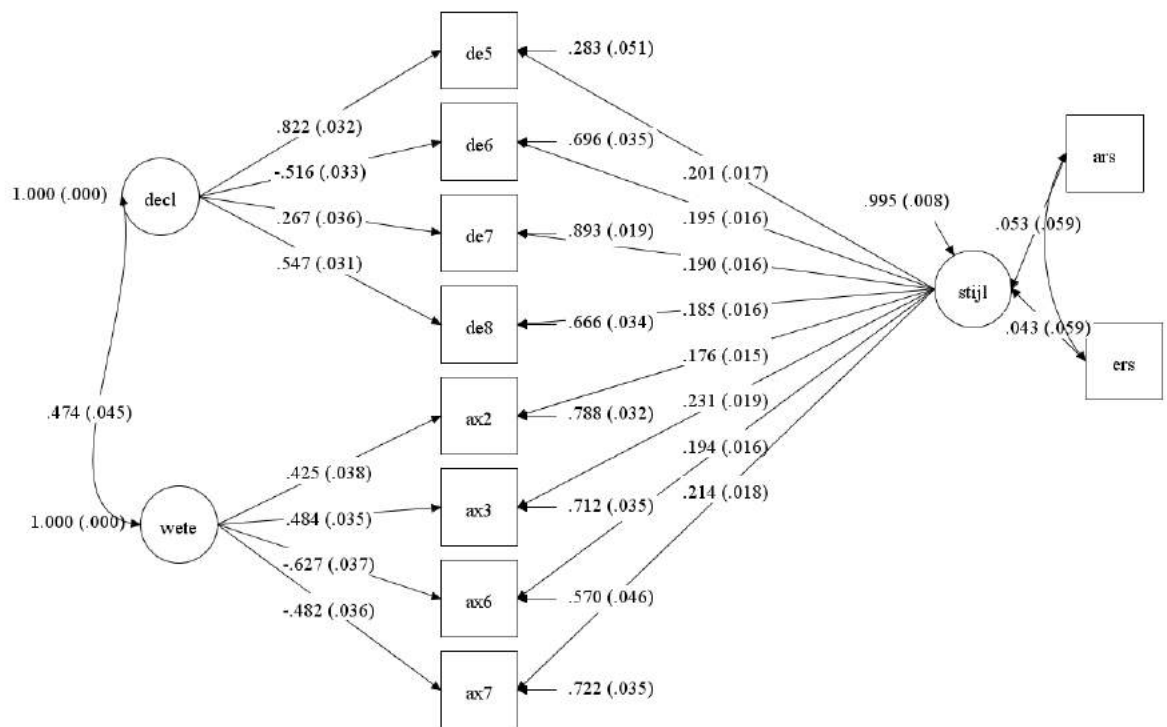
Figuur 3.5.c – inhoudelijke constructen gemodelleerd



Figuur 3.5.d – model met een stijlfactor toegevoegd

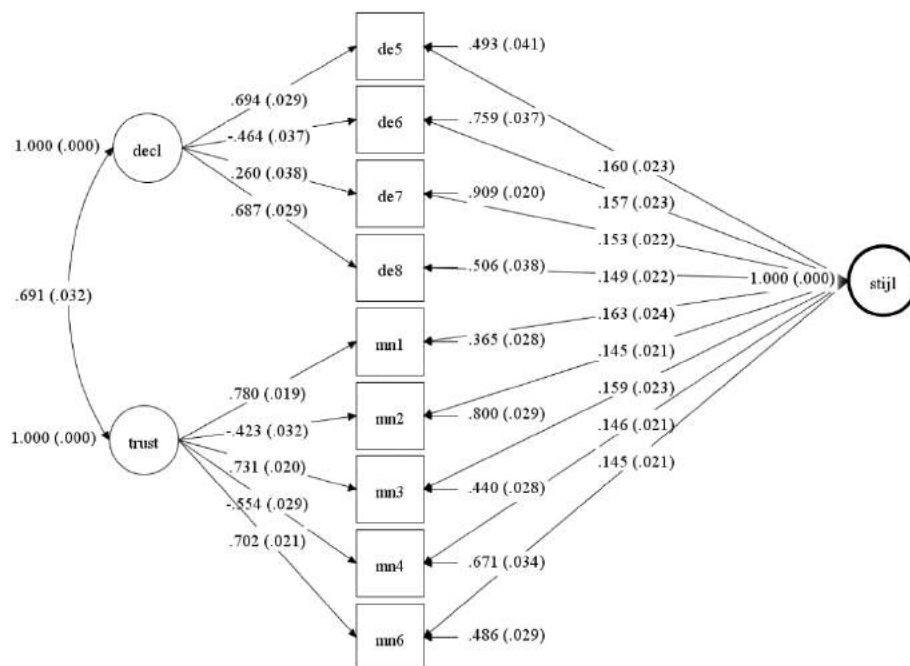


Figuur 3.5.e – model met ARS en ERS-variabele

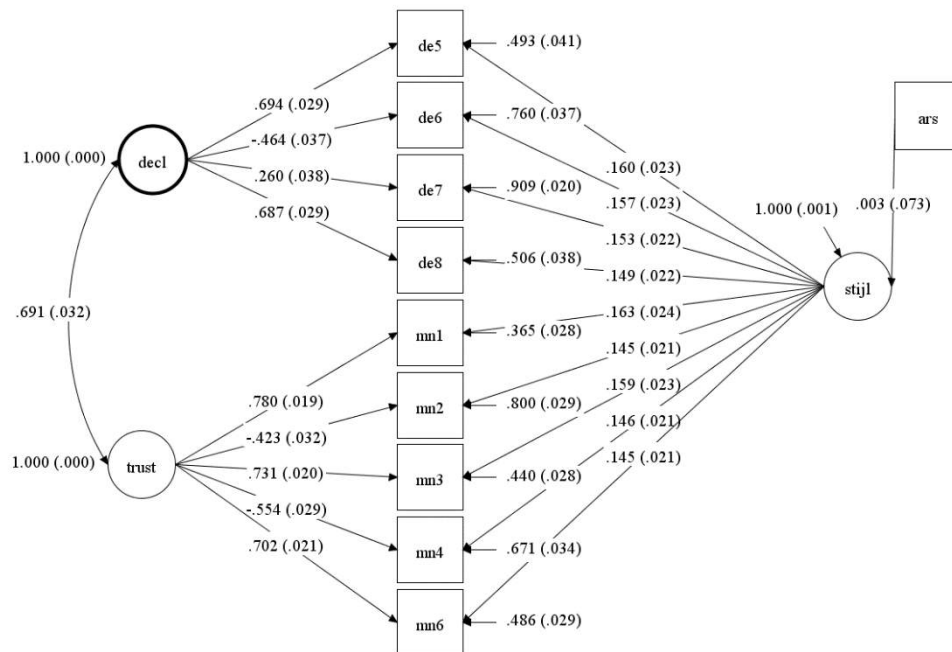


HOOFDSTUK 4

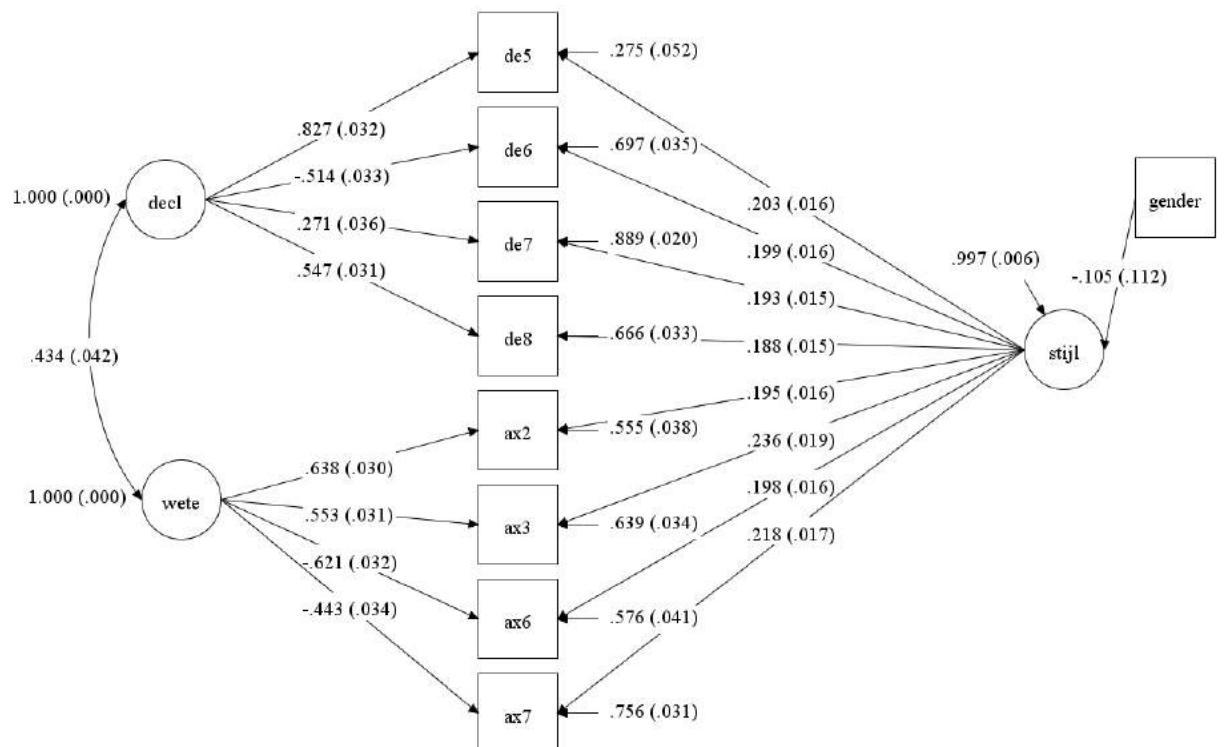
Figuur 4.2.1.a – model met ‘trust’ + stijlfactor



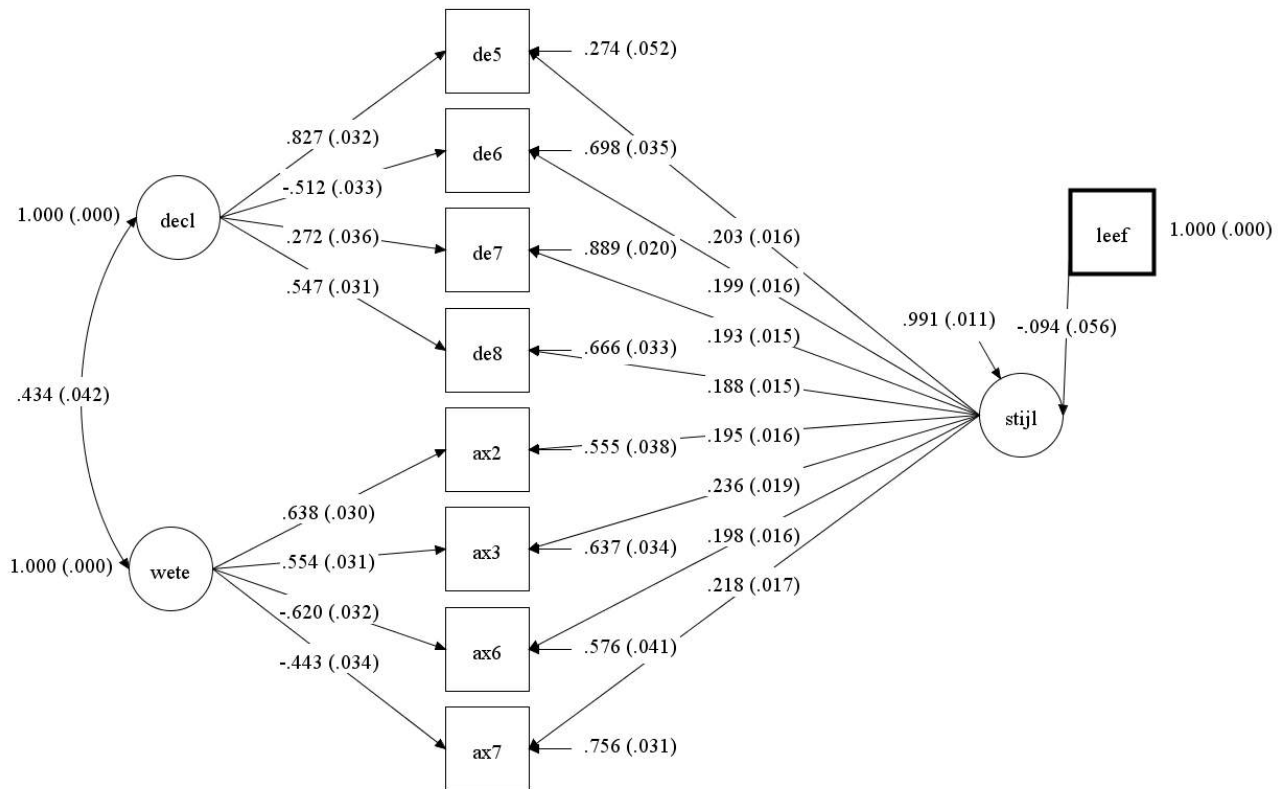
Figuur 4.2.1.b – model met ‘trust’ + ARS-variabele



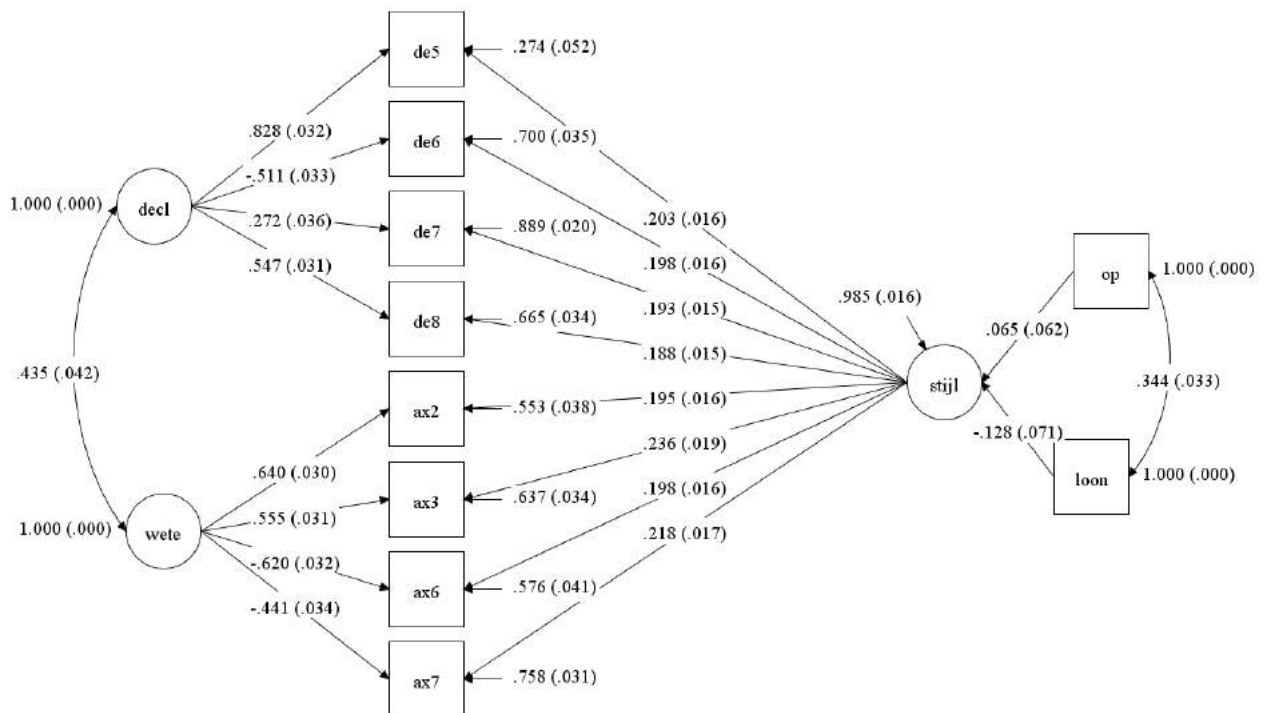
Figuur 4.2.2.a – model met de gender-variabele



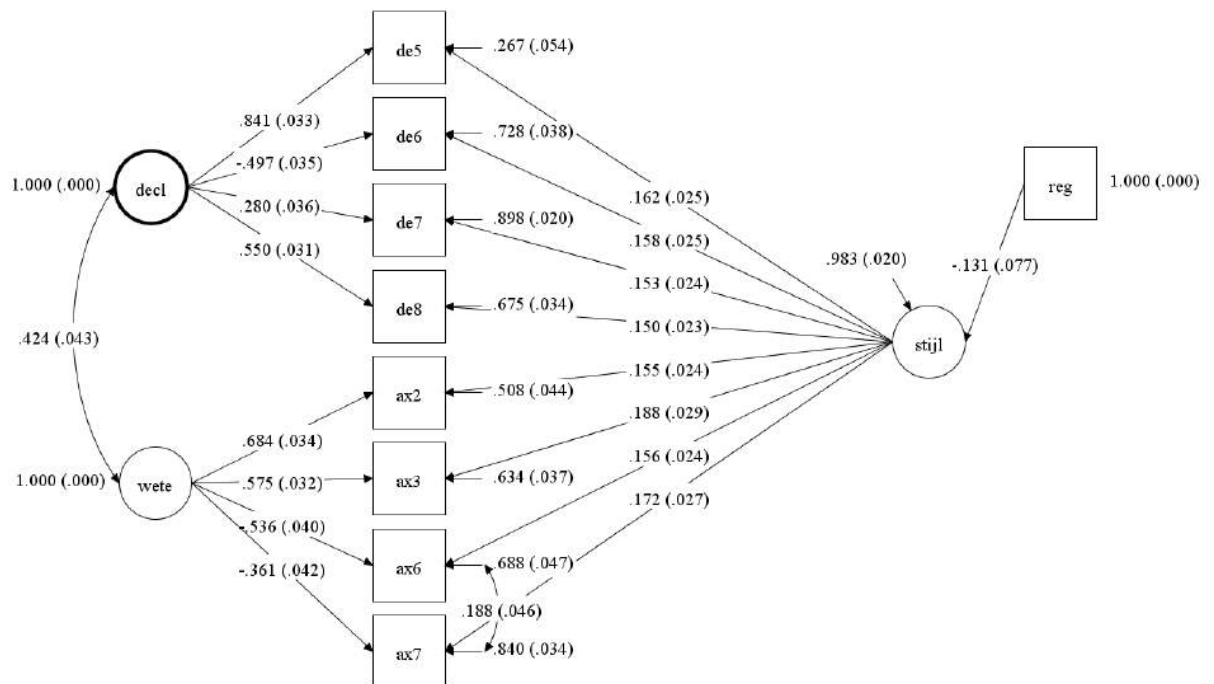
Figuur 4.2.3 – model met de leeftijd-variabele



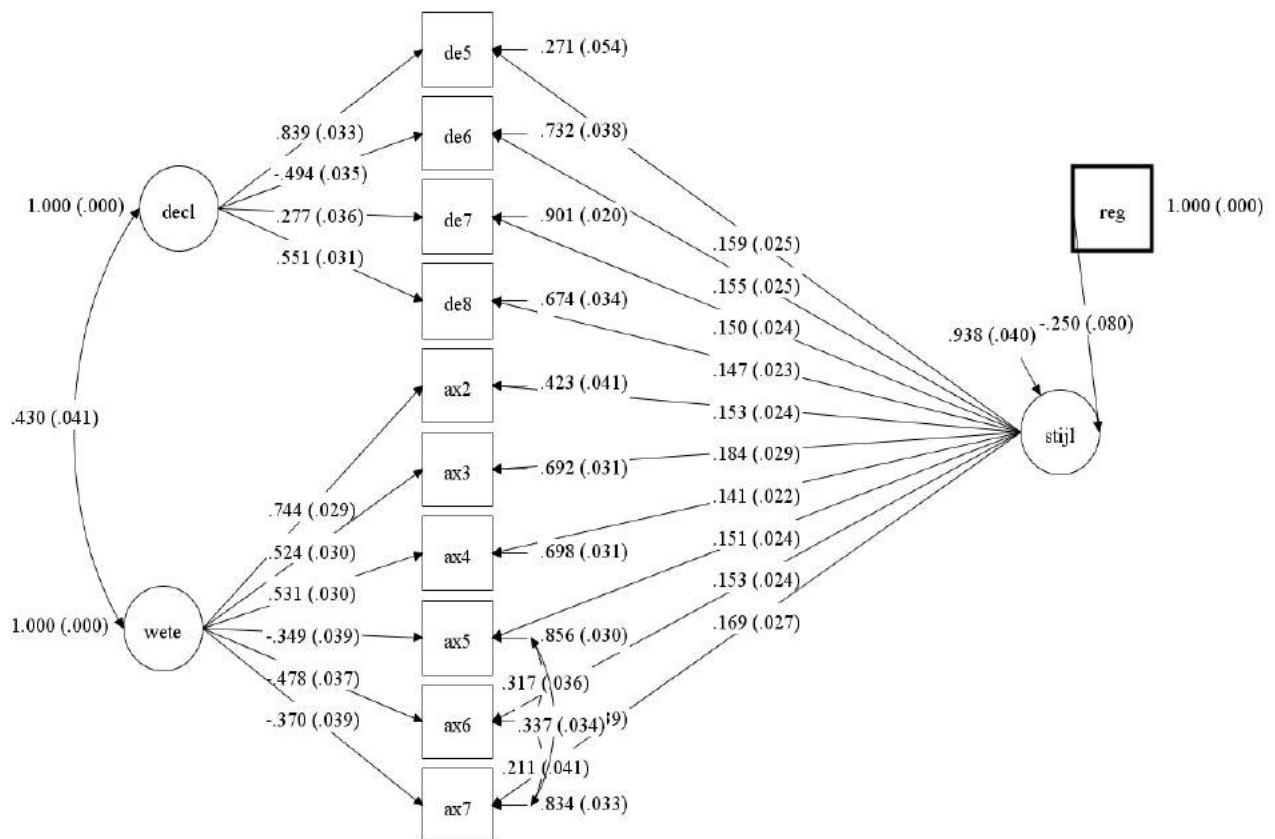
Figuur 4.2.4.a – model met loon en opleidingsvariabelen



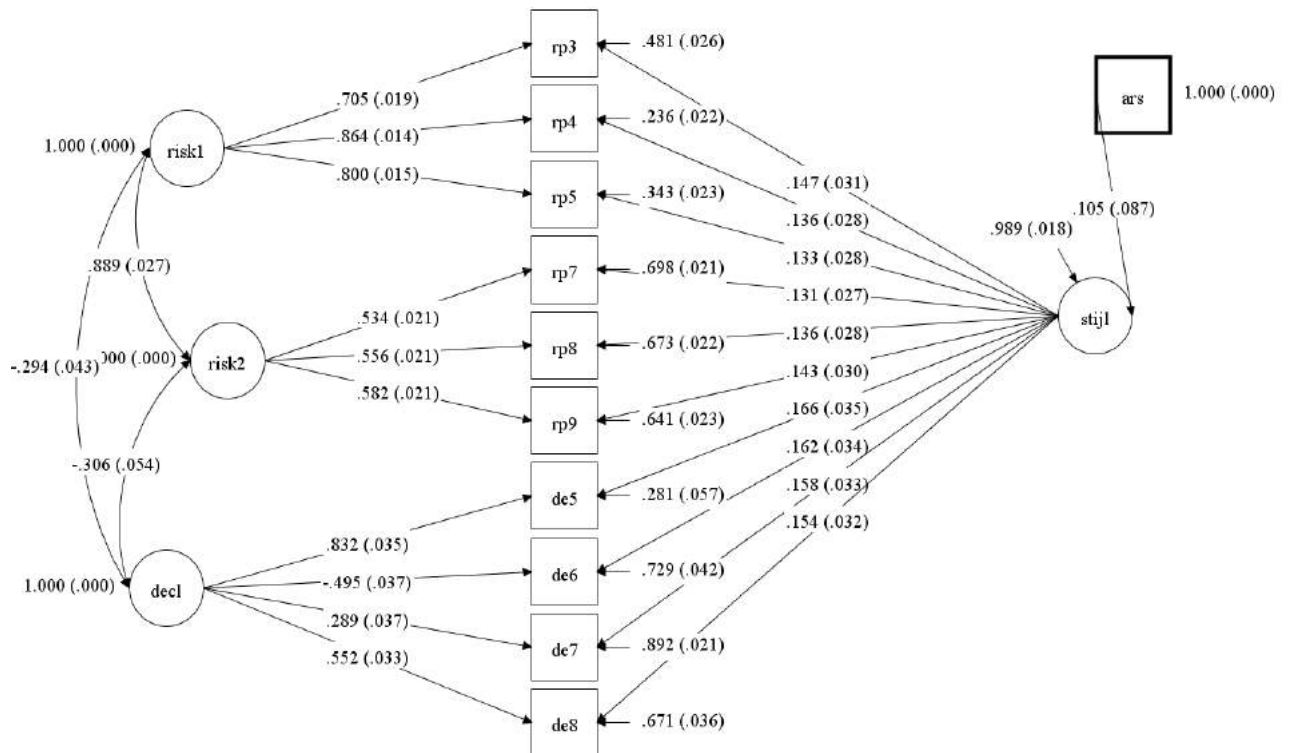
Figuur 4.2.5.a – model met de regio-variabele



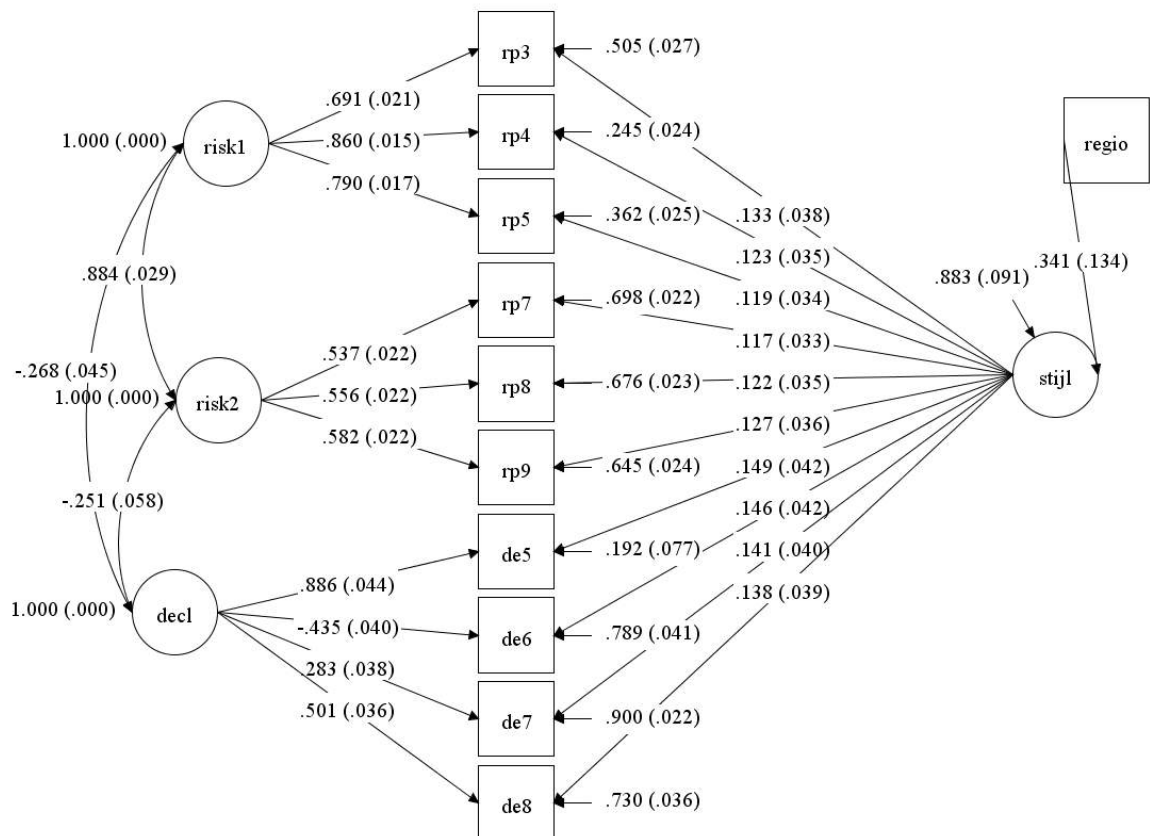
Figuur 4.2.5.b – model met de regio-variabele toegevoegd op groter model



Figuur 4.2.5.c – Model met risicoperceptie items

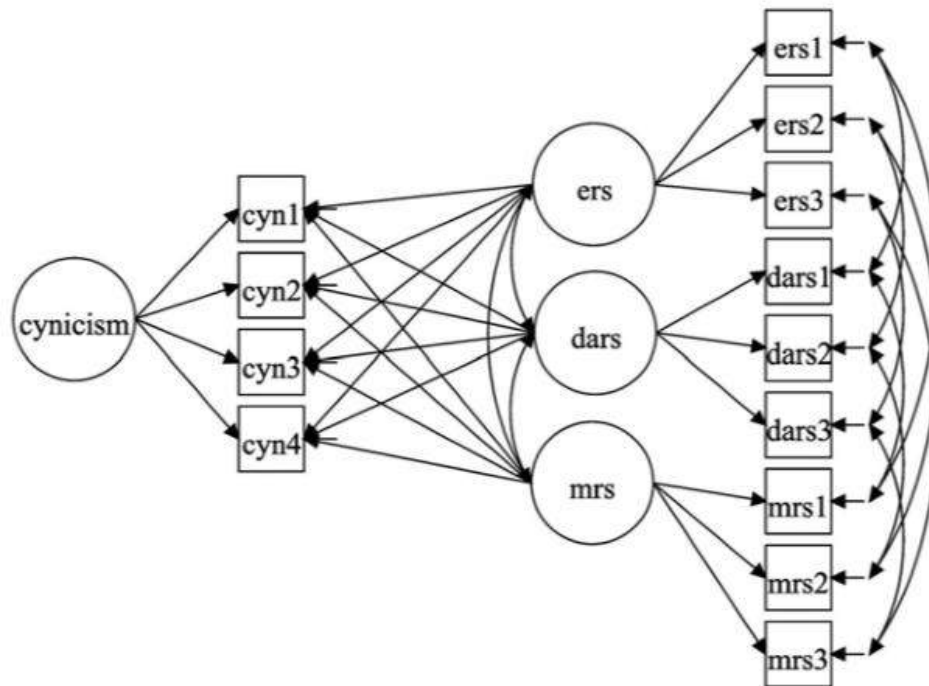


Figuur 4.2.5.d – Model met risicoperceptie surveyitems + regio-variabele

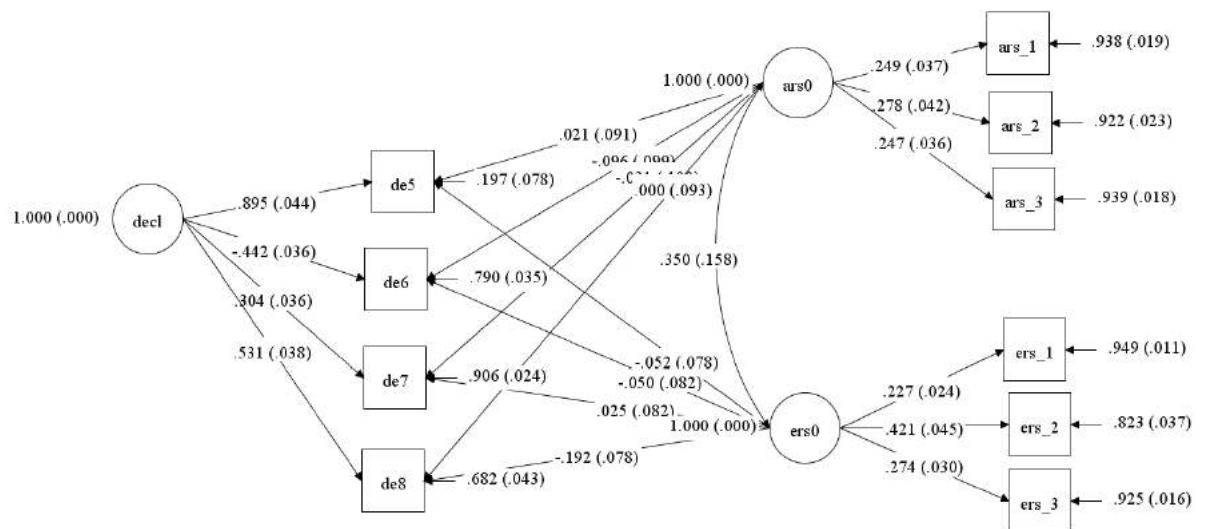


HOOFDSTUK 5

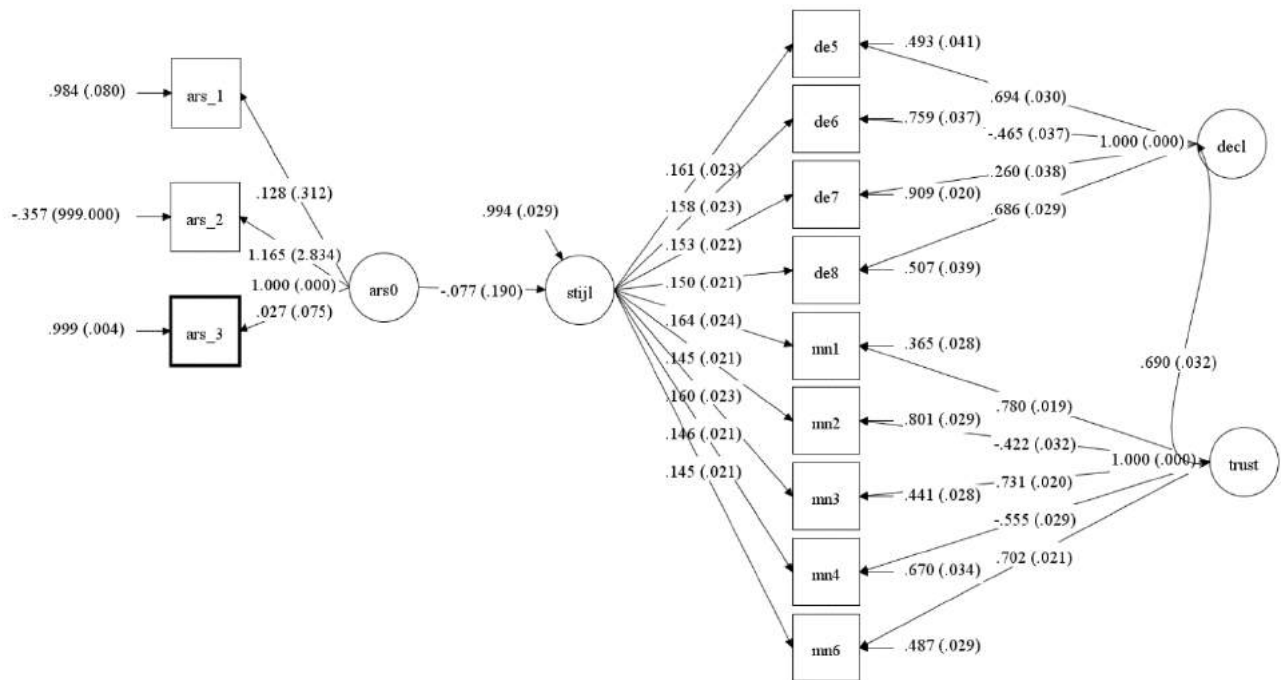
Figuur 5.1.a: RIRSMACS-methode gemodelleerd (Thomas et al., 2014)



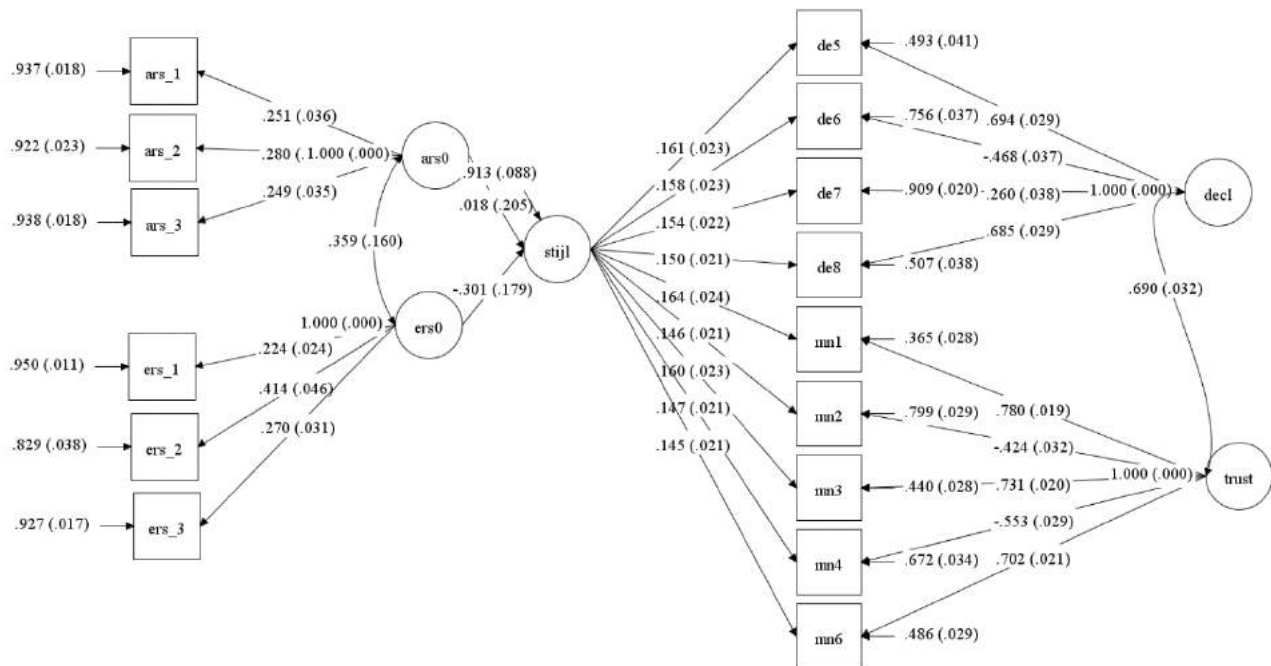
Figuur 5.2.a – RIRSMACS-model met ARS & ERS



Figuur 5.3.b – combinatie van SEM en RIRSMACS-methode



Figuur 5.3.b – Combinatie SEM en RIRSMACS-methode: ARS en ERS



HOOFDSTUK 6

/

