

Een engagerende discussie

Statistische analyse & connectionistisch modellering van de homogeniseringshypothese.

Thema 3 – Afsluitend verslag

28 Maart 2008

Tomas Zwinkels

Gilles de Hollander

Tom Aizenberg

Inhoudsopgave

1. Inleiding	3
1.1 Homogeniseringstheze	3
1.2 Computatieel model.....	3
1.3 Wetenschappelijke relevantie.....	4
1.4 Methodologisch kader.....	4
1.5 Opzet van paper	5
1.6 Dankwoord	5
2. Methode	5
2.1 Methode empirisch onderzoek.....	5
2.1.1 Onderzoeksgroep	5
2.1.2 Tokens	5
2.1.3 Verloop van de sessie	6
2.1.4 Ruwe data en verwerking.....	6
2.1.5 Video-opnames	7
2.1.6 Lijsten uit de associatietaak.....	7
2.1.7 Samentrekken vergelijkbare begrippen	7
2.1.8 Statistische toetsing	8
3. Resultaten empirisch onderzoek	9
4. Methode en opzet computationeel model	10
4.1 Vergelijking model empirie.....	12
4.1.1 Dynamische drempelmethod.....	12
5. Resultaten computationele model	13
5.1 Homogenisering in het computationele model	14
6. Discussie.....	14
6.1 De mens als associatieve machine?	14
6.2 Na-ijl effect	15
6.3 Vervolgonderzoek.....	15
7. Literatuurlijst.....	16

Samenvatting | In dit paper laten wij zien op welke manier de cognitief-psychologische stroming van het connectionisme interessant is voor symbolische interactionistisch sociologen en andere wetenschappers die geïnteresseerd zijn in sociale interactie en de invloed hiervan op denkbeelden van mensen.

Wij onderzochten de empirische houdbaarheid van de 'homogeniseringstheze'. Deze these stelt dat denkbeelden van mensen door interactie meer op elkaar gaan lijken en vormt een belangrijke onderlegger voor het werk van oa de sociologen Schutz, Giddens, Garfinkel, Mead, Beck, Parsons & Foucault.

Wij hebben gevonden dat een groepsdiscussie over Lord of the Rings leidt tot een relatieve toename van 100% in de homogeniteit van associaties met betrekking tot Lord of the Rings onder deelnemers aan deze discussie. Bovendien blijkt er, bij de door proefpersonen als belangrijk aangemerkte associaties, sprake te zijn van een sterk na-ijl effect. Twee weken na de discussie nam de homogeniteit nog eens 53% toe waarmee de totale homogeniteitstoename op de belangrijke concepten op 180% uitkomt.

Deze groepsdiscussies zijn tevens gesimuleerd met een zeer basaal, op een connectionistische paradigma gebaseerd, computationeel model. Dit basale model blijkt op een aantal specifieke parameters opvallend goed in staat op de invloed van een groepsdiscussie op de denkbeelden van de deelnemers te voorspellen.

Abstract | In this paper we intent to show in which way the connectionist paradigm is interesting for symbolic interactionists and other scientists with an interest for social interaction and its influence one people's cognitive representations.

We describe the way in which the 'thesis of homogenization' survives empirical testing. This thesis states that people's cognitive representation moves towards a more homogenous view under the influence of social interaction. Ideas of this sort form an important assumption within the work of sociologist like Schutz, Giddens, Garfinkel, Mead, Beck, Parsons & Foucault.

We found that a group discussion about 'Lord of the Rings' led to a relative increase of 100% in de homogeneity of associations considering Lord of the Rings within the participants in this discussion. Next to this there seems to be a continuing effect at least until two weeks after the discussion. Two weeks after the discussion the homogeneity increased with another 53%. The complete increase in homogeneity two weeks after the discussion comes to 180%.

We simulated these groups discussion with a very basic, connectionist based, computational model. This model is already surprisingly capable in predicting changes in people's views under the influence of group discussions.

1. Inleiding

In dit paper wordt verslag gedaan van ons onderzoek naar de invloed van een groepsdiscussie over een bepaald onderwerp op de associaties van de deelnemers met betrekking tot dit onderwerp. Tevens wordt een op een connectionisch paradigma gebaseerd computationeel model van groepsdiscussies gepresenteerd. Het onderzoek staat in theoretisch kader van wat wij de 'homogeniseringsthese' noemen. Deze hypothese vormt een belangrijk impliciet fundament onder een breed scala aan sociologische (en in minder mate psychologische) theorieën. In dit onderzoek is deze these voor het eerst empirische getoetst.

1.1 Homogeniseringsthese

De homogeniseringsthese stelt dat individuen na 'interactie' een *homogene* definitie van de situatie of een 'homogener waarheidsbeeld' hebben dan voorafgaand aan interactie.

Veel sociologen gaan er vanuit dat er sprake is van een bepaalde 'sociale gedeeldheid van kennis' die een gevolg is van interactie. Zo spreekt de Schutz over waarheidsconstructie als een 'product van collectiviteiten' (Wallace & Wolf, 2006; p264). Anthony Giddens spreekt in deze context over '*Mutual knowledge*' (Giddens, 1984; p1-25), Garfinkel over '*common sense knowledge*' (Wallace & Wolf, 2006; p267-280). Deze begrippen hebben sterke overeenkomsten met Mead zijn idee van de '*generalized other*' (Wallace & Wolf, 2006; p271). Ulrich Beck spreekt over '*handelingskennis*' (Beck, Giddens & Scott, 1994). Deze sociologen beschrijven, ieder op hun eigen manier, de aanwezigheid van een soort van 'collectief cognitief schema'. Dit schema bevat kennis over hoe men zal (of dient te) handelen in een specifieke context (Mead, Beck) of is gevuld met aannemens over 'hoe de wereld in elkaar zit'. Deze kennis is sociaal gedeeld (mutual knowledge) en is zelf een 'product van collectiviteiten'. Deze kennis wordt gevormd tijdens sociale interactie. Verschillende sociologen hebben tevens verschillende woorden gebruikt om het interactie proces waarin deze collectieve cognitieve schema tot stand komt te beschrijven. Men spreekt over '*bracketing*' (Schutz) '*framed*' of simpelweg '*definiëring*'; de functionalistische socioloog Talcot Parsons spreekt over '*internalisering*'. Ook het '*discour*' begrip van Foucault beschrijft een 'gezamenlijke manier van spreken' over een specifiek onderwerp. Dit behelst tevens een 'gezamenlijke manier van denken', oftewel een bepaalde mate van overlap tussen de cognitieve schema's die mensen hanteren om over de wereld na te denken. Foucault legt hierin overigens sterk de nadruk op de machtswerking die van deze 'talende praktijken' uitgaat (Willems, 1989). De

sociologen Berger en Luckman noemt dit proces de '*externalisation*'. Tijdens dit proces vindt "de sociale productie van symbolische structuren plaats". Deze symbolische structuren krijgen een gezamenlijke betekenis voor hen die aan het vormingsproces hebben deelgenomen (Wallace & Wolf, 2006; p287). De meer analytische netwerktheoreticus Granoveter stelt eenvoudigweg dat 'het vergaren van kennis via interactie plaatsvindt' (Granovetter, 1974).

Terugkerend element in al deze theorieën is steeds dat na interactie meer kennis 'gedeeld' wordt dan voorafgaand aan de interactie, oftewel: de kennis homogeniseert door interactie. Dit element noemen wij de homogeniseringsthese.

In de antropologie is een uitgebreide traditie van kwalitatief onderzoek naar de invloed van sociale interactie op 'issue framing' oftewel: sociale probleemdefinitie. Een ideaal typische voorbeeld van dergelijke onderzoek is bijvoorbeeld de antropologische studie: 'how issues get framed and reframed when different communities meet: A Multi-level Analysis of a Collaborative Soil Conservation Initiative in the Ecuadorian Andes' (Dewulf. et al, 2004).

Ook in de psychologie vinden we werk dat de homogeniseringsthese beschrijft. Zo spreekt Gray (2002, blz. 517-531) in 'Psychology' over de '*social-adjustive function*', wat inhoudt dat mensen om sociaal geaccepteerd te worden hun attitudes naar de wereld aanpassen aan de omgeving (lees: groep mensen) waarin zij zich bevinden. Hij koppelt dit aan 'cognitieve dissonantie'¹: mensen hebben graag een consistent wereld beeld van 'attitudes' en informatie of gedrag dat hier niet consistent is wordt dan ook als 'dissonant' en onprettig ervaren. Het gevolg hiervan is dat mensen hun attitudes/wereldbeeld steeds meer op elkaar afstemmen. Het collectieve wereldbeeld wordt dan dus steeds homogener.

1.2 Computationeel model

De homogenisering van waarheidsbeelden kan naar ons inzicht mogelijk goed gemodelleerd worden met een

¹ Sinds Leon Festinger in 1950 kwam met zijn 'cognitieve dissonantie theorie'. Dit een belangrijk begrip uit de cognitieve psychologie gebleven. Festinger stelt dat het menselijke brein een intern mechanisme kent wat zorgt voor een 'onaangenaam gevoel van dissonantie', wanneer het individu zich bewust wordt van bepaalde inconsistenties in de eigen houdingen, kennis, overtuigingen en handelingen (Gray, 2002).

connectionistisch model. In dit model staat het 'associatieve' of *connectionistische* karakter van de menselijke cognitie en het geheugen (en daarmee de cognitieve verwerking van groepsdiscussies) centraal. Na een uiteenzetting van de empirische toetsing van de homogeniseringsthese volgt dan ook toelichting op een op een connectionistisch paradigma gebaseerd computationele model van groepsdiscussies.

Onderzoek binnen het vakgebied van de cognitieve psychologie maakt het aannemenlijk dat discussies associatief verlopen. Met name het zogenaamde 'priming'-onderzoek en onderzoeksvormen die hierop voortborduren (Zie onder andere Eagle et al (1966) in Gray et al., 2002) hebben aangetoond dat mensen op individueel niveau associatief denken. Het menselijk geheugen is zo opgebouwd dat associatief sterk gelinkte concepten eerder, in het individu naar boven zullen komen dan minder sterk gelinkte concepten.

Overeenkomstig met Carley (Carley, 1993) verstaan wij in dit paper onder een concept een opzichzelfstaand idee of een 'te onderscheiden object', dit hoeft niet per definitie één woord te zijn. Voorbeeld van concepten zijn bijvoorbeeld 'groepsdiscussie' of 'connectionistische model' of 'corrumperende werking van macht'.

Het tweede deel van het verslag beschrijft een door onszelf ontwikkeld, connectionistisch model van groepsdiscussies dat de specifieke, woordelijke, veranderingen in individuele associaties naar aanleiding van de groepsdiscussie simuleert.

1.3 Wetenschappelijke relevantie

De uitkomsten van het hieronder beschreven onderzoek zijn voornamelijk interessant omdat ze een nieuw licht werpen op kennisuitwisseling tijdens menselijke interactie.

Vergelijkbaar onderzoek; waarin computationele modellen gebruikt worden om groepsdiscussies te simuleren focust zich voornamelijk op de belangen en doelen van aan de discussie deelnemende actoren. Een voorbeeld hiervan is het soort onderzoek dat bekend staat onder de noemer 'group discussion modeling'. De nadruk ligt in deze modellen niet op kennisuitwisseling maar op machtsrelaties. In het u voorliggende onderzoek is bewust gekozen voor een neutraal onderwerp (Lord of the Rings) in plaats van een discussie waarin machtsrelaties en discussieuitkomsten centraal staan. Omdat de onderzoekers van het 'group decision modeling' zich voornamelijk gericht hebben op de uitkomsten van de discussie is er maar weinig bekend over kennisuitwisseling tijdens deze discussies.

Er is bij ons geen vergelijkbaar onderzoek bekend. Emails naar verschillende kenners binnen het vakgebied van de sociale en cognitieve psychologie en een uitgebreide zoektocht in de onderzoeksdatabase van het KNAW, hebben geen onderzoek naar kennisuitwisseling tijdens groepsdiscussies met behulp van computationele modellen opgeleverd. Ook een empirische toetsing van de homogeniseringshypothese is (voor zo ver wij hebben kunnen achterhalen) niet eerder uitgevoerd.

De uitkomsten van dit onderzoek zijn interessant voor sociologen en sociaal-psychologen die onderzoek doen vanuit een symbolische interactionistisch (of anderzijds rondom interactie en taal georiënteerd) theoretische kader. Niet alleen de validiteit van de homogeniseringsthese wordt getoetst, maar ook de mate van homogenisering en het in hoeverre de effecten van interactie na een tijd stand houden, is door ons onderzocht. Tevens vormt het paper de eerste voorzichtige stapjes in een nieuw, op het connectionistisch paradigma gebaseerd, onderzoeksgebied naar kennisuitwisseling tijdens sociale interactie met behulp van computationele modellen.

1.4 Methodologisch kader

Methodologische gezien heeft het uitgevoerde onderzoek het meeste affiniteit met de geautomatiseerde 'content analyse' of 'map analyse'. Pas tijdens de afronding van ons onderzoek ontdekte wij het methodologische werk van het CASOS ²(Carley, 1993). Hieronder worden enkele interessante aanknopingspunten tussen hun werk en de door ons ontwikkelde methodologie en algoritmes uiteengezet.

Carley geeft in 'A comparison of content analysis en map analysis' een uitgebreide uiteenzetting van belangrijke praktische kennis (tacit knowledge) voor wetenschappers die met behulp van 'woordtel-software' onderzoek doen naar veranderingen in de betekenis van teksten.

Haar onderzoeksobject is dus niet, zoals bij ons, groepsdiscussies maar geschreven teksten. Toch is er een sterke overeenkomst met het door ons uitgevoerde onderzoek. Zet zij geschreven teksten veelal om in woordenlijsten om automatische analyse mogelijk te maken. In ons onderzoek maken wij ook gebruik van woordenlijsten. Deze komen echter niet uit een tekst maar worden door proefpersonen zelf opgeschreven. Aan deze benadering ligt een gezamenlijk idee ten grondslag dat deze lijsten van

² Center for Computational Analysis of Social and Organizational Systems, zie <http://www.casos.cs.cmu.edu/>

‘kernconcepten’ iets zeggen over betekenis van deze teksten. Ook wij oriënteren een ‘waarheidsbeeld’ rondom kernconcepten. Bij de ‘map- analyse’ representeren deze onderzoekers kennis aan de hand van de semantische netwerken die ‘in de tekst verborgen zitten’. Tevens onderzoeken ze betekennisverschuivingen in deze semantische netwerken over de tijd. Wij maken gebruik van een, wat recentere, benadering van kennisrepresentatie namelijk de connectionistische netwerken. Dit connectionistische paradigma is echter wel ontstaan uit de traditie van de semantische analyses. (Medler, 1998) Vanwege de overeenkomsten met dit vakgebied zullen wij tijdens de methodologische besprekingen in dit paper soms verwijzen naar analoge begrippen en methodes uit de content- en map analyse.

1.5 Opzet van paper

Om de homogeniseringsthese te kunnen toetsen hebben wij een uitgebreide onderzoeksmethodiek met bijbehorende algoritmes ontwikkeld. Voor de statistische analyse hebben we gebruik gemaakt van bestaande methodes. De manier waarop de homogeniseringsthese is getoetst, staat beschreven in het eerste deel van het tweede hoofdstuk. De resultaten van dit empirische onderzoek worden in het derde hoofdstuk uiteengezet.

Daarop volgt in het vierde hoofdstuk een beschrijving van de opzet van het computationele model en in het vijfde hoofdstuk worden ook de resultaten van het computationele model beschreven en worden deze resultaten vergeleken met de resultaten uit de empirie. Vanwege het exploratieve karakter van dit onderzoek wordt het stuk afgesloten met een uitgebreide discussie. Hierin worden de gevonden resultaten bediscussieerd en worden voorstellen gedaan voor vervolgonderzoek.

1.6 Dankwoord

Onze grote dank gaat uit naar Maarten van Someren, Jaap Murre en Bart van Heerikhuizen voor de begeleiding bij onze respectievelijke vakgebieden: kunstmatige intelligentie, psychobiologie en sociologie. Ook veel dank gaat uit naar de interdisciplinaire begeleiding vanuit het IIS in de persoon van Machiel Keestra, Martin Boeckhout en Dirk Haen. Ook willen wij graag onze proefpersonen bedanken.

2. Methode

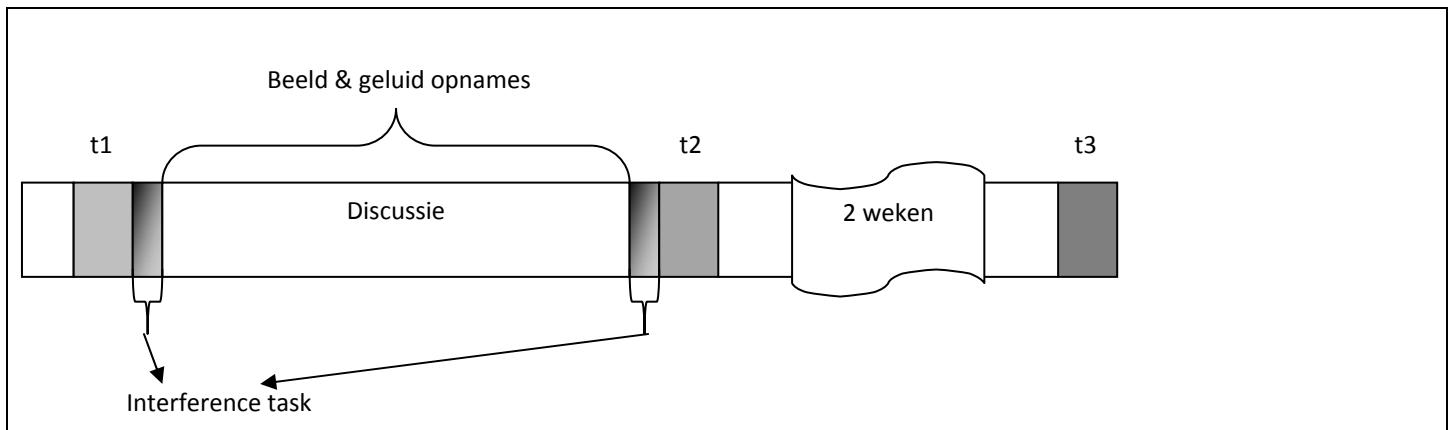
2.1 Methode empirisch onderzoek

2.1.1 Onderzoeksgroep

Aan het onderzoek deden 34 proefpersonen mee, voor het overgrote deel studenten aan de UvA tussen de 19 en 25 jaar. Deze proefpersonen zijn onderverdeeld in 7 groepen van 4 of 5 mensen. Elke groep deed zijn eigen sessie. Een sessie omhelsde een groepsdiscussie over ‘Lord of the Rings’.

2.1.2 Tokens

Bij elke discussie / sessie van het experiment zijn beeld en geluid opgenomen voor verdere analyse. Er is gebruikt gemaakt van een tweetal ‘tokens’ om de discussie te structureren en verdere analyse te vergemakkelijken. Tijdens de groepsdiscussie mochten de proefpersonen alleen praten als zij de ‘praat-token’ in bezit hadden. De proefpersonen kregen opdracht om te komen ‘met een lijst van de 5 meest karakteristieke kenmerken van Lord of The Rings’, als iemand van mening was dat er consensus was bereikt over het feit dat een woord op deze lijst geplaatst diende te worden dan kon deze hiervoor de ‘consensus-token’ pakken.



Figuur 1 | Verloop van onderzoek. Het verloop van het onderzoek onderscheidt zich in 3 tijdstipmomenten, gelabeld als t1, t2 en t3. Op deze momenten wordt de associatie-taak uitgevoerd, vlak na t1 en vlak voor t2 doen de proefpersonen een interferentie-taak. Tussen t2 en t3 zitten 2 weken. Van de discussie zijn beeld en geluidsopnames gemaakt.

2.1.3 Verloop van de sessie

De sessies die de groepen doorlopen hebben kunnen worden opgesplitst in zeven fases. We hebben deze fases uiteengezet in figuur 1.

Er is sprake van drie relevante tijdstipmomenten, namelijk vlak voor de discussie (t1), vlak na de discussie (t2) en 2 weken na de discussie (t3). In de eerste fase (t1) is aan alle proefpersonen onafhankelijk van elkaar gevraagd een 'associatie-taak'; te doen. De proefpersonen hebben een lijst ingevuld met concepten als antwoord op de vraag 'waar men aan moest denken bij *'Lord of the Rings'*', zij konden maximaal 50 concepten opschrijven (Ltotaal). Vervolgens moesten de proefpersonen met kruisjes op de eigen lijst aangeven welke 10 concepten van deze lijst zij het 'kenmerkendst' vonden (L10). Nadat zij dit gedaan hadden kregen alle proefpersonen voor 1 minuut een veelgebruikte *interference task* aangeboden³.

³ Ze moesten gezamenlijk een *willekeurige* serie van letters maken door in hoog tempo één voor één een letter te noemen. Het doel van deze *interference task* was de proefpersonen te laten vergeten wat ze precies op hun lijstjes opgeschreven hadden. Dit doet een *interference task* door het cognitieve systeem, met name het werkgeheugen, in grote mate te belasten. (Baddeley, 1998)

Vervolgens kreeg de groep de opdracht om met een *gezamenlijke lijst van de 5 meest kenmerkende begrippen voor 'Lord of the Rings'* te komen. Men kon pas een item op de *gezamenlijke lijst* zetten als alle deelnemers het hier mee eens waren. Kortom, er moest sprake van *consensus* zijn.

Na deze *gezamenlijke taak* werd nog een keer een *interference task* gedaan om herrinnering aan het verloop en de uitkomst van de discussie te onderdrukken. Vervolgens werd (op t2) de 1^e opdracht: de associatie-taak, en het daarna aangeven van de 10 belangrijkste begrippen, herhaald. Na 2 weken (t3) werd de proefpersonen per e-mail nogmaals gevraagd om de associatie taak te doen. Alle proefpersonen hebben dit gedaan.

2.1.4 Ruwe data en verwerking

Deze discussie leverde ruwe data op die zowel als input diende voor het computer model als voor de statistische analyse van de homogeniseringstheorie. Op alle tijdstipmomenten (t1, t2 en t3) hebben proefpersonen de associatietask gedaan. Er zijn in totaal per proefpersoon dus 6 lijsten beschikbaar (namelijk per tijdstipmoment de lijst met alle ingevulde begrippen en de lijst met de tien belangrijkste begrippen van deze lijst, Ltotaal en L10).

Tussen t1 en t2, oftewel: tijdens de groepsdiscussie zijn video-opnames gemaakt. Deze ruwe data is vervolgens verwerkt tot een formaat wat als input kon dienen voor statistische analyse en het computationele model.

De in het geheugenmodel en in de transcriptiematrices gehanteerde 'kernconcepten' zijn gebaseerd op de lijsten van 10 belangrijkste items op t1. In totaal besloeg het 'interpretatieve kader' van het computationele model hiermee 158 concepten. Om de invloed van eigen interpretatie te minimaliseren is er bewust voor gekozen om het interpretatieve kader van het model te baseren op de

associaties van proefpersonen in plaats van een voorgedefinieerde lijst van associaties. Binnen de content en map analyse staat deze keuze bekend als de 'interactive concept choice' (Carley, 1993; p 83).

2.1.5 Video-opnames

Met behulp van deze kernconcepten zijn de video-opnames omgezet in transcriptie-matrices. In deze matrices staat een transcriptie van de groepsdiscussie waarin de kernconcepten die proefpersonen hebben genoemd in een *turn* (de tijd dat ze onafgebroken de *praat-token* in hun handen hebben) geregistreerd staan. Een zeer klein deel van deze matrices (2 turns uit een discussie) is hieronder ter illustratie opgenomen.

Groep	1	1
Token	<i>Praat</i>	<i>Praat</i>
Wie	actor 8	actor 6
Turn	8	9
Concepten	Reisgenootschap	reisgenootschap
	drie delen	Karakteristiek
	Vriendschap	lord of the rings
	Bondgenootschap	Uiteenvallen
		Vriendschap
		Vertrouwen

Figuur 2 | Transcriptie-matrix van twee turns. Zoals te zien is voor iedere 'turn', aangegeven door het oppakken van het 'praat-token' bijgehouden welke kernconcepten zijn genoemd en door wie deze genoemd zijn. Dit betreft een zeer klein sample. In werkelijkheid bestonden iedere discussie uit ongeveer 35 turns.

2.1.6 Lijsten uit de associatietaak

Ieder proefpersoon heeft driemaal de associatietaak uitgevoerd (namelijk op t1, t2 en t3). Ook deze lijsten zijn omgezet in datamatrices om ze beschikbaar te maken voor statistische analyse en voor het computationele model. Hierbij waren de proefpersonen vrij om alles op te schrijven. Deze lijsten uit de associatietaak waren dus niet gelimiteerd tot de kernconcepten.

2.1.7 Samentrekken vergelijkbare begrippen

Om de statistische analyse van de homogeniseringsthese en de werking van het computationele model te

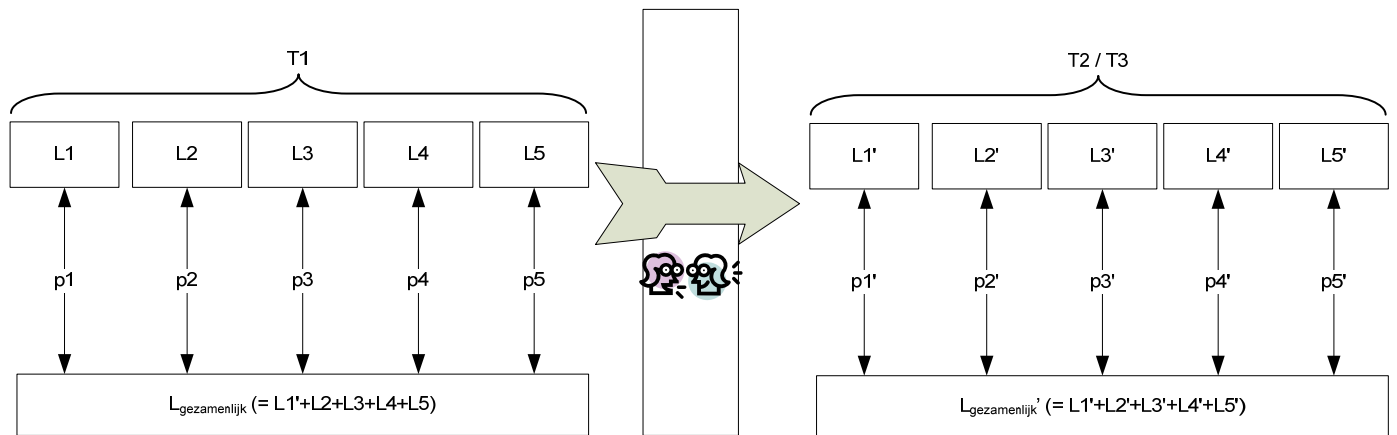
vergemakkelijken zijn begrippen die naar hetzelfde concept verwijzen 'samengetrokken'. Het computationele model gebruikt wat proefpersonen zeggen tijdens de groepsdiscussies (zie illustratie 2) . Kleine taalkundige variaties worden door de door ons gebruikte algoritmes echter automatisch herkend als een 'volledig ander concept' (de algoritmes kunnen bijvoorbeeld geen onderscheid maken tussen gollum of gollem). Daarom zijn zeer soortgelijke begrippen samengetrokken. Hier is een log van bijgehouden. De volgende regels zijn hierbij gehanteerd: verschillen in spelling zijn gelijkgetrokken (vb: *gollum* / *gollem* of *wizzard* / *wizard*), alle meervoudige woorden (vb: *paarden*) zijn omgezet naar enkelvoudige woorden (vb: *paard*) en alle Engelse woorden zijn omgezet naar Nederlandse (vb: *middle earth* = *midden aarde*).

In tegenstelling tot wat vaak in map- en content-analyse gebeurt (Carley, 1993; p 85) is de keuze gemaakt om de door proefpersonen opgeschreven concepten niet verder te 'generaliseren'. Op het moment dat de proefpersonen bijvoorbeeld 'the white wizzard' of 'die ouwe met baard' hebben opgeschreven dan hebben wij dit niet omgezet naar 'gandalf'. De reden hiervoor is dat we wilde voorkomen dat de optredende homogenisering veroorzaakt zou kunnen worden door onze eigen 'handmatige' homogenisering. Het 'meer in dezelfde woorden gaan denken' vormt immers een belangrijk deel van de homogeniseringsthese. Ook is het vaak maar de vraag of dit soort woorden aangeven dat 2 mensen echt dezelfde associatie hebben bij het kernbegrip.

Als gevolg van deze methodiek zitten er tussen de lijsten van proefpersonen waarschijnlijk nog 'verborgen overeenkomsten' die met de gehanteerde methodiek niet naar voren zijn gekomen. Er is echter geen reden om aan de nemen dat deze verborgen overeenkomst op t2 of t3 sterker is dan op t1, waarmee de eventuele toename in homogeniteit waarschijnlijk niet door deze keuze beïnvloed zal zijn. Het is echter wel waarschijnlijk dat de gevonden proportionele overeenkomst op alle tijdstipmomenten enigszins onderschat is.

2.1.8 Verwerking lijsten

Op de uitkomsten van associatietaak zijn (zelf ontwikkelde) algoritmes losgelaten om de overeenkomst tussen deze lijsten met concepten te tellen. De verschillende lijsten en hun bewerkingen en onderlinge vergelijkingen staan weergegeven in onderstaande figuur.



Figuur 3 | Schematisch overzicht van de gebruikte methode om mate van homogenisering te bepalen. Op t1 worden twee gezamenlijke lijsten opgesteld van alle individuele lijsten. Er wordt zowel van L10 en Ltotaal een gezamenlijk lijst gemaakt. Op t2 en t3 wordt dit wederom gedaan. Door de individuele lijsten en de gezamenlijke lijsten te vergelijken wordt per individu de mate van overeenkomst met de ‘groepsdefinitie’ gemeten (p). Door deze individuele parameters p_x en p_x' met elkaar te vergelijken worden kan de toename in homogeniteit bepaald worden.

De belangrijkste data die we uit deze lijsten wilde halen was de *mate van homogenisering van het waarheidsbeeld (geoperationaliseerd als toename of afname in overeenkomst tussen de uitkomst van de associatietoets)*. Om de ontwikkeling in homogenisering te bepalen is er, per meetmoment (t1, t2 en t3) per groep, een gezamenlijke lijst opgesteld van zowel de gezamenlijke lijsten ($L_{\text{gezamenlijk}}$) als de lijsten met 10 belangrijkste concepten. Op deze gezamenlijk lijst staan dus alle concepten die men in die groep heeft opgeschreven bij de associatietoets⁴. Vervolgens is per individu op alle drie de tijdstippen een parameter berekend die de percentuele overeenkomst tussen de eigen lijstjes en de gezamenlijke lijst aangeeft. Deze proportionele overeenkomst is met de volgende formule berekend:

$$\frac{H - A_{\text{individu}}}{A_{\text{groep}} - A_{\text{individu}}}$$

Formule 1 | H = getelde overeenkomst met gezamenlijk lijst.

A_{individu} : Aantal door het individu opgeschreven concepten.

A_{groep} : Aantal door de groep opgeschreven concepten.

De factor onder de streep kan omschreven worden als ‘aantal concepten die andere groepsleden bij elkaar hebben opgeschreven’, of te wel: het maximaal aantal concepten waarmee de door het individu opgegeven concepten overeen kunnen komen. De factor boven de streep is het aantal getelde overeenkomsten op de eigenlijst en de groepslijst, omdat de eigen concepten ook op de groepslijst voorkomen

moet hiervoor gecompenseerd worden. Ter verheldering: als iedereen precies dezelfde lijst zou hebben opgegeven dan is $p : 1$, als er geen enkele overeenkomst is is $p : 0$. De proportie ‘p’ varieert lineair tussen deze twee uitersten. Deze proportie is een eigenschap per proefpersoon en varieert per lijst (L10 Ltotaal, en t1,t2,t3, dus in totaal wederom zes stuks)

2.1.8 Statistische toetsing

Het verschil in percentuele individuele overeenkomsten met de gezamenlijke groepslijsten hebben we met een Repeated Measures ANOVA getoetst. De within-subjects-factor was de percentuele individuele overeenkomsten, bestaande uit 3 levels (3 meetmomenten). Tevens zijn er verschillende betrouwbaarheidsparameters berekend. De uitwerking van de statistische toetsing volgt in het volgende hoofdstuk.

⁴ Als bepaalde concepten door meerderen personen zijn opgeschreven dan komen ze dus meerderen keren op deze lijst voor.

3.Resultaten empirisch onderzoek

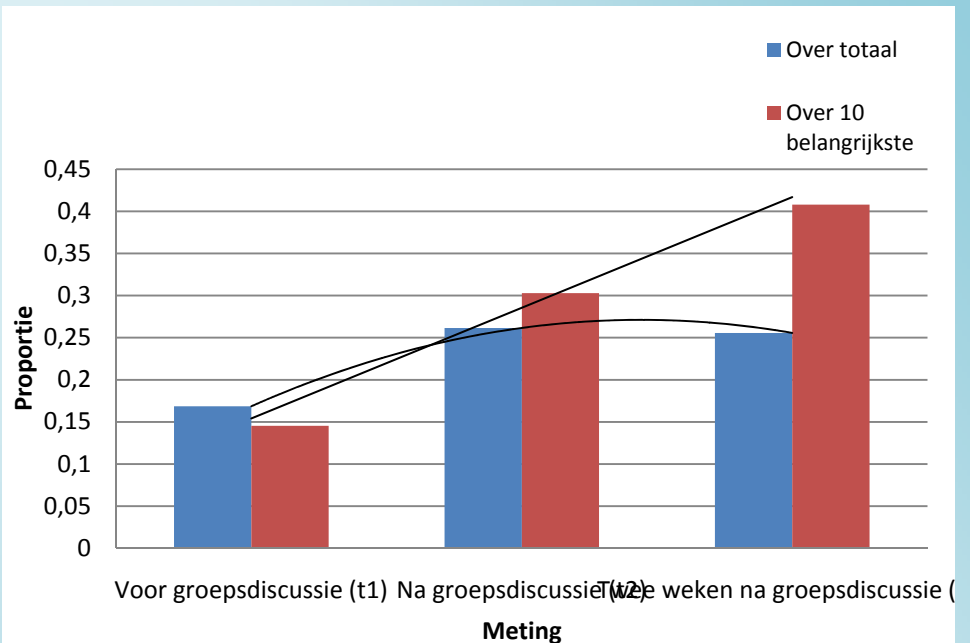
Met betrekking tot de empirische houdbaarheid van de homogeniseringstheze zijn de belangrijkste resultaten hiernaast samengevat in grafiek 1 en tabel 1.

De homogeniteit van de lijsten van '10 belangrijkste kenmerken' neemt toe. De totaallijst stijgt eerst mee, maar daalt vervolgens licht t.o.v. de toename gelijk na de groepsdiscussie. Op alle verschillen is getoetst met een repeated measures ANOVA. Met alle algoritmes van SPSS zijn beide effecten significant (Over totale lijsten: $F = 19,244$ (0,61) en $p < 0,001$, over lijsten van 10 belangrijkste kenmerken: $F(0,43) = 54,801$, $p < 0,001$).

Over de resultaten zijn contrasten gepast om te kijken met wat voor effect we te maken hebben. Voor de totale lijsten kan dit zowel een linear ($F(0,080) = 17,537$, $p < 0,001$) als een (nog wat sterker) kwadratisch contrast ($F(0,042) = 23,691$, $p < 0,001$) zijn.

Over de lijsten van 10 past alleen een (erg sterk) linear contrast ($F(0,824) = 74,474$, $p < 0,001$).

De stijging in homogeniteit tussen $t=2$ en $t=3$ bij de lijsten van 10 belangrijkste kenmerken is significant ($F(0,083) = 10,38$, $p = 0,04$). Ook de stijging in homogeniteit op $t=1$ en $t=3$ bij de totale lijsten is significant ($F(0,080) = 17,537$, $p < 0,001$). De daling in homogeniteit tussen $t=2$ en $t=3$ bij de totale lijsten is dat niet ($F(0,001) = 0,469$, $p > 0,05$). Men moet er dus van uit gaan dat de mate van homogeniteit van de totale lijsten na $t=2$ niet meer verandert.



Grafiek 1) Proportionele overeenkomsten tussen lijsten. Uit de resultaten blijkt dat de homogeniteit binnen groepen zowel voor de totaal als voor de 10 belangrijkste kenmerken na de groepsdiscussie hoger is. Voor de 10 belangrijkste kenmerken zet deze toename na twee weken verder door. Hierover kan een linear contrast gepast worden. De mate van homogeniteit tussen de totale lijsten blijft gelijk. Men kan over deze trend een kwadratisch contrast passen. Zie ook tabel 1.

	Voor discussie (t1)	Na discussie (t2)	Twee weken na discussie (t3)	Verskil t2- t1	Verskil t3 - t1
Over totale lijsten	0,179	0,307	0,272	0,128	0,093
Over 10 belangrijkste kenmerken	0,152	0,304	0,425	0,152	0,273

Tabel 1) Proportionele overeenkomsten tussen lijsten naar meting en soort lijst. Zie voor toelichting grafiek 1.

Controle op de consistentie van deze verandering, twee weken later, laat zien dat er sprake is van een duurzame toename in homogeniteit voor de totaallijsen (de totale homogeniteitstoename betreft dan 52%, de verandering t.o.v. t2 is niet significant). Opvallend is dat de homogeniteit 2 weken na de discussie op de 10 belangrijkste item nog eens 40% is toegenomen t.o.v. t2 terwijl er geen interactie tussen proefpersonen is geweest. Er is sprake van een ‘na-ijl effect’. De totale homogeniteitstoename voor de lijsten van 10 twee weken later komt hiermee op 180%.

	t1	t2	t3
Cronbach's alfa	0,647	0,823	0,715

Tabel 3 | Cronbach's Alfa over L10 en Ltotaal voor alle t's

Om te kijken in hoeverre de totale lijsten en de lijsten van 10 ‘hetzelfde meten’ hebben we voor alle 3 tijdstippen de Cronbach's alfa berekend. De resultaten hiervan staan in tabel 3. In de sociale wetenschappen wordt een Cronbach's Alfa van boven de 0,7 veelal gezien als een indicatie dat twee operationalisaties hetzelfde meten.

De gevonden resultaten ondersteunen de homogeniseringsthese. Na afloop van de discussie is er sprake van een meer homogene definitie van het onderwerp van de dan voorafgaande aan de discussie. De toename in homogeniteit door de discussie tussen proefpersonen voor L10 is gemiddeld 100% tussen t2 en t1 en 180% tussen t3 en t1. Voor Ltotaal is dit 72% tussen t2 en t1 en 52% tussen t3 en t1.

4. Methode en opzet computationeel model

De input van het computationeel model betreffen de individuele lijstjes van de proefpersonen op t1 (de associatielijsten) en de transcriptiematrices van de discussies per groep. Het model heeft dus expliciet geen kennis van de

individuele lijstjes na de discussie. Dit is immers wat het model tracht te voorspellen. Als in dit paper wordt gesproken over de proefpersonen in de empirie zal er worden gesproken over ‘actors’, als het gaat om de corresponderende individuen in het computationele model worden deze ‘agents’ genoemd.

Het ontwikkelde computationele model is in hoge mate gebaseerd op een model van Rumelhart et al. (1986, chapter 14) in het boek ‘Parallel Distributed Processing’. Rumelhart et al. hebben proefpersonen specifieke soorten kamers (bijvoorbeeld ‘badkamer’ en ‘keuken’) laten beschrijven. Proefpersonen moesten op lijsten van 40 *room descriptors* (bijvoorbeeld ‘stoel’ en ‘wastafel’) aangeven of deze wel of niet aanwezig waren in een specifiek soort kamer. Met deze informatie hebben zij een neurale netwerk, meer specifiek: een *attractor network* ⁵ gebouwd. De *nodes* in dit netwerk zijn de 40 room descriptors. De mate van connectie tussen twee verschillende descriptors houdt in dit netwerk verband met de frequentie dat 2 descriptors wel of juist niet samen op een ingevulde lijstjes voorkwamen. Het blijkt dat dit netwerk zich maar in een beperkt aantal evenwichten in kan stellen die overeen komen met schema's van bepaalde types kamer. Men zou kunnen zeggen dat als een netwerk in evenwicht is de mate van dissonantie in het netwerk een minimum heeft bereikt. Ons computationeel model van kennis, het neurale netwerkmodel, is grotendeels analoog aan dat van Rumelhart et al. opgezet. De descriptors zijn in dit geval ‘kernconcepten van Lord of The Rings’ (Bijvoorbeeld ‘ring’ of ‘Frodo’). Aan de hand van de door proefpersonen ingevulde lijsten is er een neurale netwerk opgezet dat het algemene ‘begrip’ van ‘Lord of the Rings’ van alle proefpersonen beschrijft. Ten opzichte van Rumelhart et al. is het aantal *descriptors* in ons onderzoeken stuk hoger. Rumelhart et al. had 40 descriptors, omdat proefpersonen bij ons vrij waren om eigen associaties op te schrijven (zie paragraaf 2.1.7) waren het er wij ons 158.

De eerste stap van het opzetten van ons computationele model was het vormen van een representatie van het ‘meta-begrip’ van *Lord of the Rings*. Hiervoor zijn dezelfde

	ring	frodo	hobbit	fantasie	tovenaar	orks	reis	magie	bergen	sprookjes
ring	0	0,201	0	3,321	0,143	5,198	3,321	1,099	3,321	-0,125
frodo	0,201	0	0,488	3,076	-0,799	1,179	3,076	4,634	3,076	-0,379
hobbit	0	0,488	0	3,442	0,934	0,474	3,442	1,232	3,442	4,2
fantasie	3,321	3,076	3,442	0	4,521	5,362	10,259	5,684	10,259	6,847
tovenaar	0,143	-0,799	0,934	4,521	0	-0,421	4,521	-0,087	4,521	1,099
orks	5,198	1,179	0,474	5,362	-0,421	0	5,362	2,197	5,362	1,946
reis	3,321	3,076	3,442	10,259	4,521	5,362	0	5,684	10,259	6,847
magie	1,099	4,634	1,232	5,684	-0,087	2,197	5,684	0	5,684	2,269
bergen	3,321	3,076	3,442	10,259	4,521	5,362	10,259	5,684	0	6,847
sprookjes	-0,125	-0,379	4,2	6,847	1,099	1,946	6,847	2,269	6,847	0

Tabel 2 Deel van het 'meta-netwerk' dat de lijstjes van alle proefpersonen representeert. Te zien valt dat bepaalde concepten een grote voorspellende waarde hebben voor een ander concept (bijvoorbeeld 'reis' en 'bergen'). De individuele netwerken 'zien er hetzelfde uit' als dit netwerk, maar bevatten individuele associaties.

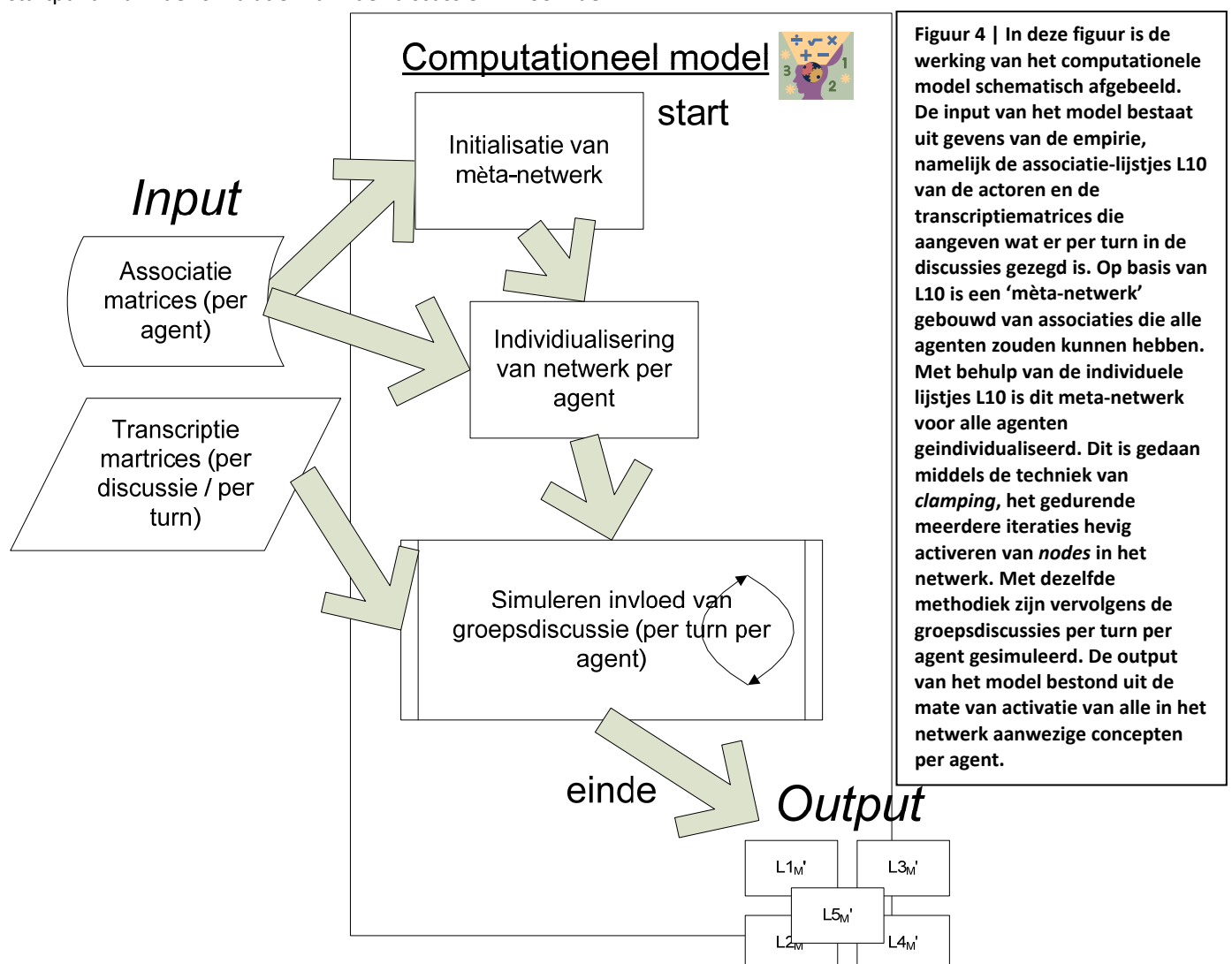
algoritmes gebruikt als door Rumelhart et al. In het hierboven behandelde onderzoek.

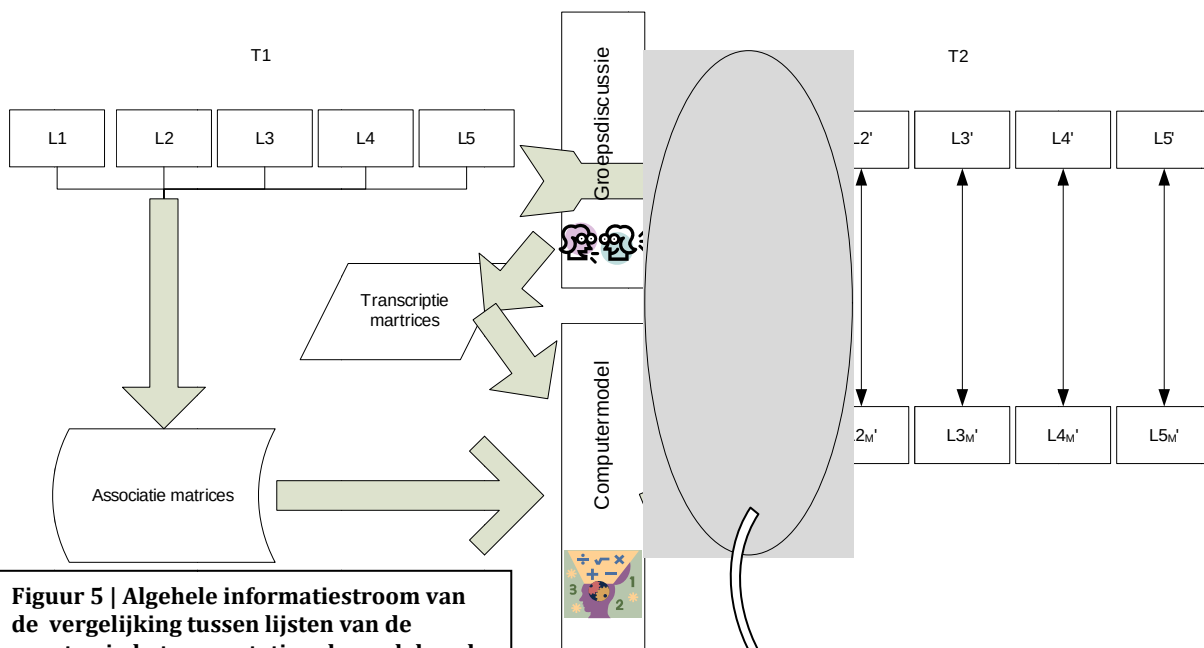
De tweede stap in de simulatie was om voor iedere actor een corresponderende agent met een eigen 'geindividualiseerd' neurale netwerk maken. Het meta-netwerk uit de vorige alinea vormt hierbij de basis voor het individuele netwerk. De individuele netwerken zijn uit het meta-netwerk geconstrueerd door de concepten op de L10 van de actor in het netwerk van de corresponderende agent te *clampen*. Als een begrip wordt *geclampt* dan wordt het gedurende meerdere iteraties maximaal geactiveerd. Omdat er een *learning rule* actief is wordt het netwerk hierdoor structureel veranderd. De concepten die door de actor op L10 zijn gekenmerkt als meest belangrijke concepten worden door deze directe ingreep nu sterker gerepresenteerd. Indirect worden ook de aan deze concepten geassocieerde concepten sterker gerepresenteerd.

Het geindividualiseerde netwerk uit stap 2 vormde het startpunt van de simulatie van de discussie. Voor de

simulatie van de discussie wordt opnieuw informatie uit de empirie gebruikt. Namelijk, de transcriptie-matrices waarin per 'turn' van de discussie aangegeven staat welke concepten genoemd zijn. In deze stap van de simulatie in het computationeel model wordt de discussie per 'turn' gesimuleerd. Analoog met de 'individualiseringsstap' zijn per turn steeds de genoemde concepten aan de netwerken 'gevoerd'. Per turn genoemde concepten worden tegelijkertijd *geclampt*. Dit gebeurde een paar iteraties, waarna weer de woorden uit de volgende *turn* geclampt worden. Iedere agent doorloopt de discussie dus in dezelfde stappen als de corresponderende actor.

Nadat de agenten in ons model hebben deelgenomen aan deze 'virtuele discussie' produceren zij, net als de actoren (in de empirie) een 'associatietaak', het model geeft als output waar de agenten 'aan moeten denken' bij Lord of the Rings na afloop van de discussie. De output van het model is de mate van activatie van alle concepten binnen het netwerk van de individuele agenten.





Figuur 5 | Algehele informatiestroom van de vergelijking tussen lijsten van de agents in het computationele model en de lijsten van de actors in de empirie. Het uiteindelijk doel van deze stroom is het bepalen van de voorspellende waarde van het model.

- Het bovenste deel van de figuur is identiek aan figuur 3. Een specifiekere toelichting van het computationele model staat uitgewerkt in figuur 4.
- Er is te zien dat de uitkomsten van het model en van proefpersonen op een viertal parameters zijn vergeleken. Namelijk de algehele lijsten, de gebleven woorden, de verdwenen woorden en de verschenen woorden.

Van deze output kunnen ‘associatie-lijstjes’ gemaakt worden door de n (bijvoorbeeld 10) nodes te selecteren met de hoogste activiteiten.

4.1 Vergelijking model empirie

Om de kwaliteit van het model te testen zijn de door de proefpersonen ingevulde associaties en de voorspelling van deze associaties door het computationele model per individu (agent en actor) op een viertal parameters met elkaar vergeleken. Namelijk: op de gehele lijst, de ‘gebleven woorden’, de ‘verdwenen woorden’ en de ‘verschenen woorden’⁶. Al deze parameters betreffen wederom proportionele overeenkomsten (op dezelfde manier vastgesteld als hierboven behandeld in paragraaf 2.1).

⁶ Respectievelijk: de woorden die zowel op t1 als t2 bij de associatietask naar voren kwamen, de woorden die wel op t1 voorkwamen maar niet op t2 en de woorden die wel op t2 voorkwamen maar niet bij t1.

Deze opsplitsing zorgt voor een beter beeld van voorspellende waarde van het computationele model (als bijvoorbeeld alleen de totaallijsten vergeleken zouden worden dan zou de volledige ‘voorspellende waarde’ veroorzaakt kunnen worden door alle gebleven concepten, een ‘computationeel model’ wat simpelweg alle concepten van de eerste lijst als voorspelling voor de tweede lijst geeft, zou dan evenveel voorspellende waarde hebben als de nu gebruikte connectionistische algoritmes). Tevens heeft deze opsplitsing de optimalisatie van het computationele model gemakkelijker gemaakt. De toetsing van de voorspellende waarde van computationele model opgenomen is het volgende hoofdstuk.

4.1.1 Dynamische drempelmethode.

De resultaten van het computationele model worden tevens op significantie getoetst. De hierbij gehanteerde methode is als volgt: het gehele ‘repertoire’ van het computationele model

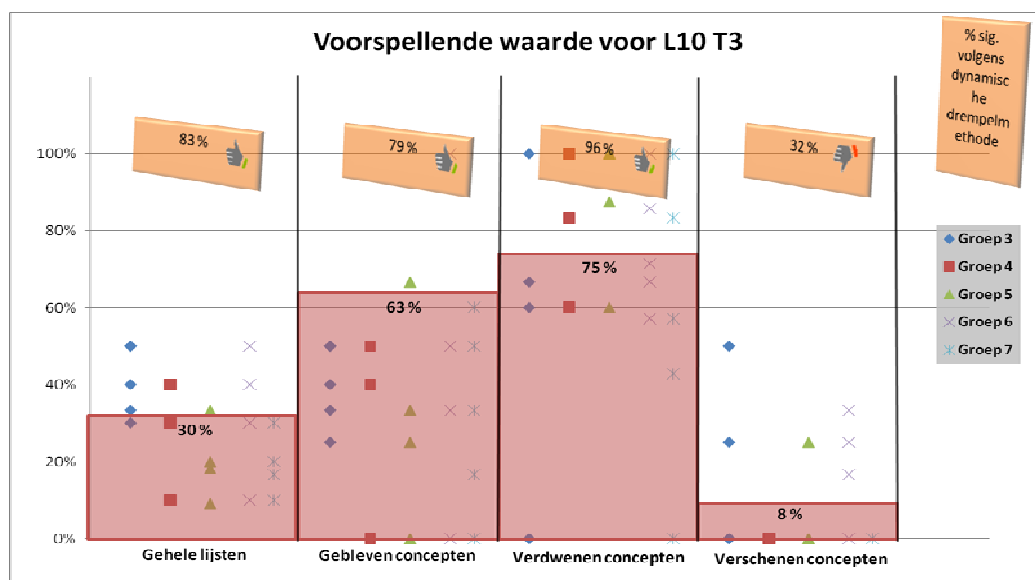
bestond uit 158 concepten, op het moment dat het computationele model deze items random aan proefpersonen zou toedelen dan maakt ieder concept even veel kans om op een lijst terecht te komen. Als het model dus random 10 items opgeeft (van de 158) dan is te verwachten dat +/- 6% van de lijst goed voorspeld wordt (=10/158). Omdat het aantal items op de verschillende lijsten (totaal, gebleven, verdwenen, verschenen) per individu varieert kan per individu een onderste drempelwaarde bepaald worden. Als de gevonden waarde niet boven deze onderste drempelwaarde uitkomt dan is het gevonden resultaat niet significant. Boven op deze berekende drempelwaarde heb wij nog een extra 5% toegevoegd alvorens de output het label 'significante voorspelling' mee te geven. Of een voorspelling 'significant' is, is dus een binaire waarde (ja/nee) per agent / per lijst. Wij zullen deze toets op significantie in dit verslag verder de 'dynamische drempelmethode' noemen.

5. Resultaten computationele model

In dit hoofdstuk worden de resultaten van het computationele model vergeleken met de resultaten uit de empirie. De uitkomsten van het computationele model zijn zowel vergeleken met de empirische uitkomsten op t2 als met de empirische uitkomsten op t3. Wederom is ook hier een splitsing gemaakt tussen L10 en Ltotaal (de lijst met 10 belangrijkste associaties en de gehele lijsten). Voor L10 in het model zijn simpelweg de 10 concepten met de hoogste activaties van een actor op zijn voorspellingslijst gezet. De gevonden voorspellende waarde voor t2 en t3 lopen niet zeer sterk uiteen, net als dat de voorspellende waarde voor L10 en Ltotaal niet erg sterk uiteenlopen. De kwaliteit van de voorspellingen voor L10 op t3 is weergegeven in grafiek 2. De (vergelijkbare) grafieken met de voorspellende waarden voor L10 op t3 en L10 en Ltotaal op t2 zijn opgenomen in bijlage 1.

Op de X is as staan de verschillende parameters waarop de lijsten zijn vergeleken. Op de Y as valt de voorspellende waarde van het model op deze vier parameters af te lezen. De verschillende punten geven de voorspellende waarde binnen valt voor de parameters in die kolom af te lezen hoeveel van de lijsten een significant voorspellende waarde hebben volgens de, in paragraaf 2.2.2 uiteengezette, dynamische drempelwaarde methode. deze parameters voor specifieke proefpersonen aan. Deze zijn, met een teken, ingedeeld naar groep. Bovenin de grafiek

Er is te zien dat het computationele model gemiddeld ruim 30% van de gehele lijsten weet te voorspellen. Dit betekent dat drie op de tien woorden die een proefpersoon op t3 op zijn lijst opschrijft ook voorkomen in de voorspelling door het computationele model hiervan. Ruim 83% van de door het computationele model voorspelde lijsten hebben een significant voorspellende waarde. Als we zowel voor de uitkomsten van het computationele model als voor de uitkomsten empirie bepalen welke concepten op de totaalijsten zijn 'gebleven' dan is het computationele model in staat om ruim 60% van deze gebleven concepten te voorspellen. Hierbij heeft 79% van de lijsten een significante voorspellende waarde. Het model blijkt het best te kunnen voorspellen welke concepten van de totaalijsten zullen verdwijnen. 75% van de woorden die van de totaalijsten verdwenen zijn wist het computationele model te voorspellen. Enkele van deze lijsten werden zelf volledig goed voorspeld. (alle woorden die in de empirie verdwenen verdwenen ook in het computationele model). 96% van de voorspellingen voor deze verdwenen concepten waren significant. Het model bleek het slechtste in staat om te voorspellen welke woorden als gevolg van de groepsdiscussie zouden zouden verschijnen. Slechts 8% van de verschenen woorden werd goed voorspeld. Tevens was slechts 32% van die gedane voorspellingen significant.



Grafiek 2 | In deze grafiek is de voorspellende waarde van het computationele model naar een aantal parameters uiteengezet. Op de x-as staan de verschillende parameters, op de y-as staat de procentuele overeenkomst tussen de lijsten van de acteurs in de empirie en agenten in het model op de betreffende parameters. Bovenaan staat het aantal significant voorspelde empirische lijstjes volgens de 'dynamische drempel'-methode.

5.1 Homogenisering in het computationele model

Ook in het computationele model doet zich homogenisering voor. Analoog aan de methode uiteengezet in paragraaf 2.1 hebben we ook voor de voorspellingen van het computationele model een homogeniseringsfactor berekend. Deze homogenisering is erg sterk (de overeenkomst gaat naar ruim 82%) dit is te verklaren uit het feit dat we voor alle actoren het op de 158 kernbegrippen gebaseerde 'meta-begrip' gebruik hebben voor hun individuele representatie. Echter, het startpunt voor de agenten was een volkomen identieke 'definitie', dit meta-begrip is vervolgens welliswaar geïndividualiseerd, desalniettemin is het, met gebruik van deze methode, onvermijdelijk dat de homogenisering in het model 'overschat wordt'.

Ook de voorspellende waarde die het model heeft voor de homogeniteitstoename in de empirie is onderzocht. Als de homogeniteitstoename in het model wordt gecorelleerd met de homogeniteitstoename in de empirie dan komt hier een zeer sterk positief verband uit. De Pearson correlatie is: 0,754 (sig. 0,001), wat neer komt op een verklaarde variatie van 57%. Oftewel, de voorspellende waarde van het model voor de homogeniteitstoename in de empirie is 57%.

6. Discussie

Het uitgevoerde onderzoek betreft de eerste voorzichtige stapjes in een potentieel zeer veelbelovend onderzoeksgebied naar de invloed van groepsdiscussies op cognitieve representaties met behulp van computationele modellen. De empirische toetsing van de homogeniseringstheorie kan een nieuw licht werpen op

kennissuitwisseling tijdens menselijke interactie. De homogeniseringstheorie stelt dat door interactie de 'waarheidsbeelden' van een groep mensen homogener worden. Uit de hier uitgevoerde statistische toetsing komt duidelijk naar voren dat dit voorspelde effect inderdaad plaatsvindt. De associaties met betrekking tot het gespreksonderwerp van de deelnemers aan een groepsdiscussie worden inderdaad homogener. Er is, onder invloed van de discussie, een verdubbeling van de overeenkomsten tussen associaties waar te nemen (van 15% naar 30%). Niet alleen ondersteunt het onderzoek de homogeniseringstheorie (en geeft zij hierdoor een steuntje in de rug aan de besproken sociaal interactionistische

sociologen), zij heeft tevens interessante nieuwe informatie opgeleverd. Er blijkt namelijk dat er niet alleen tijdens de feitelijke interactie zelf homogenisering optreedt. Er is sprake van een soort van 'na-ijl'-effect. Twee weken na de groepsdiscussie is de gemiddelde overeenkomst tussen associaties van leden van dezelfde discussiegroep nog eens een extra 12,5% gestegen (naar een homogeniteit van 42,5 %).

Als de homogeniseringstheorie inderdaad waar blijkt te zijn is een interessante vraag hoe deze homogenisering precies tot stand komt. Als mensen steeds meer hetzelfde over een onderwerp gaan nadenken, naar wiens 'waarheidsbeeld' beweegt deze homogenisering zich dan toe en welke processen liggen hier precies aan ten grondslag? Ongetwijfeld spelen factoren als sociale status, persoonlijke relaties en expertise van de deelnemers aan een interactie hierbij een belangrijke rol. In de antropologie is hier enig onderzoek naar gedaan. Zo hebben Kameda et al. (1997) in een onderzoek naar groepsdiscussies over het onderwerp 'doodstraf' uitgevonden dat personen die veel kennis met andere leden van de groep delen meer invloed hebben op de uitkomst van de discussie. Zij maken hierbij gebruik van het begrip 'sociocognitief netwerk': een model dat aangeeft hoeveel informatie specifieke leden van een groep met elkaar delen.

In het begin van dit paper vroegen wij ons af in hoeverre de invloed van groepsdiscussies op de deelnemers aan deze discussie te modelleren is met op een connectionistische paradigma gebaseerd computationele model. Op basis van bovenstaande resultaten valt te stellen dat de eerste resultaten van deze aanpak veelbelovend zijn. Vanwege het exploratieve karakter van het onderzoek heeft het gebruikte model momenteel een nogal hoog 'ducttape'-gehalte. Desalniettemin is het model, in ieder geval op een tweetal parameters (de homogeniteitstoename en de 'verdwenen woorden') opvallend goed in staat om voorspellende waarde te genereren. Ook voor het voorspellen van de 'gehele lijsten' en de 'gebleven concepten' verloopt opvallend aardig. Een belangrijk verbeterpunt vormt het voorspellen van de 'verschenen concepten'. De voorspellende waarde van het model op deze parameter blijkt nu uiterst minimaal.

6.1 De mens als associeermachine?

Een belangrijke aanname van het hier uiteengezette computationele model was het 'associatieve karakter' van de menselijke geest. Echter, in het ontwerp van het model speelden vooral praktische overwegingen een rol. De beschikbare kennis en algoritmes in de kunstmatige intelligentie vormen veelal een limiet op de mate waarin een

dergelijk model recht kan doen aan gedetailleerdere kennis omtrent de menselijke cognitie uit de psychologie.

Uiteraard werkt het menselijk brein ook niet slechts als een simpele 'associeermachine'. Bij menselijke cognitie komen een heleboel 'hogere-orde' functies kijken die niet zo makkelijk te modelleren zijn. Deze hogere-orde processen worden door cognitieve neurowetenschappers meestal voor een belangrijk deel aan de prefrontale cortex, het hersengebied vlak boven de ogen, toegeschreven. Deze hogere-orde processen zijn moeilijk te modelleren, maar spelen een belangrijke rol bij conversaties en hun invloed op cognitieve representaties. Zo laten studies naar laesies aan de prefrontale cortex zien dat mensen met een laesie moeite hebben met converseren. Deze mensen voeren bepaalde 'executive functions' (planning van acties, inhiberen van gedrag, aansturen van andere cognitieve processen) minder goed uit en hebben moeite van onderwerp te wisselen en nieuwe concepten te formeren. (Frankel et al, 2007). Het lijkt dus aannemelijk dat de executieve functies van de prefrontale cortex bij conversatie en het aanpassen van de cognitieve representaties (de zogenaamde conceptformatie) een belangrijke rol speelt. Bij schade aan de prefrontale cortex blijft men vaak steken in hetzelfde 'veld' van associaties en kan men moeilijker uiteenlopende informatie integreren.

6.2 Na-ijl effect

Opvallend is dat het effect voor de lijsten van tien 'meest kenmerkende kenmerken' twee weken na de groepsdiscussie nog verder toegenomen lijkt te zijn. De mate van homogeniteit over deze lijsten neemt in deze periode relatief nog eens 53% toe. Bij de totale lijsten vindt dit effect niet plaats. Deze lijsten blijven dezelfde mate van homogeniteit behouden als vlak na de groepsdiscussie. Hoe is dit effect te verklaren? Mogelijk vindt vooral op de waardering van bepaalde associaties een na-ijl effect plaats. Cronbach's Alfa over de homogeniteitsindicatoren van beide lijsten zijn voor (zowel $t=2$ als $t=3$) behoorlijk hoog. Dit geeft aan dat ook op $t3$ L10 en Ltotaal hetzelfde lijken te meten.

6.3 Vervolgonderzoek

De gevonden resultaten geven interessante informatie over de validiteit van de homogeniseringsthese. Om echter echt uitspraken te kunnen doen over de validiteit van deze these (en het gevonden 'na-ijl' -effect) is vervolgonderzoek nodig. Zo is het nu onbekend of de door ons gevonden homogeniseringseffecten als gevolg van interactie ook

optreden bij een andere gespreksonderwerp dan Lord of The Rings. Tevens is de invloed van 'groepstaak' tijdens deze discussie waarschijnlijk van belang.

Een belangrijke vervolgstap zou kunnen zijn om te kijken naar de invloed van herhaalde interactie. Neemt bij een tweede en derde samenkomst van de groep de homogeniteit nog steeds toe? En zo ja, is deze dan even sterk, zwakker of juist sterker? Zijn er bepaalde groepen waarbij dit wel het geval is en andere waarbij dit niet het geval is? En waar komt dit dan door? En het na-ijl effect, is daar ook bij herhaalde interactie sprake van? Voor beantwoording van deze interessante vragen is vervolgonderzoek nodig.

Deze behoefte aan vervolgonderzoek geldt ook in hoge mate voor het computationele model. Er zijn nog verschillende kleine aanpassingen en grote veranderingen aan het gehanteerde model mogelijk waarvan wij hoge verwachtingen hebben. Allereerst zal de kwaliteit van het model waarschijnlijk sterk omhoog gaan als data aan het model 'gevoerd' wordt vergaander wordt 'voorbereid'. Een weloverwogen samentrekken (generalisering) van de gebruikte concepten (zoals uiteengezet in paragraaf 2.1.7) zal de voorspellende waarde van het model waarschijnlijk sterk ten goede komen (hierbij moet uiteraard wel het subjectieve karakter en de validiteit van deze generalisering in het oog gehouden worden).

In bredere zin zijn er tijdens het programmeren van het model enkele keuzes gemaakt die later de optimalisering van het model in de weg bleken te staan. Nu deze kennis beschikbaar is zal een nieuwe versie van het model, met grotendeels dezelfde algoritmes, waarschijnlijk veel betere resultaten opleveren. Bovendien kan er 'gespeeld' worden met de hoeveelheid data die het model uit de empirie nodig heeft om voorspellende waarde genereren. Wie weet is er een model mogelijk waarin de specifiek tijdens de groepsdiscussie genoemde concepten niet meer uit de empirie gehaald hoeven te worden maar waarin het afdoende is om aan te geven 'wie er aan het woord is'.

Ondanks het basale karakter van het gebruikte model, waarin bovenstaande optimalisaties ontbraken, bleek het model op een veelbelovende manier in staat om voorspellingen te doen over de invloed van groepsdiscussies op individuele associaties.

7. Literatuurlijst

Baddeley, Alan (1998) 'Random Generation and the Executive Control of Working Memory', *The Quarterly Journal of Experimental Psychology Section A*, 51:4, 819 - 852

Beck, Ulrich & Giddens, Anthony & Lash Scott (1994) *Reflexive Modernization. Politics, Tradition and Aesthetics in the Modern Social Order*. Cambridge: Polity Press.

Carley, Kathleen (1993), Coding choices for textual analysis: A Comparison of Content Analysis and Map Analysis. *Sociological Methodology*, 23: 75-126

Dewulf, Art Craps, Marc & Dercon, Gerd. (2003). How Issues Get Framed and Reframed When Different Communities Meet: A Multi-level Analysis of a Collaborative Soil Conservation Initiative in the Ecuadorian Andes. *Journal of Community & Applied Social Psychology*. 14: 177–192

Frankel, Tali, Penn, Claire and Digby Ormond-Brown. (2007). Executive dysfunction as an explanatory basis for conversation symptoms of aphasia: A pilot study. *APHASIOLOGY*, vol 21 (6/7/8). p 814–828.

Giddens, Anthony (1984). *The Constitution of Society: Outline of the Theory of Structuration*. Polity Press: Cambridge. p 1-25.

Grannovetter, M. (1974) *Getting a job: A study of contacts and careers*. Cambridge: Harvard University Press.

Gray. (2002) *Psychology: Fourth edition*. Worth Publishers: New York.

Kameda, Tatsuya, Ohtsubo, Yohsuke & Takezawa, Masanor (1997) Centrality in Sociocognitive Networks and Social Influence: An Illustration in a Group Decision-Making Context *Journal of Personality and Social Psychology* Vol. 73, No. 2, 296-309

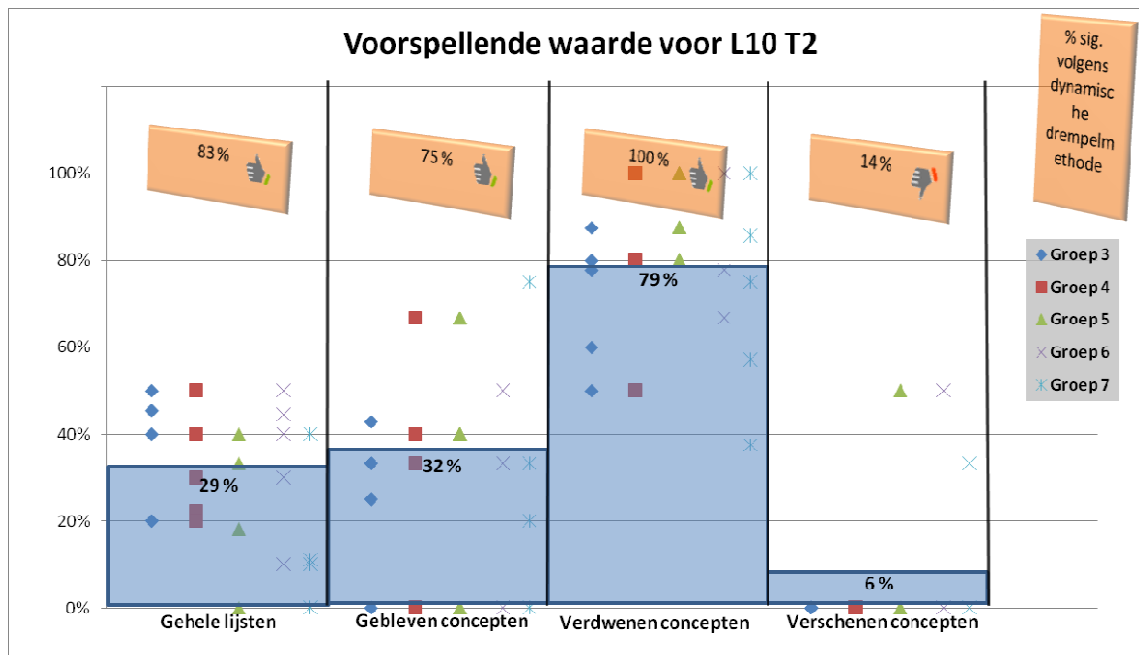
Rumelhart, D. E., Smolensky, P., McClelland, J. L., & Hinton, G. E. (1986). Schemata and Sequential Thought Processes in PDP Models. In J. L. McClelland, & D. E. Rumelhart, *Parallel Distributed Processing: explorations in the Microstructure of Cognition: Volume 2: Psychological and Biological Models* (pp. 7-57). MIT Press: Cambridge, Massachutes.

Wallace, R. Wolf, A. (2006) *Contemporary social theory: expanding the classical tradition*. New Jersey: Pearson Education.

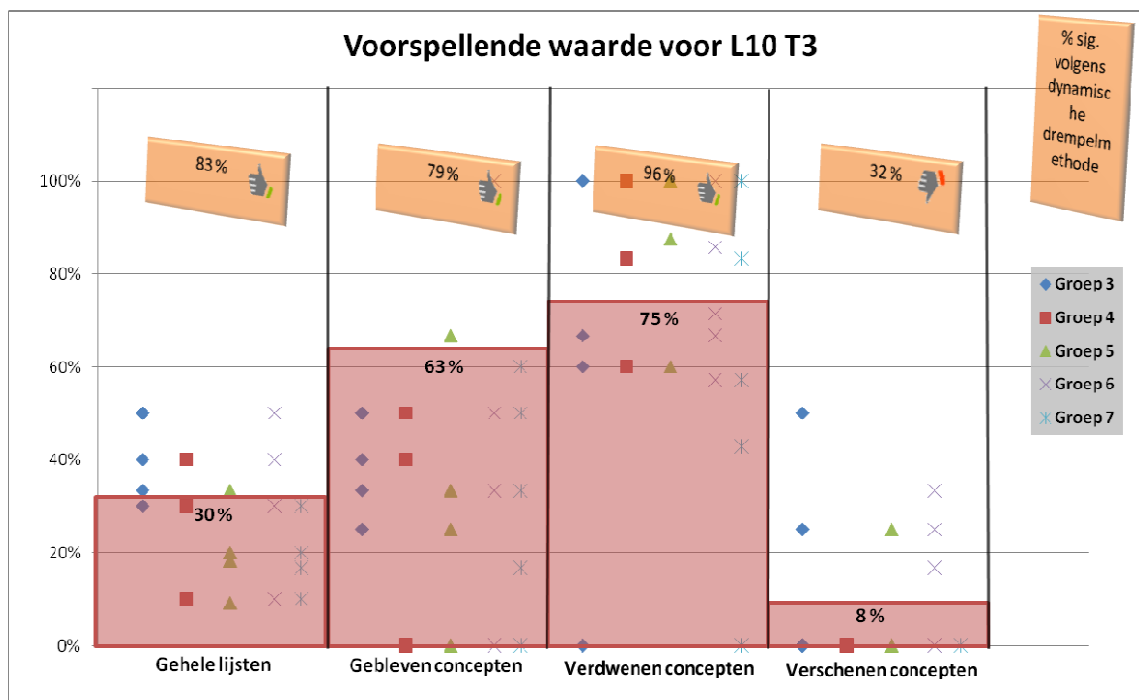
Willems, Dick (1989) "Foucault over wetenschap en Macht" in Boon Louis en Gerard de vries. *Wetenschapstheorie: de empirische wending*. Wolters Noordhof: Groningen.

Bijlage 1, voorspellende waarde computationele model voor andere lijsten

Voor lijst van 10 belangrijkste items op t2



Voor lijst van 10 belangrijkste items op t3



Voor Totaallijst op t3

Voorspellende waarde voor L Totaal T2

