

KENMERKECONOMIE IN DE GTRP-DATABASE

1. De FAND en de moderne fonologie

Uit Goeman & Taeldeman (1996),⁽¹⁾ waarin de voorgeschiedenis van de *Fonologische Atlas van de Nederlandse Dialecten* (Goossens et al., 1998, 2000; De Wulf et al., 2005) uiteen wordt gezet, blijkt dat het werk aan het Goeman-Taeldeman-Van Reenen Project (GTRP) in 1978 ontstaan is uit een gevoel van crisis: “Ondanks de enorme inspanningen en boeiende bevindingen van de voorgangers was/leek de Ndl. dialectologie onvoldoende geëquipeerd om de onderzoeksvragen van het laatste kwart van de twintigste eeuw adequaat aan te kunnen pakken.” (zie ook Weijnen, 1975)

Uit dit crisisgevoel kwam de behoefte voort om de dialectologie van het Nederlands beter in te bedden in de stand van de moderne taalwetenschap dan het geval was geweest voor eerdere atlassen zoals de *Reeks Nederlandse Dialectatlassen*. Men wilde dat de nieuw te verzamelen resultaten antwoord zouden geven op “onderzoeksvragen [waarmee] de dialectoloog in de komende decennia geconfronteerd zal worden”; het betrof hier onder andere vragen uit de “systeemlinguïstiek en taaltypologie”, waarvan Goeman & Taeldeman (1996) ook enkele voorbeelden geven:

- (1) a. In welke mate zijn klanksystemen harmonisch/economisch georganiseerd? (N.B. al een oude onderzoeksvraag uit de bloeitijd van het Praagse structuralisme)
- b. Hoe courant komen bepaalde types van klanktegenstellingen voor (b.v. bij vocalen: [voor]-[achter] en bij voorvocalen: [+gespreid]-[+rond])
- c. Zijn er al dan niet recurrente patronen van allofonisering?
- d. Is er een correlatie tussen de stabiliteitsgraad van een segment en z’n lexicale bezetting?

⁽¹⁾ De tekst van Goeman & Taeldeman (1996) wordt voor een belangrijk deel gerecapituleerd in de inleiding van De Wulf et al. (2005).

- e. Van welke aard zijn de condities op morfeemvariatie?
- f. Welke (types van) morfeemstructuurcondities (evt. syllabificatieregels) kunnen voor de Ndl. dialecten variëren? Binnen welke speelruimte kan er zich hier variatie voordoen?
- g. Wat leren ons de Ndl. dialecten m.b.t. allerlei theorieën/hypotheses over regelordening en/of regeltypologie?
- h. Welke elementen kunnen/moeten er deel uitmaken van een typologie van (mor)fonologische regels, enz.

Hoewel deze vragen inmiddels bijna dertig jaar later misschien iets anders geformuleerd zouden worden, zijn ze alle nog steeds relevant voor de moderne klankleer. In Van Oostendorp (2002) wordt betoogd dat de papieren FAND minder geschikt is om dit soort vragen te beantwoorden (zie De Wulf *et al.* (2005) voor een repliek). De GTRP-database die aan de FAND ten grondslag ligt, lijkt om een aantal redenen geschikter. Zo kan de onderzoeker gemakkelijk allerlei zelfgeformuleerde vragen loslaten op de overvloed aan gepresenteerde gegevens. Dat is de methodologie die ik in dit artikel wil volgen.

2. Kenmerkeconomie

2.1. Een oude onderzoeksvraag herleeft

De eerste vraag in (1) is voor de contemporaine theoretisch georiënteerde fonoloog op dit moment mogelijk nog relevanter dan ze dertig jaar geleden was. Ik herhaal de vraag hier voor het gemak:

- (2) In welke mate zijn klanksystemen harmonisch/economisch georganiseerd? (N.B. al een oude onderzoeksvraag uit de bloeitijd van het Praagse structuralisme)

Deze vraag behoeft mogelijk enige toelichting. We kunnen voor elke taal of elk dialect de klinkers inventariseren. Voor het standaard-Nederlands vinden we bijvoorbeeld de volgende klinkerverzameling (exclusief de diftongen [ei, œy, au] en de sjwa, [ə]):

(3)

gespannen/lang				ongespannen/kort			
	voor gespreid	voor gerond	achter		voor gespreid	voor gerond	achter
hoog	i	y	u	hoog			
midden	e	ø	o	midden	ɪ	ʏ	ɔ
laag	a			laag	ɛ		ɑ

2.1. Een oude onderzoeksvraag herleeft

Het is een belangrijke theoretische vraag wat de status is van dit soort tabellen. Er zijn bijvoorbeeld onderzoekers die volhouden dat ze de neiging hebben ‘symmetrisch’ of ‘harmonisch’ of ‘economisch’ te zijn en dat gaten die zich ergens in een tabel dreigen voor te doen onmiddellijk zullen worden gestopt.

We kunnen een dergelijke hypothese onderzoeken door de klinkersystemen van alle Nederlandse dialecten, bijvoorbeeld op de bovenstaande manier, op een rijtje te zetten. Op diezelfde manier zouden we kunnen onderzoeken in hoeverre het toevallig is dat zich in het bovenstaande rijtje gaten voordoen bij de hoge korte klinkers; zijn er dialecten die deze gaten vullen?

(Het antwoord is overigens ja; er zijn zelfs dialecten die alle gaten in het Standaardnederlandse systeem lijken te vullen; Van Oostendorp (2000)).

Bij het beantwoorden van dit soort vragen worden we overigens onmiddellijk geconfronteerd met ingewikkelde vragen over de relatie tussen fonetiek en fonologie. Zo heb ik in het bovenstaande de a als [+achter] gekarakteriseerd, maar men zou kunnen zeggen dat dit eerder is ingegeven door de wens om het systeem symmetrisch te maken dan om fonetisch accuraat te zijn. Op dezelfde manier zou men vraagtekens kunnen stellen bij de karakterisering van bijvoorbeeld de [ɪ] als midden- in plaats van hoge klinker.

Clements (2003) heeft deze ‘oude onderzoeksvraag uit de bloeitijd van het Praagse structuralisme’ recentelijk weer tot leven gebracht; ze speelt op dit moment ook een belangrijke rol omdat ze een van de belangrijkste onderwerpen van debat raakt: de relatie tussen het fonologische systeem en de fonetiek. Clements (2003) stelt dat klankinventarissen in natuurlijke taal onderworpen zijn aan een principe van ‘kenmerkeconomie’ (*feature economy*). Het idee is dat talen maximaal gebruik zullen maken van eenmaal aangewende fonologische kenmerken: als een kenmerk eenmaal gebruikt wordt om een bepaald

contrast uit te drukken, dan zal dit kenmerk ook aangewend worden voor andere contrasten. In de woorden van Clements (2003):

- (4) “A sound S will have a higher than expected frequency in languages that have another sound T bearing one of its features, and vice versa”

Met andere woorden, stel dat we twee talen hebben, T1, die onderliggend een /p/ heeft, en T2, die dit foneem niet kent. Dan is volgens het principe van *feature economy* de kans groter dat T1 (ook) een /b/ heeft, dan dat T2 die heeft (zie ook Van de Weijer & Hinskens, 2004).

Het idee is relevant voor de theorievorming over de werkverdeling tussen fonologie en fonetiek omdat kenmerkeconomie voorspellingen doet die diametraal staan tegenover (bepaalde opvattingen) van de fonetiek. Dit wordt geïllustreerd door de twee hypothetische medeklinkerinventarissen in (5):

- | | | | | | |
|-----|----|-----------------|----------|-----------|---------|
| (5) | a. | | [labial] | [coronal] | [velar] |
| | | [-son] | [-vc] p | t | k |
| | | | [+vc] b | d | g |
| | | [+son] | m | n | ŋ |
| | b. | {t, g, ŋ, φ, k} | | | |

In (5b) verschilt ieder paar van twee willekeurig gekozen segmenten van elkaar op een groot aantal manieren, terwijl veel medeklinkers in (5a) slechts in één kenmerkwaarde van elkaar verschillen. Fonetisch zou men kunnen argumenteren dat (5b) de voorkeur heeft (‘dispersie’), aangenomen dat de articulatorische inspanning die nodig is voor al deze medeklinkers niet noemenswaardig van elkaar verschilt: voor de perceptie is het immers nadelig als een taal een heleboel klanken heeft die veel op elkaar lijken. Onder de aannames van kenmerkeconomie verwachten we daarentegen een systeem als in (5a): het cognitieve systeem dat nodig zou zijn voor de representatie van (5b) is nauwelijks economisch te noemen. Wat betreft medeklinkersystemen lijken de gegevens hier gunstiger voor het kenmerkeconomische model dan voor het fonetische (klinkersystemen lijken daarentegen gevoeliger voor dispersie — op zichzelf een intrigerend feit). Een ander intrigerend, nog onbegrepen, verschil is dat medeklinkers constanter zijn en in ieder geval in het Nederlands minder onderhevig aan dialectvariatie (tegenover drie delen over klinkers in de FAND staat slechts één deel over medeklinkers).

2.2 Kenmerkeconomie in Nederlandse dialecten

De generalisatie in (4) is een statistische. Hoewel ze geïnterpreteerd kan worden als een beperking op interne, cognitieve taalsystemen in Chomskyaanse zin, kan ze eigenlijk alleen bestudeerd worden aan de hand van redelijk grootschalige databases van fonologische systemen. Op basis van een enkel fonologisch systeem valt er niets te zeggen over de juistheid van de voorspellingen die kenmerkeconomie doet, omdat deze implicatieel van aard zijn. Er zijn bijvoorbeeld vier mogelijke systemen van labiale plosieven: /p/ en /b/ (de taal heeft beide labiale plosieven), /ø/ (de taal heeft helemaal geen labiale plosieven), /b/ en /p/ (de taal heeft alleen een stemhebbende, resp. een stemloze labiale plosief). Kenmerkeconomie zegt niet dat een van deze systemen onmogelijk is, maar wel dat een systeem met alleen een /b/ veel onwaarschijnlijker is dan een van de eerste twee systemen.

In zijn eigen werk heeft Clements (2003) het onderzoek naar deze kwesties vooral gericht op macrotypologieën, dat wil zeggen naar databases van foneemsystemen zoals deze te vinden zijn in bijvoorbeeld Ladefoged & Maddieson (1996) of de UPSID database.⁽²⁾ Er is natuurlijk geen dwingende reden waarom dergelijk onderzoek niet ook op microtypologieën zou kunnen worden uitgevoerd. We kunnen ons bijvoorbeeld afvragen in hoeverre de medeklinkerinventarissen van de Nederlandse dialecten georganiseerd zijn volgens principes van kenmerkeconomie. Hiervoor herformuleren we dan Clements' oorspronkelijke hypothese in (4) tot (6):

- (6) A sound S will have a higher than expected frequency in dialects that have another sound T bearing one of its features, and vice versa

Merk op dat we hiermee de oorspronkelijke vraag van de initiators van het GTRP in (1) al bijna geoperationaliseerd hebben. Onze gegevens zijn bovendien veel efficiënter rechtstreeks te onttrekken aan de GTRP-database, waarop we nu de aandacht richten.

3. De GTRP-database

3.1 Segmentinventarissen

Uit de GTRP-database kan op een vrij eenvoudige manier een overzicht van segmentinventarissen van Nederlandse dialecten gedestilleerd worden. Deze database bestaat in verschillende versies. Ik heb gebruik gemaakt van een versie

⁽²⁾ <http://www.langmaker.com/upsidlanguages.htm>

waarin alle gegevens — in totaal ongeveer een miljoen woorden die verzameld zijn in 612 meetpunten — in een groot tekstbestand zijn ondergebracht. De structuur van dit bestand illustreer ik aan de hand van een fragment:

(7)	1.	2.	3.
	E192p	130	6n dr7a2.t
	E192p	131	dr7a2.<d6n
	E192p	132	6n dr7a.>ts2i
	E192p	133	6n dr7o5po2l4
	E192p	134	6n do7y_f
	E192p	135	do7y_v6n
	E192p	136	6n do7y_fi

In de kolom 1 vinden we de Kloeke-code van de vindplaats waar de opname is gedaan; in kolom 2 staat het vraagnummer van de vragenlijst, en in kolom 3 een transcriptie van het antwoord. Deze transcriptie is gemaakt in K-IPA, een voor het GTRP ontwikkelde versie van IPA die alleen gebruik maakt van lettertekens die met een ‘gewoon’ toetsenbord kunnen worden gemaakt — er wordt bijvoorbeeld geen gebruik gemaakt van het bijzondere teken ‘ə’, maar in plaats daarvan van een ‘6’.⁽³⁾

3.1. Segmentinventarissen

Een aardige eigenschap van het K-IPA is dat alle ‘hoofdsymbolen’ van het IPA gerepresenteerd worden door een letter, terwijl diacritische tekens worden gerepresenteerd door interpunctietekens en cijfers. De enige uitzondering is de zojuist genoemde sjwa = ’6’. Dit maakt het relatief eenvoudig om de gegevens in de derde kolom te ontleden in van elkaar onderscheiden fonetische ‘symbolen’. Een computerprogramma kan zo vrij eenvoudig vaststellen hoeveel verschillende symbolen er zijn gebruikt in de transcriptie van een dialect.⁽⁴⁾

De gemiddelde grootte van de segmentinventaris per dialect blijkt op deze manier gesteld te kunnen worden op 203. Dat is dus het gemiddelde aantal verschillende symbolen dat een transcribent voor een dialect gebruikt heeft. Een eerste conclusie die we hieruit kunnen trekken is dat de transcripties over het algemeen tamelijk ‘narrow’ of fonetisch zullen zijn geweest. Het standaard-Nederlands kent tussen de veertig en de vijftig ‘fonemen’ – afhankelijk van hoe

⁽³⁾ Meer informatie over de database, en over K-IPA, kan gevonden worden op <http://www.meertens.knaw.nl/projecten/mand/>.

⁽⁴⁾ De gebruikte scripts zijn geschreven in de scripttaal Python, en zijn op aanvraag beschikbaar bij de auteur.

men telt — en het is onwaarschijnlijk dat de Nederlandse dialecten gemiddeld meer dan vier keer zoveel taalkundig relevante onderscheidingen maken.

Interessant is ook het kaartje in (8). Hierin zijn onderscheiden symbolen gebruikt voor de inventarissen die groter zijn dan gemiddeld en inventarissen die kleiner zijn dan gemiddeld.



(8)

Een opvallend kenmerk van dit kaartje is dat de grens tussen de twee symbolen vrij nauwkeurig de landsgrens volgt. In het bijzonder hebben vrijwel alle plaatsen in Vlaanderen een kleiner aantal segmenten in hun inventaris dan gemiddeld.⁽⁵⁾

3.2 Structuur van een segment

Behalve de grootte van de segmentinventarissen kunnen we natuurlijk ook hun interne structuur bestuderen, om hiermee bijvoorbeeld de voorspellingen van de theorie over kenmerkeconomie te toetsen. Nu valt niet te verwachten dat Nederlandse dialecten op grote schaal verschillen in het al dan niet toestaan van bijv. /b/ naar gelang de afwezigheid van /p/: aan de hand van de GTRP-database stellen we zelfs in een handomdraai vast dat alle Nederlandse dialecten beide segmen-

⁽⁵⁾ Waarschijnlijk hangt dit effect sterk samen met verschillen tussen transcribenten. Er zijn transcribenten met gemiddeld 70 onderscheidingen, en anderen met gemiddeld meer dan 350 onderscheidingen per dialect. Dit bemoeilijkt natuurlijk de vergelijkbaarheid van de verschillende clusters. Zie Hinskens & van Oostendorp (2006) voor een uitgebreidere bespreking en statistische analyse van dit en verwante punten.

ten in hun inventaris hebben. We kijken daarom in eerste instantie naar fijnere onderscheidingen (Zie Hinskens & van Oostendorp, 2006, voor een toepassing op palatalisering en velarisering van -nt- en -nd-clusters.) In deze afdeling richt ik me daarom op dentalisering van coronalen. Dit secundaire kenmerk wordt in GTRP weergegeven met het diacriticum ‘{’, en wordt alleen gebruikt na een van de coronale medeklinkers {t, d, s, n, l, z}.⁽⁶⁾

De vraag die zich nu voordoet uit het oogpunt van kenmerkeconomie is de volgende (waarbij $\underset{\cdot}{d}$, $\underset{\cdot}{t}$ staan voor een gedentaliseerde d en t):

- (9) Is het waar dat de kans op d groter is in dialecten mét t, dan in dialecten zonder $\underset{\cdot}{t}$?

Soortgelijke vragen kunnen we natuurlijk voor ieder willekeurig paar van dentalen stellen, maar gezien de relatieve frequentie van $\underset{\cdot}{t}$, en de relatief oncontroversiële nabijheid van $\underset{\cdot}{d}$ en $\underset{\cdot}{t}$, is dit een goed paar om mee te beginnen.

De onderstaande tabel geeft de resultaten van een uiterst eenvoudige statistische analyse van de segmentinventarissen in het GTRP-materiaal:⁽⁷⁾

(10)		$\underset{\cdot}{t}$	$\neg\underset{\cdot}{t}$	$\chi^2=70$
	$\underset{\cdot}{d}$	14	11	$p<0,001$
	$\neg\underset{\cdot}{d}$	41	546	

Deze tabel moet als volgt gelezen worden: in de kolom onder $\underset{\cdot}{t}$ staan alle dialecten met een dentale stemloze plosief (dat zijn er in totaal $14+41=55$); in de kolom onder $\neg\underset{\cdot}{t}$ staan alle dialecten zonder zo’n segment ($11+546=557$).

In de rij achter $\underset{\cdot}{d}$ staan alle dialecten met een dentale stemhebbende plosief ($14+11=25$), achter $\neg\underset{\cdot}{d}$ alle dialecten zonder dat segment ($41+546=587$).

Het belangrijke punt is dat het aantal dialecten dat zowel $\underset{\cdot}{t}$ als $\underset{\cdot}{d}$ heeft, veel groter is dan we zouden mogen verwachten op basis van de distributie van die segmenten op zichzelf:

- (11) a. Het percentage dialecten met $\underset{\cdot}{d}$ is $(14+11) / 612 = 4$ procent.
b. Het percentage dialecten met $\underset{\cdot}{t}$ is $(11+41) / 612 = 9$ procent.

⁽⁶⁾ Tijdens de eerste keer dat ik het hier gebruikte script liet draaien bleek er ongeveer een vijftal keren een diacriticum ‘{’ geplaatst te zijn na een andere klank dan een coronale medeklinker. Bij nadere inspectie bleek het hier in alle gevallen te gaan om een verschrijving, die inmiddels hersteld is in de GTRP-database.

⁽⁷⁾ Er is geen reden om te veronderstellen dat er sprake is van een regionaal effect in dit geval: dentalen worden in een groot deel van het taalgebied genoteerd.

- c. De kans dat we *zowel* d als t vinden zou daarom 0.04×0.09 moeten zijn. We zouden dus verwachten dat er $0.04 \times 0.09 \times 612 = 2$ dialecten zijn met deze combinatie. Maar er zijn er meer, namelijk 14.

Merk op dat de statistische voorspelling *niet* is dat er geen dialecten kunnen zijn met alleen een stemhebbende dentale plosief zonder een stemloze, of omgekeerd; allebei de groepen komen ook voor — de stemloze dentaal zonder de stemhebbende is zelfs nog frequenter dan de combinatie van de twee. Toch is in de GTRP-database een neiging tot economisch gebruik van dentalisering aantoonbaar.

3.3 Problemen

Helaas kunnen we uit het voorafgaande toch nog niet zonder meer een (deel)antwoord destilleren — hoe klein ook — op de ‘oude onderzoeksvraag uit de bloeitijd van het Praagse structuralisme’, ook niet in de door Clements (2003) geactualiseerde vorm.

Het grootste probleem is dat de transcribenten veel van elkaar verschillen, en het bovenstaande verschil mogelijk aan hen moet worden toegeschreven. Het is bijvoorbeeld mogelijk dat een transcribent die een oor heeft dat de t observeert, een veel grotere kans heeft om ook de d te horen dan een transcribent die dit niet doet.

Hierbij moet worden opgemerkt dat informele observatie van het materiaal doet vermoeden dat er minstens twee ‘scholen’ zijn geweest van transcribenten, een ‘Amsterdamse’ en een ‘Gentse’, waarbij de eerste geneigd was nauwkeurig fonetisch detail te noteren, terwijl de laatste veel meer impressionistisch-fonologisch te werk ging (Taeldeman p.c.; Goeman p.c.) Nu is de fonologische status van dentalisering in de Nederlandse dialecten over het algemeen tamelijk twijfelachtig.

Een volgende stap in de bestudering van de rol van kenmerkeconomie in de structuur van segmentinventarissen zou daarom bijvoorbeeld gebaseerd kunnen zijn op het begrip ‘zuinigheidsindex’, eveneens afkomstig van Clements (2003):

- (12) “Feature economy can be quantified in terms of a measure called the economy index.

Given a system using F features to characterize S sounds, its economy index E is given by the expression $E = S/F$ ”

De zuinigheidsindex geeft onder meer een maat om talen met elkaar te vergelijken, en we kunnen ons daardoor nu afvragen of kleine, meer ‘fonologisch opgebouwde’ inventarissen economischer zijn dan grote, meer ‘fonetische’. Als kenmerkeconomie een eigenschap is van fonologische systemen, zouden we dat verwachten; maar als het samenhangt met de gevoeligheid van individuele transcribenten verwachten we dat veel minder.

Nu wijzen eerste berekeningen uit dat kleine inventarissen inderdaad economischer georganiseerd zijn dan grote: voor de laatste (inventarissen met meer dan het gemiddelde aantal segmenten >203) is de zuinigheidsindex gemiddeld 5,1; voor kleine inventarissen (<204) komt dit getal uit op 2,6. Dat lijkt een bemoedigend resultaat voor de onderzoeker van kenmerkeconomie aan de hand van het GTRP-materiaal.

Overigens blijft bij dit alles nog steeds onopgehelderd wat de precieze status is van een individuele transcriptie in GTRP: fonologisch of fonetisch. Het profijtelijkst zijn de gepresenteerde gegevens waarschijnlijk te zien als ruw materiaal, dat op basis van verder onderzoek (veldwerk, literatuuronderzoek) nader moet worden uitgewerkt. Het zou mooi zijn als er ooit een al dan niet elektronische atlas kwam waarin we bijvoorbeeld fonologische segmentinventarissen zouden kunnen naslaan. Dit zou een dialectologische tegenhanger van de al genoemde UPSID-database kunnen zijn.

Een dergelijk gegevensbestand zou in mijn ogen het GTRP-materiaal niet zozeer moeten vervangen als wel aanvullen en verrijken. Voor zover het waar is dat er twee scholen zijn geweest in de transcripties van de veldwerkopnamen zou het mooi zijn als ieder dialect ook van een transcriptie vanuit de andere school zou kunnen worden voorzien: de Nederlandse opnamen door een fonologische, de Vlaamse door een fonetische. Ook andere uitbreidingen liggen in het verschiet: verrijking van de database met de oorspronkelijke geluidsbestanden bijvoorbeeld, het compatibel maken met andere bronnen, zoals de hopelijk ooit te digitaliseren *Reeks Nederlandse Dialectatlassen*, de databases met ‘banden vrij gesprek’ zoals deze onder andere op het Meertens Instituut en in Gent te vinden zijn, het materiaal van de *Syntactische Atlas van de Nederlandse Dialecten* en dat van een nog te ontwikkelen atlas over intonatievariatie. Bovendien worden in het buitenland ook databases en corpora voor de studie van fonologische microvariatie aangelegd — zoals bijvoorbeeld het project *Phonologie du Français Contemporain* in het Franse taalgebied — waarmee aansluiting kan worden gezocht, al is het maar om te werken aan een gemeenschappelijke oplossing van technische problemen.

Het wensenlijstje in de vorige alinea is zo groot dat het weinig reëel is om te veronderstellen dat het helemaal zal kunnen worden vervuld. Misschien wordt het tijd dat, net als in 1975, een groep dialectfonologen weer bij elkaar komt om te inventariseren welke volgende stappen genomen kunnen worden na de FAND en de GTRP-database.

Bibliografie

CLEMENTS, G.N.

2003, 'Feature Economy in Sound Systems'. *Phonology*, 20, 3: 287–333. URL [http://ed268.univ-paris3.fr/lpp/publications/2003/Clements Feature economy.pdf](http://ed268.univ-paris3.fr/lpp/publications/2003/Clements%20Feature%20economy.pdf).

DE WULF, CHRIS, JAN GOOSSENS, & JOHAN TAEDEMAN

2005, *Fonologische Atlas van de Nederlandse Dialecten. Deel IV. De consonanten*. Gent: Koninklijke Academie voor Nederlandse Taal- en Letterkunde.

GOEMAN, TON & JOHAN TAEDEMAN

1996, 'Fonologie en morfologie van de Nederlandse dialecten. Een nieuwe materiaalverzameling en twee nieuwe atlasprojecten'. *Taal en Tongval*, 48, 1: 38–59.

GOOSSENS, JAN, JOHAN TAEDEMAN, & GERT VERLEIJEN

1998, *Fonologische Atlas van de Nederlandse Dialecten. Deel I*. Gent: Koninklijke Academie voor Nederlandse Taal- en letterkunde.

GOOSSENS, JAN, JOHAN TAEDEMAN, & GERT VERLEIJEN

2000, *Fonologische Atlas van de Nederlandse Dialecten. Deel II. De Westgermaanse korte vocalen in open lettergreep. Deel III. De Westgermaanse lange vocalen en diftongen*. Gent: Koninklijke Academie voor Nederlandse Taal- en Letterkunde.

HINSKENS, FRANS & MARC VAN OOSTENDORP

2006, 'Bronnen van variatie in het GTRP'. *Taal en Tongval*.

LADEFOGED, PETER & IAN MADDIESON

1996, *The Sounds of the World's Languages*. Oxford: Blackwell.

VAN OOSTENDORP, MARC

2000, *Phonological Projection*. Berlin/New York: Mouton de Gruyter.

VAN OOSTENDORP, MARC

2002, 'Recensie van Goossens, Taeldeman en De Wulf (2000)'. *TNTL*, 118:4–8.

VAN DE WEIJER, J.M. & F. HINSKENS

2004, 'Markedness and Segmental Complexity: A Cross-Linguistic Study'. In: *GLOW Phonology Workshop, Thessaloniki*.

WEIJNEN, A.

1975, 'Crisis in de dialectkunde'. *Taal en Tongval*, 27: 110–117.