

Een computermodel van verbale zelfmonitoring

Robert J. Hartsuiker¹

Herman H. J. Kolk

Nijmeegs Instituut voor Cognitie en Informatie, KU Nijmegen

Verbale zelfmonitoring verwijst naar het samenstel van cognitieve processen waarmee de spreker zijn spraak op juistheid controleert. Volgens de theorie van Levelt (1983) gebruiken wij voor deze taak het taalbegripsysteem. Dit systeem controleert zowel onze overte spraak als onze 'interne spraak'. Wij hebben een computermodel geïmplementeerd dat gebaseerd is op deze theorie. Met dit model hebben we getracht om empirisch verkregen verdelingen van de tijdsintervallen tussen de momenten van fout, interruptie en correctie te simuleren, alsmede de effecten van spraaksnelheid op deze intervallen. Het model slaagt er in dergelijke verdelingen na te bootsen, gegeven een aantal aannames over de werking van monitoring: (1) interruptie en correctie zijn twee parallelle processen; (2) in snelle spraak worden zowel taalproductie als taalbegripsprocessen versneld.

Inleiding

Sprekers bewaken continu de kwaliteit van hun eigen spraak. Hierbij streven ze naar zowel vloeiendheid als accuratesse. Het samenstel van cognitieve processen dat deze kwaliteitscontrole uitvoert, wordt aangeduid met "verbale zelfmonitoring". Dit zijn processen die zorgen voor detectie van een fout en processen die vervolgens in actie komen: zelfinterruptie en foutcorrectie. Zelfmonitoring is een belangrijk onderwerp binnen de cognitieve wetenschappen: mensen controleren immers de kwaliteit van al hun handelingen en ondernemen acties wanneer er iets fout gaat. Binnen de psycholinguïstiek wordt monitoring vaak aangehaald om te verklaren waarom sommige versprekingen (bijvoorbeeld versprekingen die leiden tot een bestaand woord) vaker voorkomen dan andere. Het idee is dat door zelfmonitoring sommige versprekingen onderschept worden voordat ze daadwerkelijk zijn gearticuleerd, maar dat de kans op onderschepping kleiner is voor versprekingen die een bestaand woord vormen (Baars, Motley, & MacKay, 1975). Binnen de taalpathologie worden sommige symptomen verklaard als het gevolg van een disfunctionerende monitor, of als gevolg van een interactie tussen gestoorde spraakplanning en normale monitoring. Hierbij kan men denken aan stotteren (Postma & Kolk, 1993), jargonafasie (Marshall, Robson, Pring & Chiat, 1998) en verbale hallucinaties bij schizofrenie

Correspondentie-adres: University of Edinburgh, Dept. of Psychology,
7 George Square, EH8 9JZ, Edinburgh, Scotland, rob.hartsuiker@ed.ac.uk

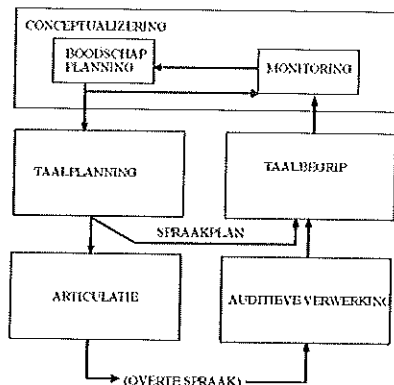
(Frith & Done, 1988).

Bij verbale zelfmonitoring zijn drie deelprocessen van belang: (1) het detecteren van een fout in de eigen spraak; (2) het interromperen van de eigen spraak; (3) het plannen van een correctie. We illustreren deze begrippen aan de hand van een voorbeeld uit de experimentele literatuur. Stel dat een proefpersoon een route beschrijft door een netwerk van gekleurde objecten (Levelt, 1983; Oomen & Postma, in druk). Het zal dan nogal eens voorkomen dat de proefpersoon zich versprekt, zijn spraak staakt en opnieuw begint, zoals in voorbeeld (1).

(1) ..dan neem je de h ^ verticale lijn

In dit voorbeeld produceert de spreker eerst het verkeerde foneem /h/. Het is goed mogelijk, gegeven de taak, dat hij op een pre-articulatoir niveau van planning "horizontale" had geproduceerd. De spreker detecteert echter een fout in het (interne) spraakplan, en hij interrumpeert zichzelf. Het moment van interruptie is hier weergegeven door de "^". Het interval tussen het begin van de fout in het spraaksignaal en het interruptiemoment, noemen we het fout/interruptie-interval. Na de interruptie corrigeert de proefpersoon zichzelf: hij produceert een "correctie", nl. *verticale*. Het interval tussen interruptie en correctie, noemen we het interruptie/correctie-interval. Deze tijdsintervallen zullen later in dit artikel van belang blijken bij het toetsen van een theorie over monitoring.

Een invloedrijke theorie van zelfmonitoring is de "perceptuele lus" theorie (PLT) van Levelt (1983). We maken deze theorie duidelijk aan de hand van Figuur 1. De figuur toont een "blauwdruk van de spreker" (Levelt, 1989). Deze blauwdruk bestaat uit een aantal verwerkingscomponenten die betrokken zijn bij het produceren en begrijpen van taal. Productie begint met het samenstellen van een boodschap. Vervolgens dient de boodschap vertaald te worden in een welgevormde zin. Dat is de taak van de component "Taalplanning". De uitvoer van deze component dient als invoer voor de "Articulator". Deze component verzorgt de motorische planningsprocessen bij articulatie. Aan de rechterkant van Figuur 1 zijn er componenten voor taalbegrip. Er is een component voor auditieve verwerking. Deze component genereert een fonetische representatie als uitvoer. De Taalbegripscomponent gebruikt de fonetische representatie voor o.a. het herkennen van woorden, het integreren van de woorden in een zin, en het afleiden van thematische rollen ("wie deed wat met wie?"). De perceptuele lus theorie berust op drie veronderstellingen. De eerste veronderstelling is dat er naast een extern monitoringkanaal (luisteren naar je eigen gerealiseerde spraak) ook een intern monitoringkanaal is, dat spraak evalueert voordat het gearticuleerd is. Dit kan aannemelijk gemaakt worden aan de hand van het eerder gegeven voorbeeld (1). In dit voorbeeld werd het verkeerde woord al afgebroken na het articuleren van één foneem. Het is onwaarschijnlijk dat deze fout via het externe kanaal werd gedetecteerd, omdat de processen van auditieve verwerking, taalbegrip, vaststelling van de fout, en interromperen meer tijd kosten dan het articuleren van het volgende foneem. De tweede veronderstelling is dat beide monitoringkanalen ge-



Figuur 1. Overzicht van verwerkingscomponenten en processen bij het spreken, luisteren en monitoren gebaseerd op Levelt (1989). Een aantal van de componenten die Levelt onderscheidt zijn, vanwege de overzichtelijkheid, weggelaten.

bruik maken van het *taalbegripssysteem*. De interne "lus" verbindt Taalplanning direct met Taalbegrip. De externe "lus" maakt een omweg via Articulatie en Auditieve verwerking (zie Figuur 1). Ten derde zouden alle spraakplanningsprocessen zo snel mogelijk stil gelegd worden nadat een fout is ontdekt. Deze laatste veronderstelling wordt de "hoofdinterruptieregel" genoemd (Nooteboom, 1980).

Hoeveel ondersteuning is er voor deze drie veronderstellingen? De theorie dat er een intern monitoringkanaal is, ondervindt brede steun (bijvoorbeeld Dell & Repka, 1992; Lackner & Tuller, 1979; Motley, Camden, & Baars, 1982; Postma & Kolk, 1992; Postma & Noordanus, 1996). Zo vroegen Lackner en Tuller (1979) aan proefpersonen om op een knop te drukken wanneer zij een fout in hun eigen spraak waarnamen. In een conditie konden de proefpersonen hun eigen spraak echter niet horen, want zij droegen een koptelefoon waaruit zeer luide ruis kwam. Toch slaagde men er in veel fouten te ontdekken. De reactietijden waren bovendien korter dan in een conditie zonder ruis. Dat is in overeenstemming met het idee van een intern monitorkanaal. Immers, in de conditie met ruismaskering heeft de proefpersoon alleen het interne kanaal tot zijn beschikking. Dat fouten sneller ontdekt worden kan verklaard worden uit het feit dat het interne kanaal uit minder verwerkingsstappen bestaat dan het externe kanaal.

Een andere bron van evidentie voor het idee van een interne monitor komt uit werk met hersenpotentialen (bijvoorbeeld Bernstein, Scheffers, & Coles, 1995). Deze auteurs vonden een component in het EEG-sigitaal die optreedt wanneer er een discrepantie bestaat tussen de bedoelde en daadwerkelijke reactie. Deze component, de Error Related Negativity (ERN) is voornamelijk gevonden in reactietijdtaken waarbij de proefpersoon de keuze had uit twee reacties. Deze component treedt zo vroeg op dat de representatie van de daadwerkelijke reactie niet via visuele of tactiele informatie geconstrueerd kan worden. Het is daarom aannemelijk dat deze representatie

via een intern monitoringkanaal geconstrueerd wordt. Merk overigens wel op dat deze studies monitoring van handelen betroffen ("druk ik wel op de juiste knop?") en niet monitoring in het verbale domein ("zeg ik wel wat ik wil zeggen?").

De tweede aanname van de perceptuele lus theorie is dat monitoring gebruik maakt van het taalbegripssysteem. Er is nog maar weinig ondersteuning voor deze aanname. Een uitzondering is een studie van McGuire, Silbersweig en Frith (1996) die activiteit in de hersenen zichtbaar maakten met Positron Emissie Tomografie (PET) terwijl proefpersonen woorden voorlasen. In de ene conditie hoorden proefpersonen via de koptelefoon hun eigen spraak, maar die spraak werd verstoord (de toonhoogte werd artificieel gemanipuleerd). In de andere conditie las men ook de woorden voor, maar via de koptelefoon hoorde men de stem van een andere spreker die dezelfde woorden las. Beide condities werden vergeleken met een controleconditie waarin de proefpersonen voorlasen en hun eigen stem hoorden. De rationale was dat in de twee experimentele condities veel monitoringactiviteit zou optreden, omdat er steeds een discrepantie was tussen wat men dacht te produceren en wat men daadwerkelijk hoorde. In beide experimentele condities werd activiteit gevonden in de temporaalschors in beide hemisferen, in gebieden die ook bij taalperceptie actief worden. Dit is geen eenvoudig bijverschijnsel van het feit dat men een taalperceptietaak uitvoerde, want deze activiteit trad op nadat het patroon van activiteit in de controleconditie van dat in de experimentele conditie was afgetrokken. Gebieden die betrokken zouden zijn bij semantische of fonologische processen in taalproductie werden echter niet actief. Deze bevindingen zijn in overeenstemming met de PLT. Merk echter wel op dat het experimentele paradigma beperkt was tot het niveau van het enkele woord, dat de discrepantie tussen doelrepresentatie en geproduceerde representatie beperkt was tot akoestische verschillen en verschillen in de stem, en dat deze discrepantie alleen gedetecteerd kon worden door de externe monitoringlus.

Het derde uitgangspunt van de PLT is de hoofdinterruptieregel. In zijn oorspronkelijke vorm behelst deze regel dat zo gauw een fout is gedetecteerd de modules voor conceptualisering, taalplanning en articulatie tegelijk zouden stoppen met hun activiteit. Het effectueren van een dergelijke interruptie zou enige tijd kosten, door Levelt geschat op zo'n 200 ms. De tijd na de interruptie zou gebruikt worden voor het plannen van de zelfcorrectie. Het interruptie/correctie-interval reflecteert deze planningstijd. Er zijn echter data van Blackmer en Mitton (1991) en van Oomen en Postma (in druk) die in strijd zijn met deze interruptieregel. Beide studies toonden aan dat interruptie/correctie-intervallen vaak zeer kort zijn en dat dikwijls de correctie onmiddellijk op de interruptie volgde (een interval van 0 ms dus). De hoofdinterruptieregel kan dergelijke intervallen niet verklaren, want het kan immers niet zo zijn dat het plannen van de correctie geen tijd kost.

Kortom, de ondersteuning voor de perceptuele lus theorie is niet geheel en al overtuigend. Er is veel ondersteuning voor de veronderstelling van een pre-articulatorische monitor. Er is echter weinig ondersteuning voor de veronderstelling dat zowel het pre-articulatorische als het externe monitoringkanaal gebruik maken van het taalbegripssysteem. Tenslotte is er duidelijke evidentie tegen de hoofdinterruptiere-

gel. Een probleem is bovendien dat de perceptuele-lustheorie niet ver genoeg uitwerkt is om duidelijke voorspellingen te kunnen afleiden. Een vooralsnog onbeantwoorde vraag is bijvoorbeeld of een pre-articulatoire monitor die gebruik maakt van taalbegrip wel snel genoeg is om de zeer korte fout/interruptie-intervallen die in spraak gevonden worden te verklaren. Onduidelijk is ook hoe de processen van interruptie en correctie in de tijd gecoördineerd worden. Een andere vraag is of effecten van spreksnelheid op de duur van fout/interruptie- en interruptie/correctie-intervallen (Oomen & Postma, in druk) uit de theorie volgen.

Om deze vragen te beantwoorden, hebben we de perceptuele-lustheorie geïmplementeerd als een computermodel. Daardoor zijn we gedwongen al onze aannames expliciet te maken, en kunnen we duidelijke voorspellingen afleiden. Het model telt schattingen van de tijdsduren van de deelprocessen bij elkaar op, en produceert voorspellingen over de timing van interruptie en reparatie. Een belangrijke wijziging ten opzichte van de oorspronkelijke theorie is de aangepaste hoofdinterruptieregel. De nieuwe regel stelt dat interrumpen en corrigeren twee parallel verlopende processen zijn, die beide meteen beginnen na de detectie van een fout. Volgens de oude interruptieregel werd de correctie pas gepland nadat de interruptie een feit was.

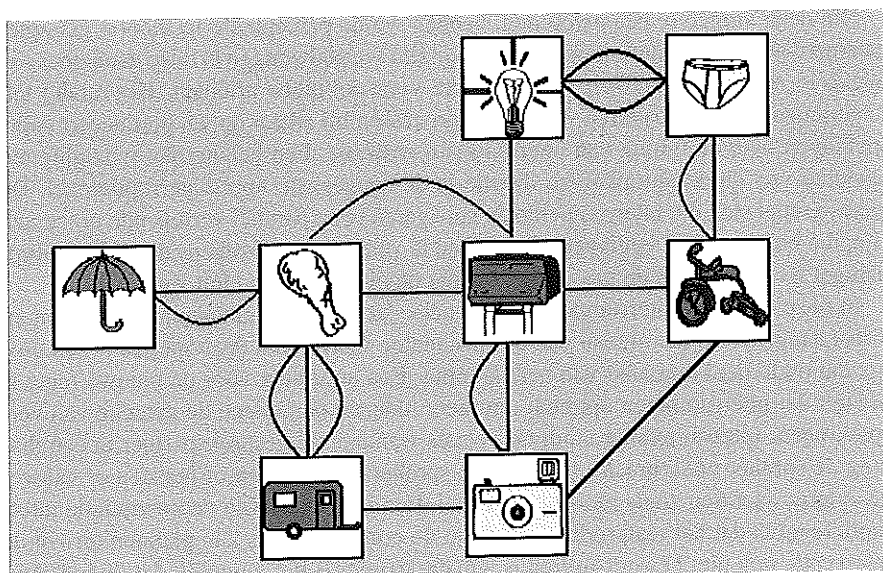
In onze optiek is interrumpen een proces dat buiten de taalplanning om gaat en direct aangrijpt op het niveau van de Articulator. Corrigeren daarentegen is een proces dat onderdeel uitmaakt van taalplanning. Een argument voor de nieuwe regel is het feit dat interrumpen optreedt om meer redenen dan alleen de detectie van versprekingen. Interrumpen vindt ook plaats wanneer de gesprekspartner in een dialoog de spreker onderbreekt, wanneer de spreker moet niezen of hoesten of vanwege externe factoren zoals plotseling lawaai. Een implicatie van de gewijzigde interruptieregel is dat het taalproductiesysteem de correctie reeds kan plannen terwijl de interruptie nog niet bewerkstelligd is. Op die manier is het mogelijk dat de tijdsduur tussen interruptie en correctie soms minimaal is.

Met het model hebben we een aantal simulaties uitgevoerd. Doel van onze simulatiestudies was om te toetsen of een perceptie-gebaseerd monitoringsysteem in overeenstemming is met metingen van fout/interruptie- en interruptie/correctie-intervallen. De belangrijkste vragen daarbij waren: (1) Hebben we werkelijk de aanname van een intern monitoringkanaal nodig om de verdeling van deze tijdsintervallen te kunnen nabootsen?; (2) Zo ja, is het taalbegripssysteem snel genoeg om deze intervallen te produceren?; (3) Kan de gewijzigde hoofdinterruptieregel de supersnelle interruptie/correctie-intervallen produceren?; (4) Onderzoek van Oomen en Postma (in druk) heeft aangetoond dat fout/interruptie-intervallen en interruptie/correctie-intervallen korter worden met spreksnelheid. Volgen deze effecten uit het model?

In de rest van dit artikel zullen we eerst de experimentele gegevens bespreken die het model dient te simuleren. Daarna beschrijven we het model in meer detail. Vervolgens rapporteren we drie simulaties.

Experimentele gegevens

De te simuleren data zijn fout/interruptie- en interruptie/correctie-intervallen. Deze data zijn vergaard in een experiment van Oomen en Postma (in druk). Deze auteurs vroegen aan 24 proefpersonen (studenten aan de Universiteit Utrecht) om routes te beschrijven door een netwerk van gekleurde objecten. Een voorbeeld van een netwerk dat Oomen en Postma hebben gebruikt is gegeven in Figuur 2.



Figuur 2. Voorbeeld van een netwerk dat proefpersonen beschreven in het experiment van Oomen en Postma (in druk).

De sprekers moesten zowel de objecten benoemen als de verbindingen daartussen. Zowel de volgorde waarin als de snelheid waarmee de objecten en verbindingen benoemd moesten worden, werd geregeld door middel van een gekleurde stip die door het netwerk bewoog. De spraak werd gemeten in twee condities: normaal spreektempo (277 ms per syllabe) en snel spreektempo (222 ms per syllabe). De spraak werd opgenomen en getranscribeerd, waarbij versprekingen en zelfcorrecties werden genoteerd. Een subset van deze incidenten werd geanalyseerd, waarbij de duur van fout/interruptie- en interruptie/correctie-intervallen gemeten werd. Uit deze subset selecteerden wij incidenten die naar onze mening ondubbelzinnig correcties van versprekingen waren, en die geen deel uitmaakten van complexe incidenten (waarbij de zelfcorrectie op zijn beurt weer onderbroken en gecorrigeerd werd).

In Tabel 1 rapporteren we verdelingskenmerken van de fout/interruptie- en interruptie/correctie-intervallen

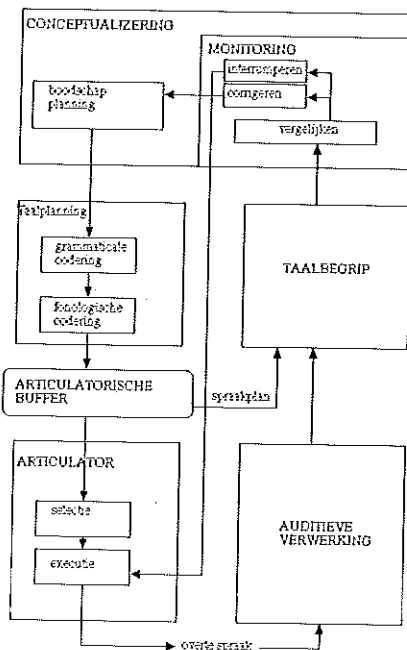
Gemiddelde	Standaard	N	Afwijking
Fout/Interruptie			
interval	321	195	95
Normaal	371	207	38
Snel	288	180	57
Interruptie/Correctie			
interval	266	237	97
Normaal	276	236	38
Snel	259	239	59

Tabel 1. Gemiddelden, Standaard Afwijkingen en Steekproefgroottes (N) in empirische gevonden fout/interruptie- en interruptie/correctie-intervallen, in normale en snelle spraak.

Zoals blijkt uit Tabel 1 zijn zowel fout/interruptie- als interruptie/correctie-intervallen korter in de conditie met snelle spraak. Oomen en Postma (in druk) vonden dat deze verschillen statistisch significant zijn.

Het Model

De architectuur van het model is weergegeven in Figuur 3.



Figuur 3. Overzicht van verwerkingscomponenten en processen in ons model.

Het model kent dezelfde globale structuur als de blauwdruk in Figuur 1. Er is echter een aantal componenten uitgewerkt. Zo wordt binnen de Taalplanning nu een onderscheid gemaakt tussen grammaticale codering (het oproepen van woorden en het construeren van een zinsstructuur met die woorden) en fonologische codering (het bepalen van de fonologische vorm van de zin). Deze subcomponenten werden ook door Levelt (1989) onderscheiden. De Articulator is verdeeld in twee subcomponenten, een voor het genereren van een "motorisch plan", een set instructies voor de motorische spraakprocessen (Selectie), en een voor het uitvoeren van dat plan (Executie). De monitoringcomponent voert drie processen uit. Het eerste proces vergelijkt de bedoelde representatie met de waargenomen representatie. Als er een discrepantie is ontdekt tussen deze representaties worden de andere twee processen uitgevoerd: het interrumperen van de spraak en het plannen van de correctie. De door ons voorgestelde interruptieregel neemt aan dat het interrumperen en het plannen van de correctie gelijktijdig aan elkaar plaats vinden. Het duurt ongeveer 200 ms voor de interruptie geëffectueerd is, en in die tijd begint de spreker al met het plannen van de zelfcorrectie.

Het model berekent fout/interruptie- en interruptie/correctie-tijden door schattingen uit de literatuur te nemen voor de duur van elk van de processen die betrokken zijn bij productie en perceptie. Aan deze schattingen wordt ruis toegevoegd: dit betekent dat de schattingen voor de tijdsduur van elk deelproces volgens een toevalsverdeling korter of langer gemaakt worden. Vervolgens worden deze tijdsduren opgeteld, en worden er schattingen afgeleid voor (1) het begin van het uitspreken van het (foute) woord; (2) het moment van interruptie als de fout gedetecteerd wordt door het interne kanaal; (3) het moment van interruptie als de fout gedetecteerd wordt door het externe kanaal; (4) het begin van de zelfcorrectie.

Met behulp van momenten (2) en (3) kunnen, voor iedere fout, twee fout/interruptie-intervallen berekend worden. Een interval als de fout gedetecteerd zou zijn via het interne kanaal, en een interval als de fout gedetecteerd zou zijn via het externe kanaal. Maar hoeveel procent van de daadwerkelijk gecorrigeerde fouten wordt door ieder kanaal gedetecteerd? We hebben dit probleem opgelost door een procedure te ontwikkelen waarmee we op basis van empirische data een schatting kunnen maken van de bijdrage van beide kanalen aan een gegeven set correcties. De empirische gegevens zijn het aantal fouten dat daadwerkelijk is gecorrigeerd en het aantal "coverte" correcties van alle correcties. Een coverte correctie is een incident waarbij een fout gerepareerd wordt voordat deze is uitgesproken. Volgens de theorie van Postma en Kolk (1993) leiden dergelijke incidenten tot niet-vloeiendheden, zoals herhalingen van fonemen of syllaben, pauzes en prolongaties. Een uitgebreide beschrijving van onze partitioneringsprocedure is in Hartsuiker en Kolk (in druk) te vinden. Voor de huidige dataset betekent het dat in ongeveer 70% van de fouten het interval gebruikt moet worden dat uitgaat van de externe lus.

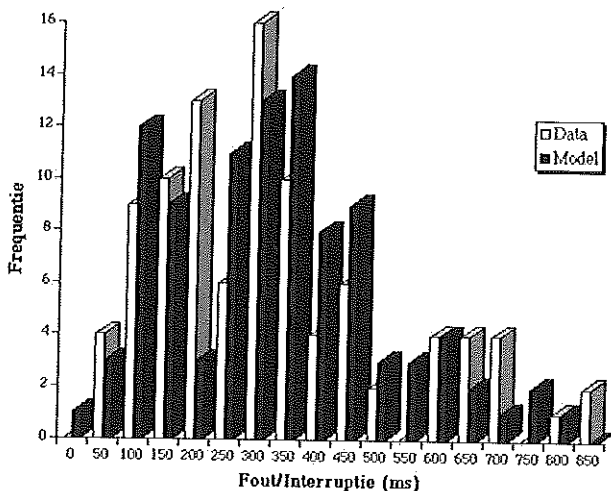
Test van het model

Hoe goed voldoet het model? In deze sectie bespreken we een drietal simulaties. De eerste simulatie betreft de vraag of het model de empirische verdeling van

fout/interruptie-intervallen, zoals aangetroffen door Oomen en Postma (in druk) kan nabootsen. Om deze vraag te beantwoorden hebben we drie modelparameters gevarieerd.

De eerste parameter was de partitionering in interne-lus en externe-lus detecties. We hebben drie waarden voor deze parameter getest: 100% externe lus, 67% externe lus (dus in overeenstemming met ons partitioneringsalgorithme) en 50% externe lus. We concludeerden dat de gemiddelde fout/interruptie-intervallen stelselmatig te lang waren, wanneer 100% van de correcties geïnitieerd zou worden door de externe lus. De gemiddelde fout/interruptie-intervallen werden adequaat gesimuleerd, wanneer 50% of 67% van de detecties door de externe lus gedaan zouden worden. We hebben dus inderdaad een intern monitoringkanaal nodig om de verdeling van fout/interruptie-intervallen na te bootsen.

De andere gevarieerde parameters waren de duur van het interruptieproces en de hoeveelheid ruis die aan elk interval werd toegevoegd. We vonden dat alleen bij een schatting van de duur van het interruptieproces die overeenkomt met schattingen uit de literatuur (150 ms -200 ms) de gesimuleerde verdeling goed overeenkwam met de waargenomen verdeling. Dit gold voor een breed bereik van de ruisparameter. In Figuur 4 tonen we de verdeling van de waargenomen fout/interruptie intervallen en een verdeling die representatief is voor de simulaties. De figuur geeft de frequentie aan waarmee deze intervallen tussen bepaalde waardes vielen.

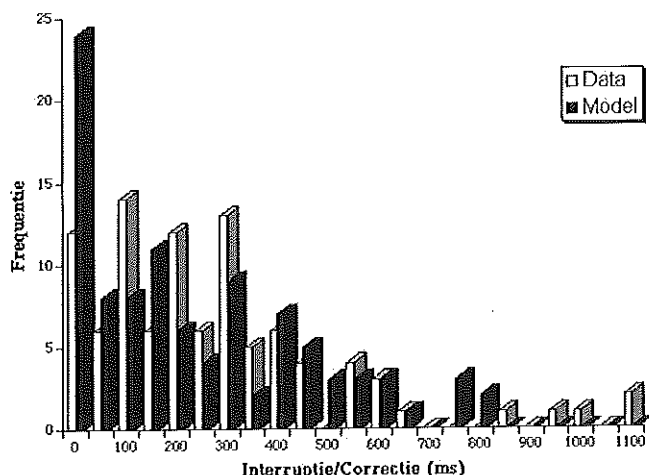


Figuur 4. Empirische en Gesimuleerde verdeling van fout/interruptie-intervallen.

Figuur 4 laat zien dat de empirische en gesimuleerde verdelingen sterk overeenkomen. De meeste intervallen, in zowel data als simulatie, vallen in het gebied tussen de 100 en 350 ms, met een lange "staart" tot 850 ms. Een goodness-of-fit test toonde aan dat er voor een groot deel van de parameter ruimte geen significante verschillen waren tussen model en data. We kunnen dus concluderen dat het model, met de

interne lus, de empirische verdeling van fout/interruptie-intervallen kan nabootsen.

In een tweede serie simulaties werd nagegaan of het model de verdeling van interruptie/correctie-intervallen kan simuleren. De duur van deze intervallen wordt mede bepaald door de gewijzigde hoofdinterruptieregel. Een typische modelverdeling en een empirische verdeling worden gerapporteerd in Figuur 5.

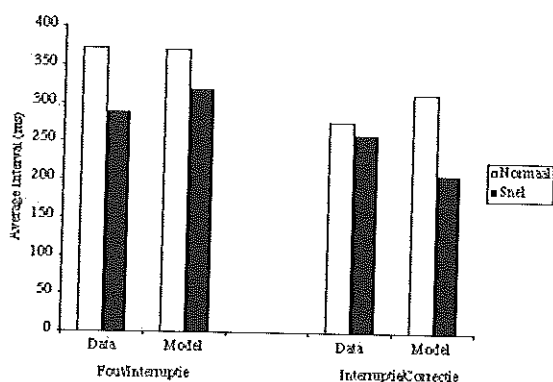


Figuur 5. Empirische en Gesimuleerde verdeling van interruptie/correctie-intervallen.

Figuur 5 laat zien dat de empirische verdeling en de modelverdeling sterk overeenkomen: De verdeling van interruptie/correctie-intervallen is scheef, met veel intervallen tussen de 0 en 200 ms. In Figuur 5 lijkt het model teveel intervallen van 0 ms te voorspellen, maar deze visuele impressie wordt niet bevestigd door statistische analyse: Model en data verschillen niet significant van elkaar. Het model, met de gewijzigde interruptieregel, blijkt dus ook in staat om de verdeling van interruptie/correctie-intervallen te simuleren.

In een derde serie simulaties werden effecten van spreesnelheid onderzocht. Oomen en Postma (in druk) vonden dat zowel fout/interruptie- als interruptie/correctie-intervallen korter werden naarmate spraak werd versneld. Voorspelt het model deze effecten? Het antwoord is bevestigend, gegeven twee aannames. Deze aannames betreffen de vraag wat er precies gebeurt met de processen die betrokken zijn bij taalperceptie, taalproductie en monitoring wanneer iemand sneller gaat spreken. De eerste aanname volgt uit de gewijzigde hoofdinterruptieregel. Die regel stelt dat interruptie een actie is die direct op articulatorische verwerking aangrijpt. Daaruit kan men afleiden dat dit proces niet gevoelig is voor manipulaties op niveaus die eerder dan de articulatie gepland worden. Een van die manipulaties betreft spreesnelheid. Daarom zal in snelle spraak de duur van het interrumperen constant blijven, maar zal de duur van het plannen van de correctie afnemen. Dit leidt tot een verkorting van de tijdsduur tussen interruptie en correctie.

De tweede aanname is dat bij het sneller spreken, alle taalproductie- en taalperceptieprocessen sneller verlopen. Dat articulatie sneller gaat in snelle spraak spreekt voor zich. Maar er zijn ook argumenten voor de stelling dat fonologische planning sneller gaat (zie bijvoorbeeld Dell, 1986). Kan taalperceptie dergelijke versnellingen bijhouden? Er is veel literatuur over de perceptie van "gecomprimeerde spraak", spraak die artificieel versneld is (bijv. Dupoux & Green, 1997). Uit die literatuur blijkt dat men spraak kan verwerken met tempi die vele malen hoger zijn dan men spraak kan produceren. Het begripssysteem kan dus grote versnellingen bijhouden in externe spraak. Het is daarom aannemelijk dat zelfgeproduceerde spraak ook sneller verwerkt wordt. Beide aannames zijn geïmplementeerd in ons model. De resultaten van simulaties waarin de snelheid gemanipuleerd wordt, zijn weergegeven in Figuur 6.



Figuur 6. Effecten van spreek snelheid op gemiddelde fout/interruptie- en interruptie/correctie-intervallen in model en data.

Zoals blijkt uit Figuur 6 voorspelt het model inderdaad een verkorting van zowel fout/interruptie-intervallen als interruptie/correctie-intervallen in snelle spraak.

Discussie

In dit artikel hebben we een geformaliseerde en uitgebreide versie van de perceptuele lus theorie gerapporteerd. Onze versie wijkt af van eerdere voorstellen met betrekking tot de relatie tussen interruptie en correctie. Wij beschouwen interruptie als een proces dat de taalplanning overslaat maar direct op articulatie aangrijpt. Dit impliceert dat interrumpen en het plannen van de correctie twee onafhankelijke, parallelle processen zijn. Daarom kan met het plannen van de correctie begonnen worden voordat de interruptie geëffectueerd is. Deze wijziging stelt het model in staat om zeer korte interruptie/correctie-intervallen te produceren, net zoals die in de werkelijkheid gevonden worden. Bovendien kan men uit de gewijzigde interruptieregel afleiden dat het interruptieproces niet gevoelig is voor talige variabelen zoals spreek-snelheid. Een andere nieuwe bijdrage is een algoritme dat op basis van empirische

gegevens correcties verdeelt in incidenten die geïnitieerd worden door de interne en door de externe monitoringlus.

De simulaties lieten zien dat het model er in slaagt om de empirische verdelingen van fout/interruptie-intervallen en interruptie/correctie-intervallen na te bootsen. Een voorwaarde voor het simuleren van de fout/interruptie-intervallen was dat de interne lus bijdroeg aan de foutdetectie. Ons decompositie-algorithme voorspelt dat het externe kanaal verantwoordelijk is voor ongeveer 2/3 van de correcties en het interne kanaal voor 1/3 van de correcties. Indien we die verhouding hanteren slaagt het model er inderdaad in om de verdeling van fout/interruptie-intervallen na te bootsen. Dit resultaat leidt tot een bevestigend antwoord op de eerste twee vragen die wij aan de orde stelden. Ten eerste is het nodig een intern monitoringkanaal te veronderstellen om de experimentele data te simuleren. Ten tweede is een intern monitoringkanaal dat gebruik maakt van het taalperceptiesysteem snel genoeg om deze data te simuleren.

Het model simuleert ook de verdeling van interruptie/correctie-intervallen adequaat. In de empirische dataset waren er veel interruptie/correctie-intervallen van 0 ms. Dit duidt erop dat de oorspronkelijke hoofdinterruptieregel niet houdbaar is (want volgens die regel weerspiegelt het interruptie/correctie-interval de duur van het plannen van de correctie). In ons model zijn het plannen van de correctie en het interromperen van de spraak twee parallelle processen. De succesvolle simulatie van interruptie/correctie-intervallen ondersteunt onze aangepaste interruptieregel.

Tenslotte lukte het ook de effecten van spreek snelheid op fout/interruptie- en interruptie/correctie-intervallen na te bootsen (namelijk een verkorting van de duur van beide intervallen in snelle spraak). Deze effecten volgen uit onze modelaannames en uit een extra aanname: de aanname dat bij snelle spraak alle taalverwerkingsprocessen sneller verlopen, inclusief taalbegrip.

Het model, een geformaliseerde en uitgebreide versie van de perceptuele lus theorie blijkt dus succesvolle voorspellingen te doen over de timing van interruptie en correctie. Uiteraard is hiermee het laatste woord nog niet gezegd over monitoring. In de toekomst zullen we het model verder moeten verfijnen. Dergelijke verfijningen betreffen onder andere de aard van het matchingproces: Welke geproduceerde representatie wordt vergeleken met welke doelrepresentatie? Dit matchingproces is vooralsnog vrijwel onuitgewerkt.

In de inleiding suggereerden we al dat een aantal taal- of spraakstoornissen verklaard kan worden als gevolg van een dysfunctionerende monitor of als gevolg van interacties tussen monitoring en stoornissen in de planning van taal. Een belangrijke vraag voor de toekomst is of met behulp van het model dergelijke stoornissen beter begrepen kunnen worden. Een mogelijkheid is bijvoorbeeld om het model toe te passen op niet-vloeiendheden in spraak van stotteraars. Het model kan voorspellingen doen over het moment van interruptie. Postma en Kolk (1993) suggereerden dat men uit het moment (in een syllabe) wanneer de interruptie geëffectueerd is, de vorm van de niet-vloeiendheid kan herleiden. Vindt de interruptie bijvoorbeeld plaats voordat begonnen is met articulatie, dan zal de niet-vloeiendheid de vorm hebben van een pauze of een blokkade. Valt de interruptie echter na het eerste foneem, dan zal het gevolg eerder een repetitie of prolongatie zijn. Het model kan dus mogelijk voorspellingen doen over de verdeling van diverse niet-vloeiendheden.

Een andere mogelijkheid is om het model toe te passen op monitorgedrag bij afasie. Oomen, Postma, en Kolk (dit nummer) onderzochten bijvoorbeeld de proporties zelf-gedetecteerde fouten in de spraak van patiënten met een afasie van Broca en van controleproefpersonen, in twee condities: normale auditieve feedback en gemaskeerde feedback door middel van luide ruis. In de ruisconditie kon men de eigen spraak niet horen. In die conditie stond de proefpersoon dus enkel het interne monitoringkanaal ter beschikking. Controleproefpersonen detecteerden significant minder fouten in de ruisconditie; patiënten met een afasie van Broca detecteerden echter evenveel fouten in beide condities. Dit impliceert dat de patiënten met de afasie van Broca gebruik maken van het interne monitoringkanaal, maar niet van het externe monitoring-kanaal (zie ook Schlenk, Huber, & Willmes, 1987). In termen van ons model betekent dit een drastische verschuiving in de partitionering in fouten die door de interne en door de externe lus gedetecteerd en gecorrigeerd worden. Toekomstig onderzoek zal moeten uitwijzen of de modelvoorspellingen in die situatie overeenstemmen met de afasiedata.

Noten

- 1 Nu verbonden aan the University of Edinburgh, Dept. of Psychology, 7 George Square, EH89JZ, Edinburgh, Scotland.
- 2 Waar 'hij' staat dient men 'hij/zij' te lezen, waar 'zijn' staat dient men 'zijn/haar' te lezen.

A computer model of verbal self-monitoring Summary

Verbal self-monitoring is the set of cognitive processes with which the speaker checks his or her speech for accuracy. According to a theory proposed by Levelt (1983) the language perception system is involved in this task. This system would check both our overt speech and our so-called inner speech. We have formalized this theory and implemented it as a computer model. We have simulated empirically obtained distributions of the intervals between the moment of error onset, interruption, and error repair, as well as the effects of speech rate on these intervals. The model was successful in simulating these intervals, given a number of assumptions about the mechanisms of monitoring. First, interrupting and repair planning should be conceived of as two parallel and independent processes. Second, in fast speech there is a speed-up of both speech production and speech perception.

Referenties

- Baars, B. J., Motley, M. T., & MacKay, D. G. (1975). Output editing for lexical status in artificially elicited slips of the tongue. *Journal of Verbal Learning and Verbal Behavior*, 14, 382-391.

- Bernstein, P. S., Scheffers, M. K., & Coles, M. G. H. (1995). "Where did I go wrong?" A psychophysiological analysis of error detection. *Journal of Experimental Psychology: Human Perception and Performance*, 21, 1312-1322.
- Blackmer, E. R., & Mitton, E. R. (1991). Theories of monitoring and the timing of repairs in spontaneous speech. *Cognition*, 39, 173-194.
- Dell, G. S. (1986). A spreading-activation theory of retrieval in sentence production. *Psychological Review*, 93, 283-321.
- Dell, G. S., & Repka, R. J. (1992). Errors in inner speech. In B. J. Baars (Ed.), *Experimental slips and human error: Exploring the architecture of volition*. New York: Plenum Press.
- Dupoux, E., Green, K. (1997). Perceptual adjustment to highly compressed speech: Effects of talker and rate changes. *Journal of Experimental Psychology: Human Perception and Performance*, 23, 914 - 927.
- Frith, C. D. & Done, D. J. (1988). Towards a neuropsychology of schizophrenia. *British Journal of Psychiatry*, 153, 437-443.
- Hartsuiker, R. J., & Kolk, H. H. J. (in druk). Error monitoring in speech production: A computational test of the perceptual loop theory. *Cognitive Psychology*.
- Lackner, J. R., & Tuller, B. H. (1979). Role of efference monitoring in the detection of self-produced speech errors. In: W. E. Cooper & E. C. T. Walker (Eds.), *Sentence Processing* (pp. 281-294).
- Levelt, W. J. M. (1983). Monitoring and self-repair in speech. *Cognition*, 14, 41 - 104.
- Levelt, W. J. M. (1989). *Speaking: From intention to articulation*. Cambridge, MA: MIT Press.
- Marshall, J., Robson, J., Pring, T., & Chiat, S. (1998). Why does monitoring fail in jargon aphasia? Comprehension, judgment, and therapy evidence. *Brain and Language*, 63, 79-107.
- McGuire, P. K., Silbersweig, D. A., & Frith, C. D. (1996). Functional neuroanatomy of verbal self-monitoring. *Brain*, 119, 907-917.
- Motley, M. T., Camden, C. T., & Baars, B. J. (1982). Covert formulation and editing of anomalies in speech production: Evidence from experimentally elicited slips of the tongue. *Journal of Verbal Learning and Verbal Behavior*, 21, 578-594.
- Nooteboom, S. G. (1980). Speaking and unspeaking: Detection and correction of phonological and lexical errors in spontaneous speech. In: V. A. Fromkin (Ed.), *Errors in linguistic performance: slips of the tongue, ear, pen, and hand*. New York: Academic Press.
- Oomen, C. E., & Postma, A. (in druk). Effects of increased speech rate on monitoring and self-repair. *Journal of Psycholinguistic Research*.
- Oomen, C. E., Postma, A., & Kolk, H. H. J. (dit nummer). Spraakmonitoring bij twee patiënten met afasie van Broca.
- Postma, A., & Kolk, H. H. J. (1992). The effects of noise masking and required accuracy on speech errors, disfluencies, and self-repairs. *Journal of Speech and Hearing Research*, 35, 537-544.
- Postma, A., & Kolk, H. H. J. (1993). The covert repair hypothesis: Prearticulatory repair processes in normal and stuttered disfluencies. *Journal of Speech and Hearing Research*, 36, 472-487.
- Postma, A., & Noordanus, C. (1996). Production and detection of speech errors in silent, mouthed, noise-masked, and normal auditory feedback speech. *Language and Speech*, 39, 375-392.
- Schlenk, K.-J., Huber, W., & Willmes, K. (1987). "Prepairs" and "repairs": Different monitoring functions in aphasic language production. *Brain and Language*, 30, 226-244.