

Faculteit der Letteren
Taalwetenschap

Studiejaar 2016 - 2017
Datum: 18/02/2017

Is het bestaan van een “woordspelingmodus” aannemelijk?

Een onderzoek naar de rol van de offset bij
het maken van woordspelingen

Hay Driessen
S4249186

Bachelorscriptie

Primaire begeleider: M. Broersma
Secundaire begeleider: R. van Hout

Radboud Universiteit



Inhoudsopgave

1. Inleiding	4
1.1 Onderzoek naar woordspelingen in het Japans	4
1.1.1 Soorten woordspelingen	5
1.2 Taalverwerking en productie.....	6
1.3 Woordactivatie.....	7
1.3.1 Woordactivatie op fonologisch niveau	9
1.4 De rol van de offset.....	10
1.5 Analyseren van woordspelingen	12
1.6 Onderzoeksvragen en hypothesen	13
2. Methoden	14
2.1 Geannoteerde Eigenschappen	15
2.2 Levenshtein distance	16
2.3 Analyses.....	17
3. Resultaten.....	18
3.1 Hypothese 1: Verschil in frequentie tussen offsetoverlap en onsetoverlap.....	18
3.2 Hypothese 2: Verschil in gemiddelde Levenshtein distance tussen offsetoverlap en onsetoverlap	18
3.3 Hypothese 3: Verschil in Levenshtein distance tussen inbeddingen en mutaties	19
4. Discussie	22
4.1 Discussie van deelvragen (hypothese 1, 2 en 3)	23
4.2 Algemene discussie	25
5. Conclusie	27
Referenties.....	28
Bijlagen	31
Bijlage 1: Gebruikte inbeddingen en mutaties voor dit onderzoek	31
Bijlage 2: Lijst van alle eigenschappen waarop de woordspelingen geannoteerd zijn (Kerkhof, 2015)	42
Bijlage 3: Frequentie van type woordspeling Otake en Cutler (2013) en dit onderzoek.....	45

1. Inleiding

Er is al veel bekend over taalverwerking en woordactivatie onder normale omstandigheden. Er is echter nog geen onderzoek uitgevoerd naar een eventuele verandering van de verwerking en activatie bij het bewust maken van woordspelingen. Dit onderzoek tracht indicaties te geven over een eventuele “woordspelingsmodus”. Deze studie suggereert dat in deze woordspelingsmodus de verwerking en activatie van woorden door het bewust maken van woordspelingen, veranderen ten opzichte van een normale situatie (waarin de intentie tot het maken van woordspelingen ontbreekt). Om een woordspelingsmodus te indiceren, is de rol die de offset van het bronwoord speelt in de totstandkoming van het doelwoord bij woordspelingen zoals inbeddingen en mutaties onderzocht. Deze rol is onderzocht door middel van een corpusstudie van woordspelingen van Gerard Ekdorf. Door de rol die de offset van het bronwoord speelt in de totstandkoming van het doelwoord te onderzoeken bij woordspelingen, kan bepaald worden hoe belangrijk deze rol is. In een normale situatie blijkt dat de onset een groter effect heeft op het activeren van concurrenten dan de offset (Huettig & McQueen, 2007; Huettig, Rommers & Meyer, 2011; Magnuson, Tanenhaus, Aslin, & Dahan (2003). Als blijkt dat bij het maken van woordspelingen de offset van het bronwoord vaker gebruikt wordt dan de onset, dan zou dit kunnen aantonen dat de offset van een woord toegankelijker is bij het maken van woordspelingen. Het sterker activeren en/of controleren van offsetconcurrenten tijdens het maken van woordspelingen kan een indicatie zijn voor het bestaan van een woordspelingsmodus. In de volgende paragraaf wordt ingegaan op eerder onderzoek naar woordspelingen en op verschillende soorten woordspelingen. Vervolgens wordt globaal toegelicht welke processen van belang zijn bij taalverwerking en –productie. Daarna wordt er literatuur over woordactivatie besproken. Tot slot wordt er een koppeling gemaakt tussen woordactivatie en woordspelingen en worden de hoofdvraag en de deelvragen van deze studie besproken.

1.1 Onderzoek naar woordspelingen in het Japans

Eerder onderzoek naar woordspelingen is onder andere uitgevoerd door Otake en Cutler (2013). Zij onderzochten verschillende soorten woordspelingen in het Japans van één Japanse spreker, genaamd Dokumumushi Sandayuu, over een periode van twee jaar. Deze Japanse spreker staat bekend om zijn woordspelingen die in zijn radioprogramma aan bod komen. In dit onderzoek suggereren zij dat een grotere overeenkomst tussen bronwoord en doelwoord leidt tot beter werkende woordspelingen. Deze suggestie is gebaseerd op resultaten uit eerder onderzoek (Cutler & Otake, 2002; Kawahara & Shinohara, 2009; Lagerquist, 1980; Otake, 2010; Otake & Cutler, 2001). De resultaten van het onderzoek van Otake en Cutler (2013) bevestigen deze suggestie. Het minimaal aanbrengen van veranderingen tussen het bronwoord en het doelwoord wordt aangeduid met het Principle of Maximal Source Preservation (MaxSP) (Otake & Cutler, 2013). Tevens concluderen Otake en Cutler (2013) dat homofonen de meest geslaagde woordspelingen zijn, omdat deze woordspelingen het meest voldoen aan MaxSP.

1.1.1 Soorten woordspelingen

In het onderzoek van Otake en Cutler (2013) zijn drie typen woordspelingen onderzocht: homofonen, inbeddingen en mutaties. Homofonen zijn woorden die hetzelfde klinken, maar een andere betekenis hebben. In (1) is een voorbeeld te zien van een woordspeling waarin gebruik gemaakt wordt van een homofoon. In (1) heeft het eerste voorkomen van *teken*, het bronwoord, een andere betekenis dan het tweede voorkomen (Kerkhof, 2015). In het eerste voorkomen wordt de parasiet bedoeld. Bij het tweede voorkomen van *teken*, het doelwoord, krijgt *teken* de betekenis van symbool of verschijnsel (Kerkhof, 2015). Het bronwoord en het doelwoord klinken hetzelfde, maar hebben een andere betekenis. In (2a), (2b) en (2c) worden er voorbeelden gegeven van woordspelingen waarin inbeddingen worden gebruikt. Bij inbeddingen overlapt het bronwoord gedeeltelijk of in zijn geheel met het doelwoord. Inbeddingen ontstaan door insertie of extractie. Bij insertie wordt een deel van het bronwoord in een ander woord of andere klanken geplakt. In (2a) is een voorbeeld van insertie weergegeven. Het bronwoord *helpen* komt in zijn geheel terug in het doelwoord *verhelpen*. Wanneer niet het gehele bronwoord terugkeert in het doelwoord maar slechts een deel van het bronwoord wordt er eveneens gesproken van insertie, zie (2b). De eerste lettergreep van het bronwoord *Knopfler*, is terug te zien in het doelwoord *knoflook*. Bij extractie bestaat het doelwoord uit een deel van het bronwoord, waarbij geen veranderingen worden aangebracht in het doelwoord. In (2c) wordt een voorbeeld gegeven van extractie. *Dracula* wordt gebruikt als bronwoord, het geselecteerde doelwoord wordt *draak*. *Draak* zit als geheel ingebed in *Dracula*. In (3) wordt een voorbeeld gegeven van een woordspeling waarin een mutatie wordt gebruikt. In (3) wordt het bronwoord *zangeres* omgevormd tot *zwangeres* om als doelwoord te fungeren. De *w* wordt als nieuwe klank ingevoegd in het oorspronkelijke bronwoord om zo een humoristisch effect te geven aan het doelwoord. Bij een mutatie worden één of meerdere klanken van het bronwoord veranderd om zo het doelwoord te vormen. Hoe meer klanken er echter veranderd worden, des te groter de kans dat de woordspeling zijn humoristisch effect verliest of hoe groter de kans dat de woordspeling niet door gesprekspartners herkend wordt (schending MaxSP) (Cutler & Otake, 2002; Kawahara & Shinohara, 2009; Lagerquist, 1980; Otake, 2010; Otake & Cutler, 2001; Otake & Cutler, 2013).

(1) Ze zijn er weer de **teken**. Is het **teken** van de maand

(2a) Nou dat komt mooi uit hè Bart, ga ik jou proberen te **helpen**. Tenminste **verhelpen** van het probleem.

(2b) Och, hoor dat toch eventjes, Marc **Knopfler**. Marc **Knoflook** hè, vanwege die geur.

(2c) Ja **dracula** uit sesamstraat, die vond ik vroeger altijd heel eng. Ja dat snap ik, maar die is het niet. Je zoekt een **draak**, dracula van een plaat eigenlijk.

- (3) De Amerikaans **zangeres**...nou zeg maar gerust **zwangeres**, die in verwachting is van haar tweede kind.

1.2 Taalverwerking en productie

Bij de verwerking van gesproken taal moet eerst het akoestische signaal waargenomen worden. Tijdens de verwerking van dit signaal worden er verschillende representaties van betekenissen gegenereerd. Vervolgens wordt de juiste betekenis geselecteerd. Fonologische patronen, het mentale lexicon en wereldkennis worden gebruikt om de betekenis te achterhalen van het akoestische signaal (Lieberman, Cooper, Shankweiler & Studdert-Kennedy, 1967). De hoorder begeeft zich tijdens een gesprek in een onzekere positie. De hoorder moet achterhalen wat de boodschap van de spreker is. De hoorder test op basis van eerdere ervaringen hypothesen over de componenten (fonemen, structuur, prosodie, betekenis, context) van de inkomende spraak, om zo de betekenis van een uiting te achterhalen. Al deze componenten worden zo snel mogelijk herkend. Volgens McQueen, Dahan en Cutler (2003) zorgt een continue stroom van informatie ervoor dat hypothesen over de inkomende spraak continu vergeleken, herbeoordeeld of verworpen worden. Daarnaast moet de hoorder nog over een eventuele respons nadenken.

Het hiervoor beschreven proces van taalverwerking is niet gelijk aan het proces van taalproductie. Bij productie probeert de spreker zijn boodschap over te dragen aan de hoorder. Dit doet hij/zij door de boodschap naar de gewenste articulatoire vorm te converteren, wat betekent dat de spreker ervoor zorgt dat de geselecteerde uiting op de juiste manier uitgesproken wordt. De eerste stap voor het produceren van de boodschap is conceptualisatie. De lexicale items die van belang voor de boodschap zijn, worden geselecteerd. De volgende stap is het construeren van de uiting, wat wordt gevolgd door het articuleren van de uiting. Na het articuleren vindt bewustwording bij de spreker plaats van hetgeen de spreker heeft gezegd. De spreker heeft vervolgens de mogelijkheid de uiting nog aan te passen of te corrigeren, ook wel zelfmonitoring genoemd (Levelt, Roelofs & Meyer, 1999).

In de meeste gevallen zijn personen die zich in een conversatie begeven zowel spreker als hoorder. De intentie van gesprekspartners is het zo duidelijk en efficiënt mogelijk maken van de boodschap en het selecteren van de interpretatie die het meest voor de hand ligt voor desbetreffende context (Grice, 1970). Deze intentie wordt echter anders wanneer mensen bewust woordspelingen gaan maken. De grappenmaker probeert naast de normale conversatie ook een humoristisch effect te geven aan de conversatie. Dit zorgt ervoor dat de woordspelingen die de grappenmaker maakt niet altijd passend zijn of niet het meest logisch zijn in desbetreffende context. In (4) vindt er een dialoog plaats tussen Gerard Ekdorf (G.E.), wie bekend staat om zijn woordspelingen, en Pascal. Pascal wil de titel en de artiest van een bepaald liedje achterhalen. Pascal geeft aan dat het liedje uit eind jaren tachtig, begin jaren negentig komt. G.E. kiest ervoor om te reageren met een woordspeling. Deze reactie van G.E. helpt Pascal niet om zijn doel te bereiken. G.E. kiest er niet voor om een reactie te geven die duidelijk maakt of het antwoord bekend of onbekend is. De intentie die gesprekspartners in het algemeen hebben, namelijk het zo duidelijk maken van de boodschap en het oplossen van een gemeenschappelijk doel, verandert in dit geval. G.E. heeft de intentie om een humoristisch effect te geven aan zijn uitingen.

(4)

(G.E.): Jij zoekt een instrumentaal liedje?

(P): Ja dat klopt, uit eind jaren tachtig, begin jaren negentig als ik het goed heb.

(G.E.): Dat is in ieder geval een taal die we allemaal spreken. De eerste mentaal.

Het proces van woordspelingen maken betreft zowel spraakherkenning als spraakproductie. De grappenmaker moet eerst herkennen wat de gesprekspartner zegt. De volgende stap is het selecteren van het bronwoord. Dit bronwoord dient als basis voor het maken van een woordspeling en wordt uit een eerdere uiting van de gesprekspartner of de spreker zelf gekozen. In (4) betreft het bronwoord *instrumentaal*. Tegelijkertijd moet in gedachte gehouden worden welk bronwoord en welke eventuele modificaties van dit bronwoord tot de meest geslaagde woordspeling leidt. Daarnaast moet de grappenmaker een indicatie hebben van de gedeelde kennis met zijn gesprekspartner, zodat hij/zij weet of de woordspeling herkend wordt. Als de grappenmaker een bepaalde stand van zaken bespreekt die niet bekend is bij de hoorder, dan bestaat er een mogelijkheid dat de hoorder de woordspeling niet kent. De grappenmaker selecteert een woord uit de gegeven conversatie waarmee hij/zij een woordspeling gaat maken. Dit gekozen woord wordt ook wel het bronwoord genoemd (Kerkhof, 2015). De intentie van de grappenmaker is om van dit woord een zo geslaagd mogelijke woordspeling tot stand te laten komen. Het bronwoord ondergaat eventuele aanpassingen en/of krijgt een andere betekenis in het uiteindelijk doelwoord waar de woordspeling tot stand komt. Het woord dat naar aanleiding van het bronwoord tot stand komt, wordt het doelwoord genoemd (Kerkhof, 2015).

Als de grappenmaker eenmaal bronwoord en doelwoord geselecteerd heeft, kan de woordspeling in een correcte uiting geïmplementeerd worden. De grappenmaker moet concurrenten geactiveerd houden die passend kunnen zijn voor een woordspeling, maar die niet per se het meest voor de hand liggend zijn in desbetreffende context. In deze studie wordt ervan uitgegaan dat het activeren van verschillende concurrenten van het bronwoord en het controleren van deze activatie een belangrijke rol speelt bij het maken van woordspelingen. Op basis van eigen intuïties en onderzoek van Otake en Cutler (2013) wordt aangenomen dat er tijdens het maken van woordspelingen een wisselwerking tussen het herkenningsproces en het productieproces plaatsvindt. Op basis van herkenning worden mogelijke doelwoorden geactiveerd en tijdens de productie wordt het daadwerkelijke doelwoord geselecteerd. Tijdens de productiefase kan een woord eventueel nog aangepast worden, omdat er bewustwording plaatsvindt dat een ander woord toch geschikter is voor de woordspeling.

1.3 Woordactivatie

In deze studie is woordactivatie gekoppeld aan het maken van woordspelingen. De grappenmaker moet bij het bronwoord een doelwoord selecteren. In deze studie is aangenomen dat de keuze van het doelwoord onder andere bepaald wordt door de concurrenten die geactiveerd worden door het bronwoord. Otake en Cutler (2013) beschrijven in hun studie dat deze activatie in enige mate te controleren is. Wanneer een woord in zijn geheel verwerkt is, verliezen concurrenten van dit woord die tijdens de verwerking van dit woord geactiveerd werden, activatie (geen intentie om

woordspelingen te maken). Door deze activatie van concurrenten te controleren, blijven concurrenten geactiveerd tijdens het maken van woordspelingen wanneer het woord in zijn geheel al verwerkt is. Ter illustratie: het woord *beker* wordt verwerkt tijdens het horen van een zin. Tijdens het verwerken van het woord *beker* kunnen zowel *bever* als *spreker* concurrenten vormen van *beker*. *Bever* vormt een onsetconcurrent van *beker* omdat beide dezelfde beginklank (onset) delen. *Spreker* vormt een offsetconcurrent van *beker* omdat beide dezelfde eindklanken (offset) delen. *Bever* verliest activatie op het moment dat *beker* en *bever* niet meer in klank overlappen. *Spreker* en *beker* kennen in het begin van het woord al een mismatch in klank. Hierdoor is de waarschijnlijkheid dat *spreker* de geschikte kandidaat zal zijn tijdens het verwerken van *beker*, erg klein. Beide kandidaten verliezen uiteindelijk activatie omdat ze niet overeenstemmen met *beker*. In deze studie wordt echter aangenomen dat kandidaten zoals *spreker* en *bever* geactiveerd blijven wanneer *beker* in zijn geheel verwerkt is en er gekozen is om *beker* als bronwoord te laten fungeren. *Bever* en *spreker* blijven beide potentiële kandidaten voor het doelwoord bij het maken van een woordspeling waardoor deze woorden geactiveerd blijven.

Het eerste onderzoek dat bewijs vond voor activatie van gerelateerde objecten/concepten tijdens het verwerken van een woord, was het onderzoek van Cooper (1974). De oogbewegingen van participanten werden gemeten bij visuele stimuli als respons op gesproken taal. Participanten luisterden naar gesproken zinnen en kregen tegelijkertijd plaatjes te zien. Het is gebruikelijk om de plaatjes al te laten zien voordat de gesproken taal aangeboden wordt, zodat linguïstische en niet linguïstische informatie geactiveerd worden. Tijdens de aanbidding van deze stimuli worden de oogbewegingen (fixaties) gemeten. Wanneer bijvoorbeeld de zin “My scatterbrained dog Scotty...” aangeboden wordt en er tegelijkertijd vier plaatjes te zien zijn, waarvan een van de plaatjes een hond betreft, zijn participanten geneigd om naar het plaatje van de hond te kijken. Ook bleken fixaties gericht te zijn op een gerelateerd object. Bij het horen van de zin: “During a photographic safari...”, waren fixaties gericht op het plaatje van een camera. Fotografie en camera zijn gerelateerd aan elkaar omdat met een camera gefotografeerd kan worden. Objecten die overeenkomen of geassocieerd worden met een deel van de aangeboden spraak, krijgen meer aandacht van participanten (Huettig, Rommers & Meyer, 2011). Het gebruik van eye-tracking en het registreren van oogfixaties bij auditieve input waarbij tegelijkertijd plaatjes te zien zijn, staat nu bekend als het visual world paradigm (Allopenna, Magnuson, & Tanenhaus, 1998).

Recenter onderzoek naar woordactivatie is onder andere uitgevoerd door Huettig en McQueen (2007). Zij hebben met behulp van het visual world paradigm aangetoond dat fonologische concurrenten, semantische concurrenten en visuele vormconcurrenten geactiveerd worden tijdens de verwerking van een doelwoord. Bij de experimentele trials in het onderzoek werden vier plaatjes getoond, waarvan drie plaatjes een fonologische, semantische of visuele gelijkenis vertoonden met het doelwoord uit de zin. Het resterende plaatje vertoonde geen overeenkomst met het doelwoord. Bij de zin *Uiteindelijk keek ze naar de beker die voor haar stond*, hoorden de volgende vier plaatjes: *bever*, *vork*, *klos* en *paraplu*. *Bever* vormt de fonologische concurrent van het doelwoord *beker* uit de zin, omdat dezelfde onset gedeeld wordt. De *vork* vertoont een semantische overeenkomst met *beker* omdat beide objecten bij een maaltijd gebruikt kunnen worden. De *klos* heeft dezelfde vorm als de *beker*. De *paraplu* vertoont geen overeenkomst in visuele vorm en vertoont geen fonologische en/of semantische overeenkomsten met *beker*. Concurrenten werden zo gekozen dat er maar op één gebied (semantisch, fonologisch of visuele vorm) sprake was van een overeenkomst. Een concurrent mocht bijvoorbeeld niet semantische én

fonologische overeenkomsten vertonen met het doelwoord. Bij het horen van de zin *Uiteindelijk keek ze naar de beker die voor haar stond*, werden de fixaties sterker bij *bever* naarmate het woord *beker* uitgesproken werd. Op het moment dat *beker* en *bever* niet meer overlapt in klank, namen het aantal fixaties af bij het plaatje van *bever* en namen de fixaties bij de plaatjes van *vork* (semantische concurrent) en *klos* (visuele concurrent) toe. Dit toont aan dat fonologische onsetconcurrenten geactiveerd worden wanneer het doelwoord nog niet in zijn geheel verwerkt is en dat andere semantisch- en vormgerelateerde woorden geactiveerd worden wanneer de luisteraar weet dat er geen match is tussen bijvoorbeeld *beker* en *bever*.

1.3.1 Woordactivatie op fonologisch niveau

De activatie van fonologische concurrenten wordt in deze studie als een erg belangrijke factor gezien voor het maken van woordspelingen. Geslaagde woordspelingen ondergaan zo min mogelijk veranderingen en fonologische concurrenten leiden in de meeste gevallen tot een grotere overlap tussen bronwoord en doelwoord dan semantische en visuele concurrenten (Otake & Cutler, 2013). Daarom spelen fonologische concurrenten een belangrijke rol bij woordspelingen.

Uit onderzoek van Allopenna et al. (1998) bleek dat offsetconcurrenten ook geactiveerd werden. Woorden zoals *speaker* werden geactiveerd bij het horen van het woord *beaker*. Echter bleek dat onsetconcurrenten sterker concurreerden dan offsetconcurrenten. Zodra de akoestische informatie van *beaker* niet meer overeenkwam met *beetle* en luisteraars een fonologische mismatch tussen *beaker* en *beetle* herkenden, begonnen oogfixaties van *beetle* af te nemen en begon de waarschijnlijkheid van *beaker* toe te nemen. Fixaties naar *speaker* namen toe naarmate het einde van *beaker* verwerkt werd. Huettig, Rommers & Meyer (2011) Magnuson, Tanenhaus, Aslin, & Dahan (2003) beweren dat bij normale spraakherkenning akoestische informatie aan het begin van gesproken woorden belangrijker is dan akoestische informatie later in het woord.

In normale spraakherkenning (wanneer er geen intentie is om een woordspeling te maken/ humoristisch effect te geven aan een uiting) worden onsetconcurrenten het sterkst geactiveerd bij het horen van een bepaald targetwoord (bronwoord). Er zou verwacht kunnen worden dat bij het maken van woordspelingen het selecteren van onsetconcurrenten voor het doelwoord makkelijker is dan het selecteren van offsetconcurrenten. Onsetconcurrenten worden geactiveerd zodra de verwerking van een woord begint (Huettig & McQueen, 2007). Het zou voor de grappenmaker dan makkelijker zijn om een van deze onsetconcurrenten geactiveerd te houden om zo een van deze concurrenten te selecteren voor het doelwoord. Als offsetconcurrenten en andere soort concurrenten (visueel en/of semantisch) naast de al geactiveerde onsetconcurrenten ook geactiveerd blijven, heeft de grappenmaker veel opties voor het selecteren van een eventueel doelwoord. Hierdoor neemt het selecteren van een doelwoord te veel tijd in beslag. Dit kan resulteren in het verlies van het humoristisch effect van de woordspeling, omdat het doelwoord te laat tot uiting wordt gebracht in de conversatie. Het geactiveerd houden van concurrenten die veel overlap hebben, zorgt voor een inhiberend effect op de herkenning van een woord (Goldinger, Luce & Pisoni 1989; Luce 1986; Luce & Pisoni 1998; Vitevitch and Luce 1998, 1999). De meest succesvolle woordspelingen overlappen juist zoveel mogelijk. Het zou dan voor een groter inhiberend effect zorgen als zowel onset- als offsetconcurrenten geactiveerd blijven.

Toch blijken onsetconcurrenten geen prominentere rol te spelen bij het maken van woordspelingen. Volgens Cutler & Otake (2013) wordt bij het bronwoord *candle*, *handle* eerder

gekozen als doelwoord voor een woordspeling dan *candy*. Dit geeft aan dat het soort concurrenten dat geactiveerd en geselecteerd wordt, situatie-afhankelijk kunnen zijn. Onsetconcurrenten worden niet in elke situatie sterker geactiveerd.

Activatie van inbeddingen

Inbeddingen worden geactiveerd wanneer de onset van de syllabe van een targetwoord (bronwoord) overeenkomt met de onset van een syllabe van een eventuele concurrent (als deze onset kandidaten heeft die dezelfde onset delen). Uit onderzoek van Vroomen & Gelder (1997) bleek dat het lexicale item *boos* geactiveerd werd bij het horen van *framboos*. *Boos* vormt de tweede syllabe van *framboos*. Lexicale representaties van inbeddingen die niet in de onset van de eerste syllabe overeenkomen worden dus eveneens geactiveerd. Dat de tweede syllabe van een woord concurrenten activeert, komt overeen met de bevindingen van Norris, McQueen en Cutler (1995) en Vroomen en De Gelder (1995a, 1995b). De activatie van inbeddingen verklaart het verschijnen van inbeddingen bij woordspelingen. *Boos* ontstaat door extractie, en kent overlap met de offset van *framboos*. Doordat *boos* een bestaand woord is en *fram* niet, is het makkelijker om een woordspeling te maken met *boos* dan met *fram*. Het type overlap bij inbeddingen is echter afhankelijk van het woord dat geselecteerd wordt en van de opbouw van het woord. Het eerste deel van een woord kan een bestaand woord zijn en het tweede deel niet en vice versa. Bij *framboos* en *boos* is er overlap in de offset, bij *beeldende* en *beeld* is er overlap in de onset. Door de bevindingen van Otake en Cutler (2013) te koppelen aan de activatie van inbeddingen, wordt verwacht dat inbeddingen vaker zullen resulteren in offsetoverlap dan in onsetoverlap.

Op basis van de eerdere beschreven literatuur wordt in deze studie gesuggereerd dat offsetoverlap tussen bronwoord en doelwoord frequenter plaats zal vinden dan onsetoverlap bij het maken van woordspelingen (Otake & Cutler, 2013; Norris, McQueen & Cutler, 1995, Vroomen & de Gelder, 1995a, 1995b).

1.4 De rol van de offset

In dit onderzoek is onderzocht hoe belangrijk de rol is die de offset van het bronwoord speelt in de totstandkoming van het doelwoord bij het maken van woordspelingen zoals mutaties en inbeddingen. Op basis van woordspelingen gemaakt door Gerard Ekdorf, zijn inbeddingen en mutaties onderzocht. Bij deze woordspelingen is onderzocht wat het verschil was in frequentie tussen de volgende drie groepen: onsetoverlap, offsetoverlap en zowel onset- als offsetoverlap tussen het bronwoord en het doelwoord. De categorie zowel onset- als offsetoverlap is onderzocht omdat de omvang van deze categorie groter bleek te zijn dan voorheen verwacht werd. De nadruk in deze studie lag echter bij het vergelijken van onsetoverlap met offsetoverlap. Daardoor waren er geen specifieke verwachtingen over de categorie zowel onset- als offsetoverlap. Ook werd onderzocht wat de hoeveelheid aan overlap was van de drie overlapcategorieën. Met behulp van de Levenshtein distance is berekend wat de orthografische overeenkomst was tussen bronwoord en doelwoord (Levenshtein, 1966). Met de Levenshtein distance wordt berekend hoe groot de overeenkomst in letters is tussen bronwoord en doelwoord; het is een valide maat om uitingen met een hoge orthografische overlap te onderscheiden van uitingen met een lage orthografische overlap

(Schepens, Dijkstra, & Grootjen, 2012). The Levenshtein distance telt het minimale aantal substituties, inserties en deleties die nodig zijn om een string te veranderen in een andere. De formule van de Levenshtein distance wordt in de methodensectie besproken.

Recente studies hebben succesvol gebruik gemaakt van de Levenshtein distance om orthografische gelijkenis tussen woorden te berekenen (Heeringa, 2004; Kondrak & Sherif, 2006; Yarkoni et al., 2008). Met behulp van de Levenshtein distance kon in deze studie berekend worden welke van de drie overlapcategorieën het hoogste gemiddelde aan proportionele overlap had.

Op basis van bevindingen van Otake en Cutler (2013) was de verwachting dat de offset frequenter zou overlappen tussen bronwoord en doelwoord dan de onset. Voor deze typen overlap, onset of offset, werd verwacht dat er geen groot verschil zou zijn qua gemiddelde Levenshtein distance. De score van de Levenshtein distance wordt bepaald door welk doelwoord bij het bronwoord geselecteerd wordt door de grappenmaker en in hoeverre dit doelwoord in aantal letters verschilt ten opzichte van het bronwoord. Wanneer de grappenmaker bijvoorbeeld het bronwoord *pak* selecteert en daarbij het doelwoord *pad* of *mak* selecteert, verandert de score van de Levenshtein distance niet, terwijl *mak* geclassificeerd zou worden als offsetoverlap en *pad* als onsetoverlap. Er werd wel verwacht dat het type woordspeling een invloed zou hebben op de Levenshtein distance, waarbij werd verwacht dat mutaties een hoger gemiddelde Levenshtein distance zullen hebben. Wanneer het bronwoord ingebed wordt in het doelwoord, is de kans groot dat dit doelwoord bestaat uit een of meerdere lettergrepen of zelfs een heel woord dat niet overlapt met het bronwoord. Hierdoor is het verschil in letters tussen bronwoord en doelwoord bij inbeddingen al snel groter dan bij mutaties, waardoor de gemiddelde Levenshtein distance bij mutaties groter zal zijn. Ter illustratie: een inbedding waarvan het bronwoord *instrumentaal* betreft met als doelwoord *mentaal* en een mutatie waarbij het bronwoord *Donners* betreft met als doelwoord *Donders*. Wanneer de inbedding en mutatie vergeleken worden, is er bij de inbedding een verschil van zes letters en bij de mutatie een verschil van één letter. De overeenkomst in letters is dan bij de mutaties groter dan de overeenkomst in letters bij inbeddingen. De grotere overeenkomst bij de mutatie resulteert daarom in een grotere Levenshtein distance.

Deze verwachtingen zijn gebaseerd op het onderzoek van Otake en Cutler (2013) maar ook op eigen intuïties. Er is nog geen onderzoek gedaan naar de activatie van concurrenten bij woordspelingen. De verwachting dat de offset van het bronwoord en het doelwoord vaker overeenkomen dan de onset van het bronwoord en het doelwoord, is gebaseerd op de bevindingen van Otake en Cutler (2013). Normaal gesproken zorgt de onset voor sterkere activatie van concurrenten dan de offset (Huettig & McQueen, 2007; Huettig, Rommers & Meyer, 2011; Magnuson, Tanenhaus, Aslin, & Dahan (2003). Maar zoals Otake en Cutler (2013) aangaven, wordt bij het bronwoord *candle*, *handle* eerder gekozen als doelwoord dan bijvoorbeeld *candy*. Dit laat zien dat onsetconcurrenten niet altijd sterker concurreren om activatie dan offsetconcurrenten. De concurrenten die normaal gesproken activatie verliezen, omdat ze niet overeenkomen met datgene wat gezegd wordt, zullen langer geactiveerd blijven omdat het mogelijke kandidaten zijn voor het selecteren van het doelwoord in de woordspeling. In combinatie met het onderzoek van Kawahara en Shinohara, (2009) en het onderzoek van Lagerquist (1980) wordt gedacht dat fonologische en/of articulatorische gelijkheid tussen bronwoord en doelwoord de keuze kan beïnvloeden. Een woord dat meer articulatorische en fonologische gelijkheid vertoont met het bronwoord, is makkelijker te produceren. Op basis van de bevindingen van Otake en Cutler (2013), Kawahara en Shinohara

(2009) en Lagerquist (1980) en intuïties is de aanname in deze studie dat het activatieproces van concurrenten anders verloopt bij het maken van woordspelingen en dat er een wisselwerking tussen spraakherkenning en -productie plaatsvindt bij het maken van woordspelingen. De intentie van een persoon verandert ten opzichte van “normale” spraakherkenning. Door deze verandering van intentie is de spreker in staat om woordactivatie deels te controleren. Otake en Cutler (2013) geven aan dat er criteria zijn voor het afwijzen of behouden van concurrenten, dat geactiveerde concurrenten in enige mate te controleren zijn door de spreker en dat de tolerantie (beschikbaar houden van bepaalde concurrenten) van deze activatie situatie-afhankelijk is. De controle van geactiveerde concurrenten vindt plaats door het behouden of afwijzen van mogelijke kandidaten tijdens het horen van een bronwoord. De tolerantie van het behouden of afwijzen van mogelijke kandidaten tijdens het horen van het bronwoord komt hoger te liggen bij het maken van woordspelingen. Doordat de tolerantie voor mogelijke kandidaten bij de inkomende spraak tijdens het maken van woordspelingen hoger ligt, blijven meerdere kandidaten geactiveerd. Hierdoor ontstaat er een asymmetrische verhouding in activatie tussen onsetconcurrenten en offsetconcurrenten. Tijdens normale spraakherkenning worden onsetconcurrenten sterker geactiveerd. Tijdens het maken van woordspelingen zullen offsetconcurrenten sterker geactiveerd worden. Het zou zo kunnen zijn dat het behouden van offsetconcurrenten een voordeel biedt bij het herkennen van een woordspeling. Deze aanname is gebaseerd op het onderzoek van Lupker en Colombo (1994), waarbij rijmprimen een positief effect hadden op woordherkenning tijdens lexicale decisietaken en benoemingstaken.

Door het aantal voorkomens van de offset te onderzoeken en deze te vergelijken met het aantal voorkomens in onset, kunnen indicaties gegeven worden dat het activatieproces bij het maken van woordspelingen anders is dan normaal gesproken, als blijkt dat de offset van het bronwoord vaker terugkomt in het doelwoord dan de onset. Ook het onderzoeken van de proportie aan overlap kan meer inzichten geven in de rol die offset en onset hebben bij het maken van woordspelingen. De intentie van de persoon verandert, waardoor er een soort woordspelingmodus geactiveerd wordt. Door het bewust bezig zijn met het maken van een geslaagde woordspeling, houdt de grappenaker meer concurrenten geactiveerd die normaal gesproken al activatie verloren zouden hebben. De offset van het bronwoord zal een belangrijkere rol gaan spelen in de activatie van concurrenten, omdat het verwerken van de offset in het doelwoord bij woordspelingen een voorkeur heeft (Otake & Cutler, 2013). De offset zal hierdoor frequenter terug te zien zijn in het doelwoord. Hierdoor kunnen er indicaties gegeven worden dat het activatieproces van concurrenten tijdens het maken van woordspelingen anders is dan bij normale spraakverwerking en dat het aannemelijk is dat sprekers in staat zijn om activatie van concurrenten in enige mate te controleren.

1.5 Analyseren van woordspelingen

Voor het analyseren van de woordspelingen zijn in dit onderzoek woordspelingen van Gerard Ekdorf (G.E.) gebruikt. G.E. staat bekend om zijn woordspelingen die hij gebruikt in zijn radio-uitzendingen van Effe Ekdorf. G.E. maakt woordspelingen van bronwoorden die door hemzelf en door anderen genoemd zijn. Zowel de woordspelingen waarbij het bronwoord genoemd wordt door G.E. als de woordspelingen waarbij de bronwoorden genoemd zijn door anderen, zijn geanalyseerd in dit onderzoek. In dit onderzoek zijn alleen de inbeddingen en mutaties gebruikt voor verdere

analysen. Door beide oorsprongen te analyseren, kon meer data in het corpus opgenomen worden en konden zoveel mogelijk inbeddingen en mutaties onderzocht worden. Otake en Cutler (2013) analyseerden enkel woordspelingen waarvan het doelwoord door Dokumumushi Sandayuu genoemd werd en het bronwoord door iemand anders dan Dokumumushi Sandayuu genoemd werd. Dit werd gedaan zodat de woordspelingen die zij analyseerden zo natuurlijk mogelijk waren. Spontane woordspelingen geven een betrouwbaarder beeld weer van het proces van woordspelingen maken dan ingestudeerde woordspelingen.

Net zoals in het onderzoek van Otake en Cutler (2013) zijn in dit onderzoek drie soorten woordspelingen onderzocht: homofonen, inbeddingen en mutaties. De nadruk lag echter op inbeddingen en mutaties. Bij homofonen is er altijd sprake van een gehele overeenkomst tussen bronwoord en doelwoord. Inbeddingen en mutaties zijn dan alleen bruikbaar voor onderzoek naar onsetoverlap en offsetoverlap. De inbeddingen en mutaties die gebruikt zijn voor de analyses, zijn in Bijlage 1 te vinden.

1.6 Onderzoeksvragen en hypothesen

Hoofdvraag

Hoe belangrijk is de rol die de offset van het bronwoord speelt in de totstandkoming van het doelwoord bij het maken van woordspelingen zoals mutaties en inbeddingen?

Bij deze hoofdvraag horen de volgende deelvragen met de daarbij horende hypothesen:

Deelvraag 1

Hoe vaak komt offsetoverlap tussen het bronwoord en het doelwoord voor in vergelijking met onsetoverlap tussen het bronwoord en het doelwoord?

Hypothese 1

De hypothese die bij deze vraag gesteld wordt, is dat offsetoverlap tussen bronwoord en doelwoord frequenter zal zijn dan onsetoverlap tussen bronwoord en doelwoord. Deze verwachting is gebaseerd op de bevinding van Otake en Cutler (2013). Zij gaven aan dat combinaties zoals *candle* en *handle* (offsetoverlap) eerder gebruikt worden bij het maken van een woordspeling dan combinaties zoals *candle* en *candy* (onsetoverlap).

Deelvraag 2

Is er een verschil in de gemiddelde Levenshtein distance tussen onsetoverlap en offsetoverlap?

Hypothese 2

Er wordt geen verschil verwacht tussen de gemiddelde Levenshtein distance wanneer de categorieën onsetoverlap en offsetoverlap met elkaar vergeleken worden. De Levenshtein distance wordt berekend op basis van de lengte van het langste woord en het aantal letters dat verschilt tussen bronwoord en doelwoord. De Levenshtein distance blijft hetzelfde ongeacht of *pak* en *mak* of

pak en *pad* als paar geselecteerd wordt voor een woordspeling. De classificatie in onsetoverlap of offsetoverlap doet er dan niet toe. Ondanks dat er geen verschil verwacht wordt, wordt de gemiddelde Levenshtein distance per type overlap onderzocht omdat het mogelijk is dat de offset of de onset resulteert in een grotere Levenshtein distance. Op deze manier kan geconcludeerd worden of het type overlap een rol speelt bij de hoeveelheid overlap die plaatsvindt tussen bronwoord en doelwoord.

Deelvraag 3

Is er een verschil in de gemiddelde Levenshtein distance tussen onsetoverlap en offsetoverlap wanneer er onderscheid gemaakt wordt tussen inbedding en mutatie?

Hypothese 3

Er wordt een verschil in Levenshtein distance verwacht wanneer de Levenshtein distance per type woordspeling gemeten wordt. Er wordt geen verschil verwacht in de gemiddelde Levenshtein distance tussen het type overlap binnen inbeddingen en mutaties. Hiervoor geldt dezelfde argumentatie als bij hypothese 2. Het type woordspeling op zichzelf zal wel effect hebben op de Levenshtein distance. Mutaties zullen een grotere Levenshtein distance kennen dan inbeddingen. De onsetoverlap en offsetoverlap zullen bij mutaties groter zijn dan bij inbeddingen. Dit wordt verwacht omdat het verschil in letters tussen bronwoord en doelwoord bij mutaties in de meeste gevallen kleiner zal zijn dan het verschil in letters tussen bronwoord en doelwoord bij inbeddingen. Wanneer het bronwoord *instrumentaal* en het daarbij horende doelwoord *mentaal* (inbedding) vergeleken wordt met het bronwoord *Donners* en het daarbij horende doelwoord *Donders* (mutatie), is er bij de inbedding een verschil van zes letters en bij de mutatie een verschil van één letter. De gemiddelde Levenshtein distance ligt dan hoger bij de mutatie. Er is een mogelijkheid dat zowel het type overlap als het type woordspeling effect heeft op de Levenshtein distance. Ook al is de verwachting dat het type overlap geen rol speelt bij de totstandkoming van de Levenshtein distance, het mag niet worden uitgesloten dat het type overlap wel een rol speelt. Het niet spelen van een rol in de totstandkoming van de Levenshtein distance kan gedefinieerd worden als de rol ‘geen invloed’. Het is dus in beide gevallen van belang om de rol van het type overlap te onderzoeken.

2. Methoden

Voor dit onderzoek zijn 30 opnames van Gerard Ekdom uit het radio-programma Effe Ekdom getranscribeerd en geannoteerd. De opnames van Effe Ekdom zijn in een periode van 14-04-2014 t/m 31-07-2015 opgeslagen. Het radioprogramma werd elke week van maandag t/m vrijdag uitgezonden. Feestdagen en soortgelijke dagen kwamen te vervallen; op deze dagen werd het programma niet uitgezonden. De laatste opname die voor dit onderzoek getranscribeerd en geannoteerd is, is de opname van 29-05-2015.

De opnames zijn in mp3-formaat opgeslagen op de site <http://mirjambroersma.nl/ekdom/ekdom.html>. In dit onderzoek werden de opnames vanaf 27-03-2015 t/m 29-05-2015 opgeslagen. Voor het beluisteren van de opnames is in dit onderzoek het programma VLC gebruikt. Het programma VLC maakt het mogelijk om toetscombinaties te koppelen aan acties die de opname een aantal seconden vooruit- of achteruitspoelen. Er is gekozen

om vier toetscombinaties te linken: twee voor vijf seconden voor- en achteruit en twee voor tien seconden voor- en achteruit. Het linken van deze toetsen aan het programma VLC is met de applicatie BetterTouchTool uitgevoerd. Op deze manier kon sneller en efficiënter getranscribeerd worden. Telkens als er een woordspeling te horen was, werd deze getranscribeerd.

Voor het analyseren van de woordspelingen is een al eerder opgezet corpus van woordspelingen van Gerard Ekdome uit het radioprogramma Effe Ekdome gebruikt (Kerkhof, 2015). Het annotatieschema van Cutler en Otake (2013), waarbij verschillende eigenschappen van woordspelingen geanalyseerd werden, werd voor eerder onderzoek naar Nederlandse woordspelingen aangepast (Kerkhof, 2015). Dit aangepaste annotatieschema (Bijlage 2) werd in dit onderzoek gebruikt om woordspelingen te analyseren. Het al eerder opgezette corpus van de woordspelingen van Gerard Ekdome werd met dit onderzoek uitgebreid. Het corpus bestond voor de huidige situatie uit 410 woordspelingen. Deze 410 woordspelingen zijn door twee transcribenten opgenomen in het corpus. Het aantal werd uiteindelijk uitgebreid tot 612 woordspelingen. De verdeling van homofonen, inbeddingen en mutaties was als volgt: 284 homofonen, 195 inbeddingen en 133 mutaties. In dit onderzoek zijn enkel de inbeddingen en mutaties voor verdere analyses gebruikt. In totaal waren dit 328 woordspelingen

2.1 Geannoteerde Eigenschappen

Verschillende eigenschappen van de woordspelingen van Gerard Ekdome werden geannoteerd. De belangrijkste eigenschappen voor dit onderzoek zullen kort besproken worden. De volledige lijst aan geannoteerde eigenschappen is te zien Bijlage 2.

Soort woordspeling

Een woordspeling werd geannoteerd wanneer er sprake was van een homofoon, inbedding of mutatie.

Bronwoord

Het woord dat door de grappenmaker uit een uiting geselecteerd wordt om een woordspeling mee te maken.

Doelwoord

Het woord dat een humoristisch effect moet geven naar aanleiding van het gekozen bronwoord.

Positie overlap

De positie waarin het bronwoord en doelwoord overlappen. De positie van overlap kon onset (begin), offset (einde) of zowel onset als offset (begin en einde) betreffen. In Tabel 1 en Tabel 2 zijn voorbeelden te zien van onset-, offset- en zowel onset- als offsetoverlap tussen bronwoord en doelwoord bij inbeddingen en mutaties.

Woorden die zowel in het begin als aan het einde overlappen, zijn eveneens opgenomen in dit onderzoek voor verdere analyses. Voorafgaand aan dit onderzoek werd geen rekening gehouden

met deze categorie. Deze categorie betrof echter een redelijk groot aantal, waardoor deze niet genegeerd kon worden. Van de 328 items waren er 42 items waarbij zowel de onset als de offset tussen bron- en doelwoord overlaptten. Bij de inbeddingen waren er vier items waar zowel de onset als de offset overlaptten. Bij mutaties waren dit er 38.

Tabel 1. Voorbeelden van de drie typen overlap bij inbeddingen

bronwoord	doelwoord	positie overlap
belastingdienst	belastingvrije	begin (onset)
spoorwegen	aanwegen	einde (offset)
ledendag	ledematendag	begin en einde (zowel onset als offset)

Tabel 2. Voorbeelden van de drie typen overlap bij mutaties

bronwoord	doelwoord	positie overlap
Berger	bergen	begin (onset)
gek	track	einde (offset)
blender	blunder	begin en einde (zowel onset als offset)

2.2 Levenshtein distance

Voor het berekenen van de Levenshtein distance is het verschil in letters tussen bronwoord en doelwoord vastgesteld. Ook werden het aantal letters van het bronwoord en het doelwoord berekend. De Levenshtein distance werd als volgt berekend: $1 - (\text{het verschil in letters tussen bronwoord en doelwoord gedeeld door het aantal letters van het langste woord})$, zie Formule 1. Bij het bronwoord *last* met het daarbijbehorende doelwoord *lastig*, werd er een verschil van twee genoteerd. Wanneer er sprake was van substitutie zoals bij *retriever* en *receiver* werd de substitutie van *t* in *c* geteld als een verschil van één en de deletie van de *r* ook als één. De substitutie van *i* in *e* telde als een verschil van één en de substitutie van *e* in *i* ook. De drie substituties en één deletie resulteren in een verschil van vier tussen het bronwoord *retriever* en het doelwoord *receiver*.

De karakterlengtes van het bronwoord en het doelwoord zijn in Excel berekend met de formule =LENGTE(). Vervolgens zijn de lengtes van het bronwoord en het doelwoord met elkaar vergeleken en werd er een kolom aangemaakt voor het woord met de langste lengte. Het verschil in aantal letters tussen bronwoord en doelwoord werd handmatig geteld. Met het verschil in letters tussen bronwoord en doelwoord en de waarden van de langste lengte in karakters van het bronwoord of doelwoord (afhankelijk van welke van de twee meer karakters telt), kon de Levenshtein distance berekend worden. Als voorbeeld het bronwoord *eronder* en het doelwoord *onder*. De overlap bij *eronder* en *onder* werd geclassificeerd als offsetoverlap (hetzelfde geldt voor paren die vanaf de onset overlappen). Het verschil in letters tussen de twee woorden is twee. De lengte van het langste woord, van het woord *eronder*, betreft zeven. De formule voor het berekenen van de Levenshtein distance wordt dan $1 - (2/7) = 0.71$. De Levenshtein distance voor de woorden *onder* en *eronder* betreft dus 0.71, dit betekent dan dat er een overlap van 71% is tussen de offset

van het bronwoord en het doelwoord. Dit hele deel werd in dit onderzoek als offset beschouwd om zo een betere vergelijking te kunnen maken tussen onsetoverlap en offsetoverlap. Om woorden met kortere lengtes te kunnen vergelijken met langere lengtes werd de formule van Schepens, Dijkstra & Grootjen (2012) gebruikt.

$$score = 1 - \frac{distance}{length}$$

$$length = \max(length\ of\ source\ expression,\ length\ of\ destination\ expression)$$

$$distance = \min(number\ of\ insertions,\ deletions\ and\ substitutions)$$

Formule 1. Levenshtein distance berekening voor orthografische overlap (Schepens, Dijkstra, & Grootjen (2012).

De lengtes van de woorden en het verschil tussen bronwoord en doelwoord zijn in letters geteld (orthografisch), als benadering voor het aantal klanken. Deze manier leek het meest haalbaar. Het tellen van klanken kan erg subjectief zijn, er bestaat veel variatie in uitspraak van Nederlandse woorden. Daarnaast verschijnen bij het maken van woordspelingen non-woorden. Non-woorden kennen geen officiële uitspraak, waardoor IPA-notaties ontbreken. Ook is het automatisch tellen van klanken (IPA) erg moeilijk omdat er bij IPA-notatie veel verschillende karakters betrokken zijn die nuances aan klanken geven waardoor het moeilijk is om dit soort tekens te koppelen aan klanken en te verwerken in een programma. Om deze redenen is ervoor gekozen om de lengtes op basis van orthografische gelijkheid vast te stellen en niet op basis van fonologische gelijkheid.

2.3 Analyses

Voor het vergelijken van het aantal frequenties van de overlapcategorieën (onset, offset en zowel onset als offset) werden met behulp van SPSS Chi-kwadraattoetsen uitgevoerd. De groep offsetoverlap werd eerst vergeleken met de groep onsetoverlap. Daarna werd de groep offsetoverlap vergeleken met de groep zowel onset- als offsetoverlap. Vervolgens werd de groep offsetoverlap vergeleken met onsetoverlap en zowel onset- als offsetoverlap als één groep. Ook werd de groep zowel onset- als offsetoverlap vergeleken met de andere twee groepen.

Voor het onderzoeken van het verschil in Levenshtein distance tussen de verschillende typen overlap en typen woordspelingen, werd er met behulp van SPSS een Analysis of Variance (ANOVA) uitgevoerd. Op deze manier kon onderzocht worden wat de effecten van het type woordspeling en van de typen overlap waren op de Levenshtein distance. Ook kon een eventueel interactie-effect onderzocht worden.

De afhankelijke variabele betrof de Levenshtein distance waarde per woordspeling (per bronwoord en doelwoord). Er waren twee onafhankelijke variabelen: type woordspeling en positie overlap. De variabele type woordspeling was verdeeld in twee niveaus: inbedding en mutatie. De variabele positie overlap was verdeeld in drie niveaus: onset, offset en zowel onset als offset. Op

deze manier konden de gemiddelde waarden van de Levenshtein distance per type overlap en per type woordspeling met elkaar vergeleken worden.

3. Resultaten

In dit onderzoek werden in totaal 328 woordspelingen, waarvan 195 inbeddingen en 133 mutaties, onderzocht. Het verschil in frequentie tussen de typen overlap werd onderzocht. Ook werd het verschil in gemiddelde Levenshtein distance per type overlap en per type woordspeling onderzocht.

3.1 Hypothese 1: Verschil in frequentie tussen offsetoverlap en onsetoverlap

Van de 328 items waren er 178 items waarbij er sprake was van offsetoverlap, 108 items waarbij er sprake was van onsetoverlap en 42 items waarbij er sprake was van zowel onset- als offsetoverlap. Deze resultaten zijn weergegeven in figuur 1. De frequenties zijn terug te vinden in tabel 3. Het is duidelijk te zien dat het aantal woordspelingen waarbij bronwoord en doelwoord in offset overlappen groter is dan het aantal woordspelingen waarbij bronwoord en doelwoord in onset overlappen. Zowel onset- als offsetoverlap tussen bronwoord en doelwoord kent de laagste frequentie.

Uit Chi-kwadraattoetsen bleek dat het aantal woordspelingen waarbij bronwoord en doelwoord in offset overlappen, significant groter was dan het aantal woordspelingen waarbij bronwoord en doelwoord in onset overlappen $\chi^2(1) = 17,13$, $p < .001$. Tevens was er een significant effect wanneer offsetoverlap vergeleken werd met zowel onset- als offsetoverlap $\chi^2(1) = 84,07$, $p < .001$. Er werd geen significant effect gevonden wanneer offsetoverlap vergeleken werd met onsetoverlap en zowel onset- als offsetoverlap (offsetoverlap als groep vergeleken met onsetoverlap en zowel onset- als offsetoverlap als één groep) $\chi^2(1) = 2,39$, $p = .12$. Zowel onset- als offsetoverlap bleek significant minder frequent te zijn dan de andere twee typen overlap $\chi^2(1) = 318,00$, $p < .001$.

3.2 Hypothese 2: Verschil in gemiddelde Levenshtein distance tussen offsetoverlap en onsetoverlap

De woordspelingen werden in drie overlapcategorieën (onset, offset, zowel onset als offset) opgedeeld. Van elk van deze categorieën werd de gemiddelde Levenshtein distance berekend. De gemiddelde Levenshtein distance voor offsetoverlap was .63. De gemiddelde Levenshtein distance voor onsetoverlap was .60. De Gemiddelde Levenshtein distance voor zowel onset- als offset overlap bedroeg .79. De waarden van de gemiddelde Levenshtein distance per type overlap zijn terug te vinden in Tabel 4. Het type overlap tussen het bronwoord en het doelwoord had een significant effect op de gemiddelde Levenshtein distance ($F(2,322) = 4.91$, $p < .05$).

Met behulp van Bonferroni procedures werd het effect van het type overlap op de Levenshtein distance verder onderzocht. Deze Bonferroni procedures lieten zien dat het verschil in gemiddelde Levenshtein distance tussen onsetoverlap en offsetoverlap niet significant was ($p > .05$). De gemiddelde Levenshtein distance van zowel onset- als offsetoverlap was echter significant hoger dan de gemiddelde Levenshtein distance van alleen onsetoverlap ($p < .001$) en alleen offsetoverlap ($p < .001$). Deze resultaten zijn weergegeven in figuur 2.

3.3 Hypothese 3: Verschil in Levenshtein distance tussen inbeddingen en mutaties

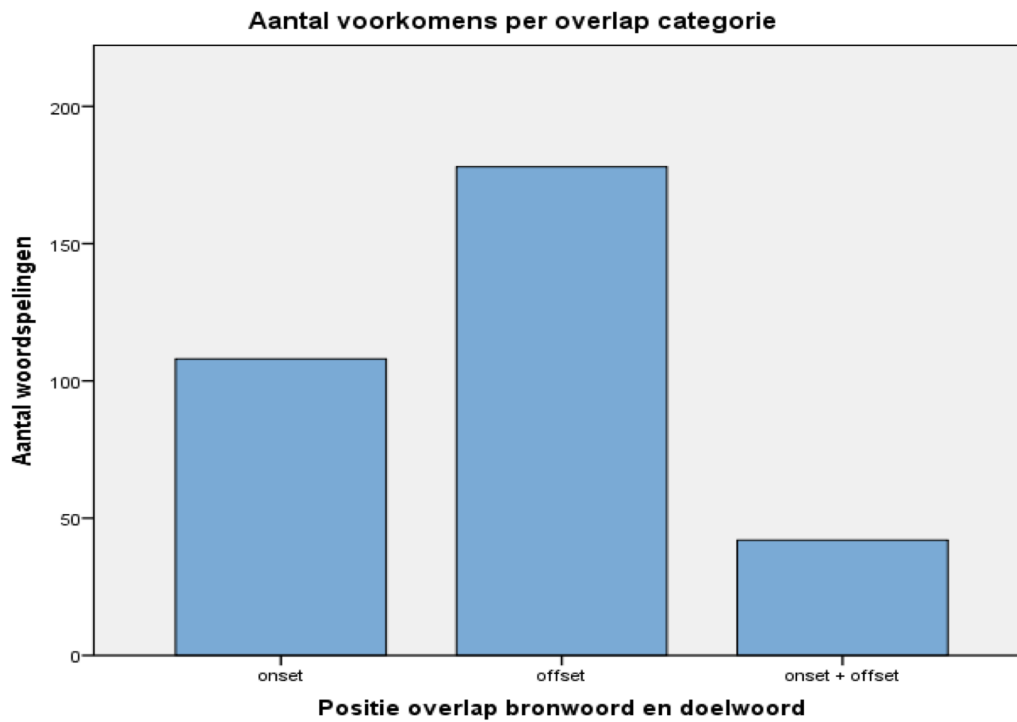
De gemiddelde Levenshtein distance per overlapcategorie werd ook per soort woordspeling berekend. Voor de inbeddingen waarbij de onset van het bronwoord en het doelwoord overlappen, was er een gemiddelde Levenshtein distance van .57. Voor de inbeddingen waarbij de offset van het bronwoord en het doelwoord overlappen, was er een gemiddelde Levenshtein distance van .58. Voor de inbeddingen waarbij zowel de onset als de offset van het bronwoord en het doelwoord overlappen, was er een gemiddelde Levenshtein distance van .66.

Voor de mutaties waarbij de onset van het bronwoord en het doelwoord overlappen was er een gemiddelde Levenshtein distance van .67. Voor de mutaties waarbij de offset van het bronwoord en het doelwoord overlappen was er een gemiddelde Levenshtein distance van .73. Voor de mutaties waarbij zowel de onset als de offset van het bronwoord en het doelwoord overlappen, was er een gemiddelde Levenshtein distance van .79. Deze waarden van de Levenshtein distance zijn eveneens te zien in tabel 4.

Uit de ANOVA bleek dat het type woordspeling een significant effect had op de gemiddelde Levenshtein distance ($F(1,322) = 25.23$, $p < .001$). De gemiddelde Levenshtein distance was significant hoger bij mutaties dan bij inbeddingen. Er werd geen significant interactie-effect tussen type woordspeling en type overlap gevonden ($F(2,322) = 1.26$, $p > .05$). In figuur 3 is de normaalverdeling van de Levenshtein distance per type woordspeling en per type overlap weergegeven. Het is duidelijk te zien dat bij elk type overlap de normaalverdeling bij mutaties meer naar rechts ligt dan bij inbeddingen (gemiddelden liggen hoger bij mutaties). Ook is te zien dat de hogere Levenshtein distance waarden frequenter zijn bij mutaties dan bij inbeddingen.

Tabel 3. Frequentie per type overlap en type woordspeling

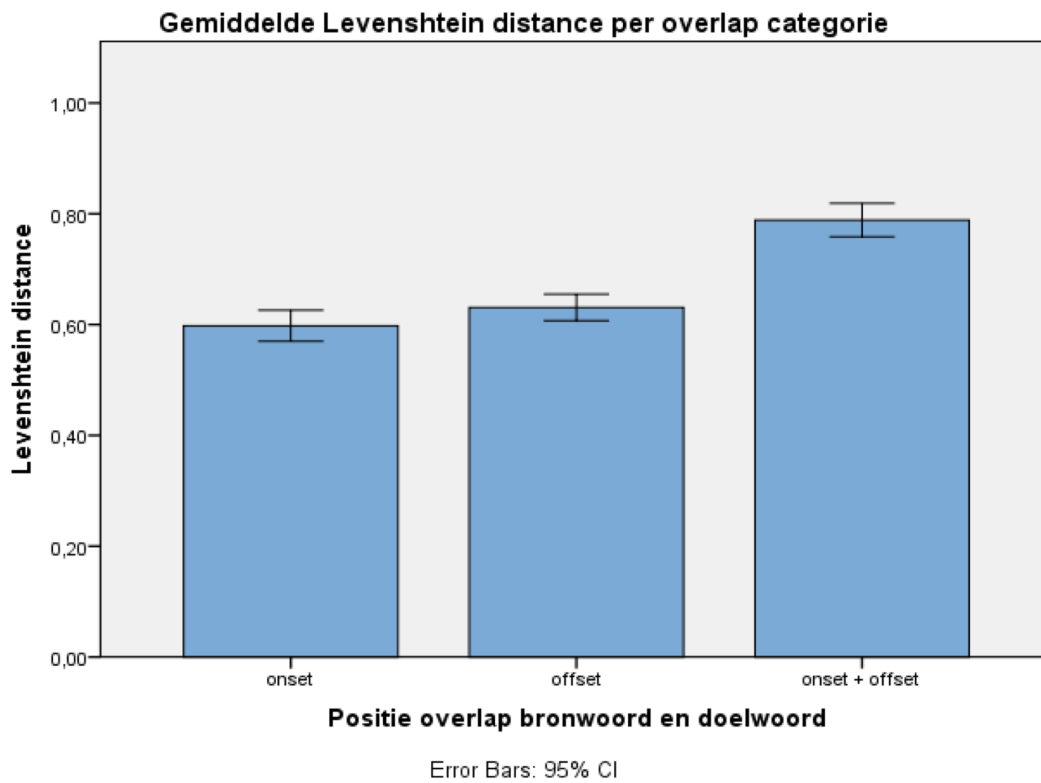
		Type woordspeling		
		inbedding	mutatie	totaal
Type overlap	onset	76	32	108
	offset	115	63	178
	zowel onset als offset	4	38	42
	totaal	195	133	



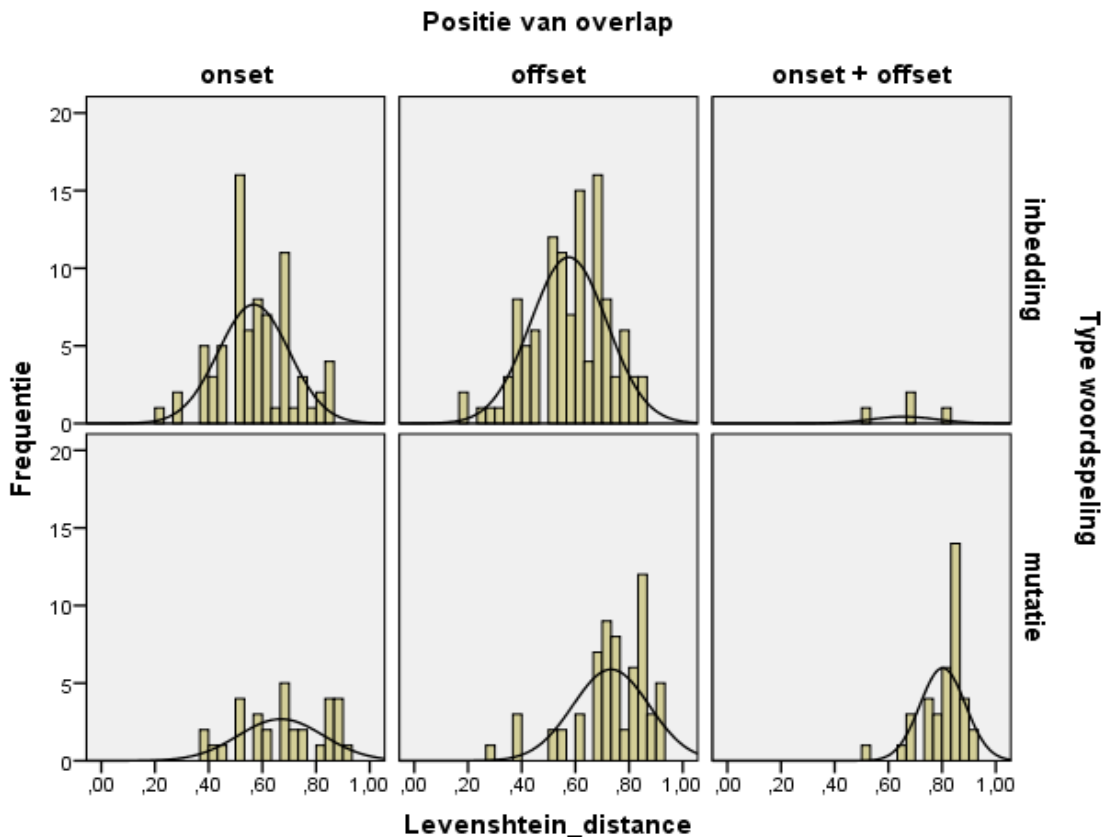
Figuur 1. Frequentie per type overlap: onset overlap, offset overlap en zowel onset als offset overlap

Tabel 4. *Gemiddelde Levenshtein distance per type overlap en type woordspeling*

		Type woordspeling		
		inbedding	mutatie	gemiddelde score per type overlap
Type overlap	onset	.57	.67	.60
	offset	.58	.73	.63
	zowel onset als offset	.66	.80	.79
gemiddelde score per type woordspeling		.57	.74	



Figuur 2. Gemiddelde Levenshtein distance per type overlap: onset, offset en zowel onset als offset



Figuur 3. Normaalverdeling van Levenshtein distance per type overlap en type woordspeling

4. Discussie

In deze studie is onderzocht hoe belangrijk de rol is die de offset van het bronwoord speelt in de totstandkoming van het doelwoord bij het maken van woordspelingen zoals mutaties en inbeddingen. Door dit te onderzoeken is geprobeerd om indicaties te geven voor het bestaan van een “woordspelingsmodus”. De aanname in dit onderzoek is dat de activatie van concurrenten van een targetwoord (bronwoord) door de grappenmaker in enige mate te controleren is (Otake & Cutler, 2013). Om de rol van de offset te onderzoeken, is onderzocht hoe vaak onsetoverlap tussen bronwoord en doelwoord in vergelijking met offsetoverlap tussen bronwoord en doelwoord voorkwam. Ook is onderzocht wat de gemiddelde Levenshtein distance per type overlap was. Op deze manier kon onderzocht worden of het ene type overlap in een grotere overeenkomst tussen bronwoord en doelwoord resulteert dan het andere type overlap. Daarnaast is onderzocht of het type woordspeling een verschil in Levenshtein distance kent. Ook is nog onderzocht of er een verschil was tussen de gemiddelde Levenshtein distance van de typen overlap (onset, offset en zowel onset als offset), als er een onderscheid tussen inbedding en mutatie gemaakt werd (interactie-effect).

Op basis van de eerdere besproken literatuur van Otake en Cutler (2013) werd verwacht dat de offset tussen bronwoord en doelwoord vaker zou overlappen dan de onset (zie hypothese 1). Voor het type overlap, onset of offset, werd verwacht dat er geen verschil zal zijn in de gemiddelde Levenshtein distance (zie hypothese 2). Daarnaast werd verwacht dat het type woordspeling invloed

zou hebben op de Levenshtein distance, waarbij mutaties een hogere Levenshtein distance hebben dan inbeddingen (zie hypothese 3).

4.1 Discussie van deelvragen (hypothese 1, 2 en 3)

Als eerste kan geconcludeerd worden dat homofonen (N=284) de meest frequente woordspelingen van Gerard Ekdome zijn, gevolgd door inbeddingen (N=195) en mutaties (N=133). Ook in het onderzoek van Otake en Cutler (2013) waren homofonen het meest frequent, tevens gevolgd door inbeddingen en dan mutaties. De conclusies dat homofonen de meest frequente woordspelingen zijn en dat een gewenste woordspeling voldoet aan MaxSP blijven overeind (Otake & Cutler, 2013; Otake, 2010; Kawahara & Shinohara (2009); Cutler & Otake, 2002; Otake & Cutler, 2001, Lagerquist, 1980). De suggestie dat een woordspeling voldoet aan MaxSP wordt ondersteund doordat homofonen de meest frequente woordspelingen vormen. Homofonen zijn woorden die hetzelfde klinken en dus qua klank geen veranderingen ondergaan. Homofonen zijn volgens MaxSP het meest optimaal omdat er geen veranderingen tussen bronwoord en doelwoord plaatsvinden. Dit wordt door deze studie en door de studie van Otake en Cutler (2013) ondersteund doordat homofonen in beide studies het meest frequent zijn.

Uit interesse zijn de frequenties van de woordspelingen uit de twee onderzoeken nog onderzocht. Er is onderzocht of de frequenties van de woordspelingen uit het onderzoek van Otake en Cutler (2013) significant van elkaar verschillen. Dit is ook gedaan voor de frequenties van de woordspelingen uit dit onderzoek. De resultaten zijn terug te vinden in Bijlage 3. Er is gekozen om dit in de bijlage te verwerken omdat de frequenties uit het onderzoek van Otake en Cutler (2013) niet volledig vergelijkbaar zijn met de frequenties uit dit onderzoek. In het onderzoek van Otake en Cutler (2013) zijn de woordspelingen niet geanalyseerd wanneer het bronwoord door Dokumumushi Sandayuu zelf benoemd werd. In dit onderzoek zijn de woordspelingen waarbij Gerard Ekdome het bronwoord zelf benoemde wel geanalyseerd (naast de woordspelingen waarbij alleen het doelwoord door Gerard Ekdome benoemd werd).

Hypothese 1: verschil in frequentie tussen offsetoverlap en onsetoverlap

Uit dit onderzoek blijkt dat offsetoverlap (N=178) tussen bronwoord en doelwoord het meest frequent is. Onsetoverlap tussen bronwoord en doelwoord is minder frequent (N=108) en zowel onset- als offsetoverlap is het minst frequent (N=42) van de drie typen overlap. Het verschil tussen offsetoverlap en onsetoverlap is significant ($p < .001$). Ook het verschil tussen offsetoverlap en zowel onset- als offsetoverlap is significant ($p < .001$). Het verschil tussen offsetoverlap en onsetoverlap en zowel onset- als offsetoverlap (offsetoverlap als groep vergeleken met onsetoverlap en zowel onset- als offsetoverlap als één groep) is echter niet significant ($p = .12$). Hiermee kan geconcludeerd worden dat offsetoverlap tussen bronwoord en doelwoord significant vaker plaatsvindt dan onsetoverlap of zowel onset- als offsetoverlap. De gestelde hypothese 1 wordt bevestigd. Er lijkt sprake te zijn van een voorkeur om de offset te laten overlappen tussen bronwoord en doelwoord, omdat offsetoverlap tussen bronwoord en doelwoord significant vaker plaatsvindt dan onsetoverlap of zowel onset- als offsetoverlap. De bevinding van Otake & Cutler (2013) dat paren zoals *candle* en *handle* eerder gekozen worden als doelwoord voor een woordspeling dan paren zoals *candle* en *candy*, wordt door dit onderzoek ondersteund.

Hypothese 2: Verschil in gemiddelde Levenshtein distance tussen offsetoverlap en onsetoverlap

De gemiddelde Levenshtein distance is bij offsetoverlap ($M = .63$) hoger dan bij onsetoverlap ($M = .60$). Dit verschil is niet significant ($p > .05$). Echter, de Levenshtein distance bij zowel onset- als offsetoverlap ($M = .79$) is significant hoger dan bij alleen onsetoverlap ($p < .001$) en bij alleen offsetoverlap ($p < .001$).

De verwachting dat de gemiddelde Levenshtein distance tussen bronwoord en doelwoord niet significant zou verschillen tussen onsetoverlap en offsetoverlap wordt ondersteund door het onderzoek (zie Hypothese 2). Over de gemiddelde Levenshtein distance bij zowel onset- als offsetoverlap zijn geen verwachtingen opgesteld, omdat de nadruk in eerste instantie lag bij de vergelijking onset- tegenover offsetoverlap. De significant hoger gemiddelde Levenshtein distance bij zowel onset- als offsetoverlap is echter geen verrassing. Als op twee plaatsen overlap is, is de kans groter dat deze overlap groter is dan wanneer er maar op één plaats overlap is.

Hypothese 3: Verschil in Levenshtein distance tussen inbeddingen en mutaties

Er is een significant effect ($p < .001$) gevonden van het type woordspeling op de gemiddelde Levenshtein distance. Mutaties ($M = .74$) hebben een significant hogere gemiddelde Levenshtein distance dan inbeddingen ($M = .57$). Deze resultaten komen overeen met de gestelde hypothese 3. Er is geen interactie-effect tussen type overlap en type woordspeling gevonden ($p > .05$).

Methodologische overwegingen

Voor de berekening van de Levenshtein distance moet opgemerkt worden dat een berekening op basis van fonologische gelijkheid representatiever zou zijn dan een berekening op basis van orthografische gelijkheid. Zoals in de methodensectie al besproken is, wordt er tussen *retriever* en *receiver* een verschil van vier letters genoteerd. Op basis van fonologische gelijkheid zou het verschil in klanken maar twee zijn omdat de *ei* in *receiver* en de *ie* in *retriever* hetzelfde klinken. Wanneer de fonologische gelijkheid gemeten zou worden, wordt de Levenshtein distance groter bij dit soort paren. Ook bij Engelse of Franstalige bron- en/of doelwoorden is de berekende overlap niet representatief voor de data uit dit onderzoek wanneer de overlap orthografisch berekend wordt. Het meten van de fonologische gelijkheid zou representatiever zijn geweest, maar zoals in de methodensectie al besproken is, is het meten van orthografische gelijkheid het meest haalbaar. Het handmatig tellen van klanken kan erg subjectief zijn. Wat door de ene persoon als één klank beschouwd wordt, kan door een andere als bijvoorbeeld twee klanken beschouwd worden. Daarnaast worden bij woordspelingen niet bestaande woorden gebruikt en zal er geen uitspraak bekend zijn van deze woorden. Het automatisch berekenen van klanken is erg moeilijk omdat in fonetisch schrift veel verschillende tekens gebruikt worden. Het is moeilijk om deze tekens te verwerken in een programma. Er is gekozen om de orthografische gelijkheid tussen bronwoord en doelwoord te berekenen, omdat de haalbaarheid hiervoor groter was dan het berekenen van de fonologische gelijkheid.

4.2 Algemene discussie

Nu de drie hypothesen besproken zijn, kan er ook antwoord gegeven worden op de hoofdvraag: Hoe belangrijk is de rol die de offset van het bronwoord speelt in de totstandkoming van het doelwoord bij het maken van woordspelingen zoals mutaties en inbeddingen?

De rol van de offset van het bronwoord speelt een belangrijke rol bij de totstandkoming van het doelwoord. Er vindt significant vaker offsetoverlap dan onsetoverlap. Ook vindt er significant vaker offsetoverlap plaats dan zowel onset- als offsetoverlap. De offset speelt echter geen andere rol dan de onset bij de totstandkoming van de gemiddelde Levenshtein distance. De gemiddelde overlap is nagenoeg hetzelfde wanneer onsetoverlap en offsetoverlap met elkaar vergeleken worden. De offset speelt dus geen rol bij de mate van overlap.

Doordat offsetconcurrenten vaker gekozen worden om als doelwoord te fungeren, lijken offsetconcurrenten van het bronwoord geschiktere kandidaten te zijn voor het doelwoord dan onsetconcurrenten. Dit bevestigt de bevindingen van Otake en Cutler (2013). Op basis van dit onderzoek en het onderzoek van Otake en Cutler (2013) kan gemotiveerd worden dat offsetconcurrenten de voorkeur krijgen bij het maken van woordspelingen. Het is aannemelijk dat de intentie van woordspelingen maken ervoor zorgt dat er een soort woordspelingmodus geactiveerd wordt en dat woordactivatie in enige mate te controleren is. Tijdens het maken van woordspelingen blijken offsetconcurrenten geschiktere kandidaten te zijn voor een woordspeling omdat deze frequenter voorkomen. In veel gevallen is het niet zo dat offsetconcurrenten de enige mogelijkheid vormen als doelwoord. In de meeste situaties waarin er gekozen werd voor een offsetconcurrent als doelwoord, had er ook een onsetconcurrent kunnen verschijnen. Toch koos Gerard Ekdorf vaker voor een offsetconcurrent. Alle mogelijke concurrenten van het bronwoord die het doelwoord kunnen vormen, zouden onderzocht kunnen worden. Op deze manier kan bekeken worden wat de waarschijnlijkheid is van het verschijnen van een onset- of offsetconcurrent.

Het zou kunnen dat het selecteren van offsetconcurrenten makkelijker wordt voor de herkenning van woordspelingen wanneer er intentie is tot het bewust maken van woordspelingen. Ter illustratie de volgende twee paren: *haast – blaast* en *haast – haar*. Het paar *haast – blaast* kent offsetoverlap. Het paar *haast – haar* kent onsetoverlap. Omdat offsetoverlap significant vaker plaatsvindt dan onsetoverlap wordt in deze studie aangenomen dat paren zoals *haast – blaast* door zowel de grappenmaker als de hoorder eerder herkend worden als woordspeling dan paren zoals *haast – haar*. Rijm zou dan een positief effect kunnen hebben op de herkenning van woordspelingen (Lupker & Colombo, 1994). Een andere mogelijkheid voor het kiezen van offsetconcurrenten is het vergroten van het bewustzijn van de woordspeling. Het bewustzijn van de gemaakte woordspeling bij de hoorder zou groter kunnen zijn wanneer het “humoristische effect”, dat wat de woordspeling een woordspeling maakt, in het einde van het woord zit omdat het einde van een woord als laatste verwerkt wordt. Het zou echter ook zo kunnen zijn dat de frequentere verschijning van offsetoverlap tussen bronwoord en doelwoord berust op toeval.

Om sterkere conclusies te kunnen trekken over activatie van offsetconcurrenten bij het maken van woordspelingen, zal er onderzoek gedaan moeten worden waarbij de woordactivatie direct gemeten wordt. Op basis van de corpusdata uit dit onderzoek kunnen er slechts indicaties gegeven worden over een mogelijkheid dat offsetconcurrenten sterker concurreren tijdens het maken van woordspelingen. Dit geeft aanleiding voor verder onderzoek naar het controleren van woordactivatie en het hele proces van het maken van een woordspeling.

Vervolgonderzoek

Uit dit onderzoek blijkt dat sommige woordspelingen niet passen in het annotatieschema dat gebruikt werd om de woordspelingen te analyseren (Kerkhof, 2015). In Tabel 5 staan voorbeelden van woordspelingen die een combinatie zijn van inbedding en mutatie. In het Japans leken dit soort woordspelingen niet voor te komen (Otake & Cutler, 2013). In het Nederlands worden dit soort woordspelingen regelmatig gemaakt, in ieder geval door Gerard Ekdom. Voor vervolgonderzoek waarin woordspelingen geanalyseerd worden, zou er een categorie kunnen worden toegevoegd aan het type woordspeling: een combinatie van een inbedding en een mutatie. De eigenschappen die bij alle typen woordspelingen geannoteerd worden kunnen gecombineerd worden met de eigenschappen die alleen bij inbeddingen geannoteerd worden en de eigenschappen die alleen bij mutaties geannoteerd worden (zie Bijlage 2). Op deze manier kan de combinatiecategorie van inbedding en mutatie geanalyseerd worden.

Een optie zou zijn om een woordspeling waarbij inbedding plus insertie, deletie, en/of substitutie plaatsvindt, te bestempelen als een combinatie tussen inbedding en mutatie. Ter illustratie het bronwoord *Trainor* en het doelwoord *hometrainer* uit Tabel 5. De *o* in *trainor* wordt *e* in *hometrainer* (substitutie). Daarnaast wordt *trainor* samengevoegd met een ander woord, namelijk *home*. Er is dus sprake van zowel een inbedding als een mutatie. In sommige gevallen is het echter moeilijk om een indeling te maken, zelfs als er een extra categorie aan het annotatieschema wordt toegevoegd. De grens tussen inbedding en mutatie is soms moeilijk te onderscheiden. Neem nou het bronwoord *slavink* en het doelwoord *afvinken*. Het is duidelijk dat *vink* tussen bronwoord en doelwoord overlapt en dat dit gedeelte ingebed wordt in het doelwoord. De stukken die echter niet overlappen zorgen voor problemen bij de indeling. Het stukje *sla* in het bronwoord dat in het doelwoord *af* wordt, kan vanuit twee perspectieven bekeken worden. Het ene perspectief is dat deze verandering door zowel substitutie als insertie plaatsvindt en het andere is dat er sprake is van een inbedding, omdat er een nieuwe samenstelling gemaakt wordt met het woord *af* en *vink*. Ditzelfde probleem treedt op bij het bronwoord *trommelvliezen* en het doelwoord *trommel te verliezen*. Het is duidelijk dat *trommel* en *liezen* overlappen. Maar verschijnen *te* en *ver* door mutatie of door inbedding of een combinatie van beide? In het vervolg zouden er duidelijkere richtlijnen vastgesteld kunnen worden zodat het indelen in type woordspeling makkelijker en minder subjectief wordt.

Tabel 5. Voorbeelden van paren die gecategoriseerd kunnen worden als combinatie van mutatie en inbedding

Bronwoord	Doelwoord
trommelvliezen	trommel te verliezen
festijn	festival
trainor	Hometrainer
spinnende	Vogelspin
geominie	iniminie
Afrojack	Afrotrosjack
slavink	afvinken

Effect van rijm en frequentie

In het vervolg zou het effect van rijm en frequentie onderzocht moeten worden bij het maken van woordspelingen. Het blijkt dat rijmprimen ervoor zorgen dat het herkenningsproces vergemakkelijkt wordt tijdens lexicale decisietafen en benoemingtafen. Hoogfrequente onregelmatige rijmparen (wel fonologische overlap maar geen orthografische overlap) blijken echter voor een inhiberend effect te zorgen (Lupker & Colombo 1994). Het zou zo kunnen zijn dat dit ook effect heeft tijdens het maken van woordspelingen. Het aandeel van mutaties en inbeddingen dat overlapt in de offset van bronwoord en doelwoord is groter dan het aandeel van onset of zowel onset- als offsetoverlap. Bij de overlap in offset is het effect van rijm niet onderzocht, omdat niet alle woorden een dusdanige overlap hebben en er bij kleine overlap geen sprake is van rijm. Het is mogelijk dat het selecteren en het produceren van rijmconcurrenten makkelijker en efficiënter is voor de grappenmaker. Verder onderzoek zou een verklaring kunnen geven voor de voorkeur voor offsetconcurrenten boven onsetconcurrenten bij het maken van woordspelingen (Otake & Cutler, 2013).

Daarnaast is in deze studie slechts data van één persoon geanalyseerd. Het is moeilijk om conclusies te trekken over woordspelingen in het algemeen wanneer alle geanalyseerde woordspelingen van een persoon afkomstig zijn, in dit geval van Gerard Ekdorf. Een ander iemand zou bijvoorbeeld een voorkeur kunnen hebben voor andere type woordspelingen. Om sterkere uitspraken te kunnen doen, zullen er woordspelingen, die van verschillende bronnen afkomstig zijn, onderzocht moeten worden.

5. Conclusie

Uit de geanalyseerde woordspelingen van Gerard Ekdorf blijkt dat offsetoverlap tussen bronwoord en doelwoord het meest frequent is. Offsetoverlap tussen bronwoord en doelwoord vindt significant vaker plaats ($p < .001$) dan onsetoverlap tussen bronwoord en doelwoord bij het maken van woordspelingen. Ook vindt offsetoverlap significant vaker plaats dan zowel onset- als offsetoverlap ($p < .001$). Offsetoverlap vindt echter niet significant vaker plaats dan onsetoverlap en zowel onset- als offsetoverlap. De gestelde hypothese 1 wordt bevestigd.

Het type overlap heeft geen significant effect op de gemiddelde Levenshtein distance. De gestelde hypothese 2 wordt bevestigd. Het type woordspeling heeft wel een significant effect op de gemiddelde Levenshtein distance. Mutaties hebben een significant hoger gemiddelde Levenshtein distance dan inbeddingen. Er is geen interactie-effect van het type overlap en het type woordspeling gevonden. Ook de gestelde hypothese 3 wordt bevestigd. De bevinding dat offsetconcurrenten betere kandidaten lijken te zijn voor het maken van een woordspeling (Otake & Cutler, 2013) wordt door deze studie gesteund.

Deze studie geeft indicaties dat er een asymmetrie is tussen de activatie van concurrenten in normale omstandigheden en tijdens het maken van woordspelingen. Onsetconcurrenten lijken onder normale omstandigheden sterker te concurreren om activatie dan offsetconcurrenten (Huettig, Rommers & Meyer 2011; Magnuson, Tanenhaus, Aslin, & Dahan, 2003). Bij het maken van

woordspelingen lijkt het tegenovergestelde effect plaats te vinden: offsetconcurrenten lijken sterker te concurreren om activatie omdat offsetconcurrenten significant vaker geselecteerd worden als doelwoord dan onsetconcurrenten.

In deze studie zijn echter geen experimenten uitgevoerd om deze activatie te meten. Verder onderzoek zal dus moeten uitwijzen of offsetconcurrenten ook daadwerkelijk sterker concurreren om activatie tijdens het maken van woordspelingen. Er kan in ieder geval geconcludeerd worden dat de rol van de offset belangrijker wordt wanneer er een intentie is om woordspelingen te maken. De offset van het bronwoord komt het meest frequent terug in het doelwoord ten opzichte van de andere twee typen overlap (onset en zowel onset als offset). In deze studie wordt ervan uitgegaan dat de grappenmaker door een verandering in intentie, namelijk het bewust bezig zijn met woordspelingen, de activatie van concurrenten tijdens de spraakverwerking en –productie in enige mate kan controleren. Het zou kunnen zijn dat door deze switch van intentie het makkelijker wordt om woordspelingen te herkennen en te maken wanneer rijmconcurrenten geactiveerd blijven (Lupker & Colombo, 1994). Vanuit deze gedachte wordt aangenomen dat personen in staat zijn om een “woordspelingmodus” te activeren. Ook hierover kunnen geen sterke conclusies getrokken worden op basis van dit corpusonderzoek. Deze studie heeft echter wel aangegeven dat het bestaan van een woordspelingmodus aannemelijk is, omdat offsetconcurrenten een belangrijkere rol spelen tijdens het maken van woordspelingen dan in een normale situatie.

Woordspelingen geven in ieder geval een aanleiding tot verder psycholinguïstisch onderzoek. De mogelijkheden van personen om activatie in enige mate te controleren is een erg interessant onderwerp en kan inzichten geven in de capaciteiten om bewust activatie te onderdrukken of juist uit te lokken in verschillende omstandigheden.

Referenties

- Allopenna, P. D., Magnuson, J. S., & Tanenhaus, M. K. (1998). Tracking the time course of spoken word recognition using eye movements: Evidence for continuous mapping models. *Journal of Memory and Language*, 38, 419–439.
- Cooper, R. M. (1974). The control of eye fixation by the meaning of spoken language: A new methodology for the real-time investigation of speech perception, memory, and language processing. *Cognitive Psychology*, 6, 84–107.
- Cutler, A., & Otake, T. (2002). Rhythmic categories in spoken-word recognition. *Journal of Memory and Language*, 46, 296–322.
- Goldinger, S. D., Luce, P. A., & Pisoni, D. B. (1989). Priming lexical neighbors of spoken words: Effects of competition and inhibition. *Journal of memory and language*, 28(5), 501-518.
- Grice, H. P. (1970). *Logic and conversation* (pp. 41-58). na.
- Heeringa, W. (2004). Measuring dialect pronunciation differences using Levenshtein distance. Ph.D. dissertation, University of Groningen.
- Huetting, F., Rommers, J., & Meyer, A. S. (2011). Using the visual world paradigm to study language processing: A review and critical evaluation. *Acta psychologica*, 137(2), 151-171.
- Huetting, F., & McQueen, J. M. (2007). The tug of war between phonological, semantic and shape information in language-mediated visual search. *Journal of Memory and Language*, 57(4), 460-482.
- Kawahara, S., & Shinohara, K. (2009). The role of psychoacoustic similarity in Japanese puns: A

- corpus study. *Journal of Linguistics*, 45, 111–138.
- Kerkhof, N. (2015). Woordspelingen van hier tot Tokyo: Een onderzoek naar de verschillen in structuur en context tussen Nederlandse en Japanse woordspelingen.
- Kondrak, G., & Sherif, T. (2006). Evaluation of several phonetic similarity algorithms on the task of cognate identification. Proceedings of the Workshop on Linguistic Distances Sydney: Association of Computational Linguistics, pp. 43–50. [ACL Anthology Network, archive at <http://aclweb.org/anthology-new/>.]
- Lagerquist, L. M. (1980). Linguistic evidence from paranomasia. In Papers from the 16th Regional Meeting, Chicago Linguistic Society (pp. 185–191). Chicago, IL: Chicago Linguistic Society.
- Levelt, W. J., Roelofs, A., & Meyer, A. S. (1999). A theory of lexical access in speech production. *Behavioral and brain sciences*, 22(01), 1-38.
- Levenshtein, V. I. (1966). Binary codes capable of correcting deletions, insertions and reversals. *Soviet Physics Doklady*, 10 (8), 707–710. [Russian original (1965) in *Doklady Akademii Nauk SSSR*, 163 (4), 845–848.]
- Liberman, A. M., Cooper, F. S., Shankweiler, D. P., & Studdert-Kennedy, M. (1967). Perception of the speech code. *Psychological review*, 74(6), 431.
- Luce, P. A., & Pisoni, D. B. (1998). Recognizing spoken words: The neighborhood activation model. *Ear and hearing*, 19(1), 1.
- Luce, P. A. (1986). Neighborhoods of Words in the Mental Lexicon. Research on Speech Perception. Technical Report No. 6.
- Lupker, S. J., & Colombo, L. (1994). Inhibitory effects in form priming: Evaluating a phonological competition explanation. *Journal of Experimental Psychology: Human Perception and Performance*, 20(2), 437.
- Magnuson, J. S., Tanenhaus, M. K., Aslin, R. N., & Dahan, D. (2003). The time course of spoken word learning and recognition: studies with artificial lexicons. *Journal of Experimental Psychology: General*, 132(2), 202.
- McQueen, J. M., Dahan, D., & Cutler, A. (2003). Continuity and gradedness in speech processing. *Phonetics and phonology in language comprehension and production: Differences and similarities*, 39-78.
- Norris, D. G., McQueen, L. M., & Cutler, A. (1995). Competition and segmentation in spoken word recognition. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21, 1209-1228.
- Otake, T., & Cutler, A. (2013). Lexical selection in action: Evidence from spontaneous punning. *Language and speech*, 56(4), 555-573.
- Otake, T. (2010). *Dajare* is more flexible than puns: Evidence from word play in Japanese. *Journal of the Phonetic Society of Japan*, 14, 76–85.
- Otake, T., & Cutler, A. (2001). Recognition of (almost) spoken words: Evidence from word play in Japanese. In P. Dalsgaard, B. Lindberg, H. Benner & Z.-H. Tan (Eds.), *Proceedings of EUROSPEECH 2001*, Aalborg (pp. 465–468). Aalborg, Denmark: Kommunik Grafiske Løsninger.
- Schepens, J., Dijkstra, T., & Grootjen, F. (2012). Distributions of cognates in Europe as based on Levenshtein distance. *Bilingualism: Language and Cognition*, 15(01), 157-166.
- Vitevitch, M. S., & Luce, P. A. (1999). Probabilistic phonotactics and neighborhood activation in

- spoken word recognition. *Journal of Memory and Language*, 40(3), 374-408.
- Vitevitch, M. S., & Luce, P. A. (1998). When words compete: Levels of processing in perception of spoken words. *Psychological science*, 9(4), 325-329.
- Vroomen, J., & de Gelder, B. (1997). Activation of embedded words in spoken word recognition. *Journal of Experimental Psychology: Human Perception and Performance*, 23(3), 710.
- Vroomen, L., & de Gelder, B. (1995a). Lexical inhibition in spoken word recognition. Proceedings of the Fourth European Conference on Speech Communication and Technology, Madrid, Spain, 1711-1714.
- Vroomen, L., & de Gelder, B. (1995b). Metrical segmentation and lexical inhibition in spoken word recognition. *Journal of Experimental Psychology: Human Perception and Performance*, 21, 98-108.
- Yarkoni, T., Balota, D., & Yap, M. (2008). Moving beyond Coltheart's N: A new measure of orthographic similarity. *Psychonomic Bulletin & Review*, 15 (5), 971– 979.

Bijlagen

Bijlage 1: Gebruikte inbeddingen en mutaties voor dit onderzoek

Inbeddingen

Bronwoord	Doelwoord
beer	berenplaat
Blinde	blindlegger
Last	lastig
Zweedse	zweet
Chris	Christmas
Beks	bekt
Montell	mond
Teske	task
Jeline	Jelinek
monster	monsterhit
shower	sjouwt
Sue	sushi
Lakers	leek
non	non-actief
logistiek	logisch
Iman	iemand
Paaspop	passen
Yara	jaren
Kos	kosten
Kaiser	kaiserbroodjes
Dornan	doorn
Bittergarnituur	bittere
sandalen	sandalig
pretzel	pret
beenham	beenhamvraag
Dominique	dominicaanse
Genemuiden	geen
Sil	silly

beren	beregoed
app	epic
kruis	kruispunt
Valtifest	valti
marsepein	marsepeinboompitten
staart	staartjes
Bade	Badedas
brillen	brillenkoker
Regelink	regel
top	toplijst
lakje	slakje
Sandé	zand
Lars	Larstiger
neerslag	neerslachtig
waggelder	waggelt
Ja	ja-woord
harten	hard
feest	feestplaten
kat	katten
paassfeer	paastakken
genomineerd	genoom
dracula	draak
uitsmijter	uitsmijt
huidkleurig	huid
facturen	factureluurs
rekstok	rek
giet daar	gitaaroverlast
Verkoelen	verkoeldown
bliebjes	blieblitheek
handknak	hand
Sziget	zie
vergaderzaal	vergadert
spanband	spannend

op	opstand
Segett	zeg ik
eighties	ezel
bril	briljant
Knopfler	knoflook
monden	mondig
hout hakken	hakt
Han	handig
belastingdienst	belastingvrije
plankenkoorts	plank
Raider	rede
te	tegoedbon
Disneyland	dis nie
festijn	festival
onderbroeken	onder
kedendag	ledematendag
Tom	Tomtom
trommelvliezen	trommel te verliezen
sleedoorn	doorn
actief	radioactief
rik	viezerik
beer	masturbeer
kommer	komkommer
loftrompet	witloftrompet
corr	hardcore
sleehak	arresleehak
standaard	fietsstandaard
cent	Vincent
pijn	marsepein
bek	spaubek
mary	nachtmerrie
bucketlist	icebucketlist
dendrofiel	rododendrofiel

ballen	voetballen
webbe	website
pizzapuntjes	punt
vieren	klavieren
bakt	pottenbakt
zoen	Blaudzun
lak	Polak
druk	drugs
gieren	Schieren
rijst	parijst
lak	haarlak
Bakermat &	
Watermat	matten
Barend	schrikbarend
deur	ambassadeur
Eldom	Rekdom
ballen	pingpongballen
no	Renault
Haim	geheim
zak	doedelzak
nee	Chardonnay
bakker	koekenbakker
defector	tor
Rian	riante
boontje	Koffieboontje
transporteren	Sport
Roel	ruet
teen	meteen
ring	tering
Sef	beseft
taal	instrumentaal
loot	malloot
hamburger	Burger

Beckham	Breedbackham
gifje	vergif
stoter	betonstoter
logisch	biologisch
Ta fête	estafette
theezakje	satézakjes
saté	hupsaté
paracetamol	mol
Oosterschelde	schelden
koude	verkouden
suikerspinnen	spinnen
keys	uitkiezen
bade	zonnebaden
vlak	zandvlakte
Messi	Stanleymessi
Becky	douchebecky
Trainor	hometrainer
paal	mijlpaal
domtoren	ekdomtoren
erger	Kerger
Jet	pakket
tijd	zomertijd
selfiestick	stick
been	een
Geerling	centrifugeerling
Angenent	Henk
lachen	glimlach
Erk	vingerwerk
spinnende	vogelspin
instrumentaal	mentaal
riptide	sparerib tijd
Snijder	eiersnijder
mei	ei

haarbanden	zijbanden
leun in	rugleuning
Einden	ondermijnden
Wieman	relikwiman
spoorwegen	aanwegen
broek	korte broek
berg	hooiberg
pop	op
Kees	Irakees
hoek	geodriehoek
plaatsen	lamphangplaatsen
roquefort	voor
ring	tering
Ommen	goudviskommen
therapeute	neuspeuteren
Sheperds	paard
jong	hondenjong
vos	sloddervos
Lodder	Flodder
Geominie	iniminie
Raboeffie	boeffie
haan	sprinkhaan
Kramer	Marskramer
Soer	amour
Kahn	gedaan
Schipper	binnenvaartschipper
Pietzel	kanariepietzel
helpen	verhelpen
eronder	onder
Roelofsen	Letlofroelofsen
Marnix	niks
Maarten	Sint Maarten
onderbroeken	broek

Afrojack	Afrotrosjack
Hazenberg	paashazenberg
Geerling	corrigeerling

Mutaties

Bronwoord	Doelwoord
benen strekken	penis trekken
prosecco	antisecco
beladen	ballade
puik	pruik
struik	kruik
rolmaat	holmaat
dankbare	drankbare
denk	tank
saving	shaving
feestje	vestje
handen	honden
Ibiza	eat pizza
seksleven	snacksleven
zin	gin
Pharall	viral
koeien	queen
passe-partouts	Paas-partouts
hopen	heupen
grey	gay
Santana	Sultana
slavink	afvinken
Keukenhof	kuikenhof
Eldom	Eggdom
past	plast
beter	Peter
bussen	blussen

reading	redding
Baker	bakker
vis	pis
kralen	koraal
Yoshi	sushi
zangeres	zwangeres
surfen	smurfen
fall	vol
paarden	pardon
Raaijmakers	Haaijmakers
in bloom	in boom
goed	bloed
Ekka	lekker (lekka)
Shoupet	toupet
straks	Kwaks
vak	fack
wijn	chagrijn
save	shave
te vet	ta fête
Veneverklo	jeneverklo
Dree	Andre
holy	folie
cilinder	zie linder
brillen	billen
drie	die
ontarmd	omarmd
karaat	paraat
Ferry	very
suède	zwued
to wreck	trek
huis en tuin	gruis en puin
gek	track
zangeres	zwangeres

Nash	zash
ovaal	oh faal
Vaasen	fase
Tunstall	tuinstoel
Vriezeveen	schizofreen
tatoeëren	tutoyeren
fagot	van god
hapert	haperen
Swift	swiffer
as	Arcen
Murs	meurt
lokken	lokkertje
lasser	lastig
Zweden	zweet
leven	lever
Zegger	zegt
Barend	barond
Vermaat	vermaakt
Alanis	ellende
draak	drank
Oetan	oetang
even	Eva
gillen	giller
Mensa	mensen
opsteken	opsteker
loef	Louvre
remmers	remmen
Kosters	kosten
logo	logisch
populair	populier
Bagoo	toe
Kosters	kosten
gefeliciteerd	gefelicitaart

Wim	win
	Evert
even afleiding	afleiding
Berger	Bergen
sign	saaï
june	juni
Little	Lidl
blender	blunder
pop	plop
Sheppard	sheppaard
lot	lid
genageld	genaveld
Donners	Donders
Souberg	Soepberg
stelselmatig	stencilmatig
Brugman	Bregman
part	paard
lot	lid
kosten	korsten
uitgeteld	uitgejarreteld
people	piemel
konsten	kosten
Toeve loe	toedeloë
overzien	aubergine
tables	tepels
Zoab	soap
Sting	string
retriever	receiver
geknuffeld	getruffeld
pret	pritt
Hammer	hamer
dealer	deler
fris	fresh

kat	Kate
run	Ryan
pardon	paardon
lode dakje	leide dokje
balin	bowlen
kotsen	kosten
tables	tepels
buitenkeuken	buitenkuiken
Maui	mooi

Bijlage 2: Lijst van alle eigenschappen waarop de woordspelingen geannoteerd zijn (Kerkhof, 2015)

Naam kolom	Beschrijving	Annotatie
Nummer woordspeling	Nummer van de woordspeling, op volgorde van annotatie	Getal
Datum opname	Datum van de uitzending waarin de woordspeling zich bevindt	jjjjmmdd
Begintijd woordspeling	Tijdstip (uur, minuut, seconde) in de uitzending waar de voorafgaande zin aan de woordspeling begint	uu:mm:ss
Eindtijd woordspeling	Tijdstip (uur, minuut, seconde) in de uitzending waar de woordspeling eindigt	uu:mm:ss
Gesprekspartner	De gesprekspartner van de grappenmaker	Naam
Soort woordspeling	Wat voor soort woordspeling wordt er gemaakt	Homofoon, mutatie of inbedding
Bronwoord	Wat is het bronwoord	Woord
Doelwoord	Wat is het doelwoord	Woord
Wie zegt bronwoord	Door wie wordt het bronwoord uitgesproken	Naam
Wie zegt doelwoord	Door wie wordt het doelwoord uitgesproken	Naam
Hoe vaak bronwoord gezegd	Hoe vaak is het bronwoord (door verschillende)	Getal
Bestaand doelwoord	Is het doelwoord een echt bestaand woord? Bij twijfel, gekeken op www.vandale.nl	0 (nee) of 1 (ja)
Doelwoord voor bronwoord	Wordt het doelwoord eerder uitgesproken dan het bronwoord?	0 (nee) of 1 (ja)
Woordsoort bronwoord	Wat is de woordsoort van het bronwoord	Afkorting desbetreffende woordsoort
Woordsoort	Wat is de woordsoort van het doelwoord	Afkorting desbetreffende

doelwoord		woordsoort
Over woordgrens	Gaat het bron- en/of doelwoord over een woordgrens heen	0 (nee) of 1 (ja)
Lengte bronwoord	Wat is de lengte van het bronwoord in lettergrepen	Getal
Lengte doelwoord	Wat is de lengte van het doelwoord in lettergrepen	Getal
Klemtoon	Hebben bron- en doelwoord hetzelfde klemtoonpatroon	0 (nee) of 1 (ja)
Klemtoonpatroon bronwoord	Wat is het klemtoonpatroon van het bronwoord	1 = beklemtoonde lettergreep, * = onbeklemtoonde lettergreep

Klemtoonpatroon doelwoord	Wat is het klemtoonpatroon van het doelwoord	1 = beklemtoonde lettergreep, * = onbeklemtoonde lettergreep
Collocatie	Maakt de woordspeling deel uit van en collocatie	0 = nee, 1 = correcte collocatie, 2 = vervormde collocatie
Taal bronwoord	Wat is de taal van het bronwoord	Taal
Taal doelwoord	Wat is de taal van het doelwoord	Taal
Inflectie	Is de woordspeling vervoegd	0 (nee) of 1 (ja)
Voorafgaande zin	Wat is de zin voorafgaand aan de woordspeling	Zin
Zin met doelwoord	Wat is de zin met doelwoord erin	Zin
Opmerkingen	Eventuele bijzonderheden bij de woordspeling	Tekst
Extra bij mutaties:		
Aantal klanken veranderd	Hoeveel klanken van het bronwoord zitten niet in het doelwoord	Getal
Aantal nieuwe	Hoeveel klanken van het doelwoord zitten niet in	Getal

klanken	het bronwoord	
Vershil (nieuw-oud)	Wat is het verschil tussen de kolommen ‘aantal klanken veranderd’ en ‘aantal nieuwe klanken’	Getal
Klanken bronwoord	Welke klanken van het bronwoord zitten niet in het doelwoord	Letter(s)
Klanken doelwoord	Welke klanken van het doelwoord zitten niet in het bronwoord	Letter(s)
Positie in woord	Waar vindt de mutatie plaats?	Begin, midden of einde
Homofoon mogelijk	Is er een homofoon mogelijk? Geverifieerd op www.vandale.nl	0 (nee) of 1 (ja)
Extra bij inbeddingen		
Soort inbedding	Wat voor soort inbedding is het	Insertie of extractie
Positie overlap	Waar in het langere woord bevindt het overlappende deel van de inbedding zich	Begin, midden of einde
Homofoon mogelijk	Is er een homofoon mogelijk? Geverifieerd op www.vandale.nl	0 (nee) of 1 (ja)

Voor de combinatie categorie, moeten alle eigenschappen geannoteerd, dus zowel de eigenschappen die bij *extra bij mutaties* en *extra bij inbeddingen* staan.

Bijlage 3: Frequentie van type woordspeling Otake en Cutler (2013) en dit onderzoek

Uit interesse zijn naast de genoemde hoofdvraag en deelvragen de frequenties van homofonen, inbeddingen en mutaties uit dit onderzoek en de frequenties van de homofonen, inbeddingen en mutaties uit het onderzoek van Otake en Cutler (2013) onderzocht. Zowel in dit onderzoek als in het onderzoek van Otake & Cutler (2013) waren homofonen het meest frequent, dan inbeddingen en dan mutaties. De frequenties van de woordspelingen uit beide onderzoeken zijn weergegeven in Tabel 6. Met behulp van Chi-kwadraattoetsen werd het verschil in frequenties tussen homofonen, inbeddingen en mutaties voor beide onderzoeken onderzocht.

Uit deze Chi-kwadraattoetsen bleek dat de frequenties uit het onderzoek van Otake & Cutler (2013) niet significant van elkaar verschillen $\chi^2(2) = .80$, $p = .67$.

De frequenties uit dit onderzoek bleken echter wel significant van elkaar te verschillen $\chi^2(2) = 56.48$, $p < .001$. Homofonen kwamen significant vaker voor dan inbeddingen $\chi^2(1) = 16.34$, $p < .001$. Tevens kwamen homofonen significant vaker voor dan mutaties $\chi^2(1) = 54.68$, $p < .001$. Homofonen kwamen echter minder frequent voor dan zowel inbeddingen als mutaties, maar niet significant minder frequent $\chi^2(1) = 3.16$, $p = .08$.

Het verschil tussen de twee onderzoeken zou verklaard kunnen worden doordat woordspelingen in de twee onderzoeken niet op eenzelfde manier geanalyseerd werden. In dit onderzoek werden alle woordspelingen die door Gerard Ekdorf gemaakt werden, geanalyseerd (zowel bronwoord en doelwoord als alleen doelwoord door Gerard Ekdorf benoemd). In het onderzoek van Otake en Cutler (2013) werden de woordspelingen van Dokumumushi Sandayuu alleen geanalyseerd wanneer het bronwoord door iemand anders genoemd werd en niet door Dokumumushi Sandayuu zelf. Wanneer het bronwoord gebruikt wordt dat door de grappenmaker zelf geuit werd, kan het zo zijn dat de grappenmaker meer controle heeft over zijn woordspeling. Deze grotere controle kan resulteren in het frequenter gebruik van homofonen omdat dit de meest geslaagde woordspelingen zijn (Otake & Cutler, 2013). Alle woordspelingen gemaakt door Dokumumushi Sandayuu (dus ook de woordspelingen waarin hij zelf het bronwoord genoemd heeft) zouden onderzocht moeten worden om zo een betere vergelijking te kunnen maken tussen Gerard Ekdorf en Dokumumushi Sandayuu. Een andere manier is om de woordspelingen van Gerard Ekdorf waarin hij het bronwoord zelf benoemd, niet te analyseren. Door een van deze twee mogelijkheden toe te passen, zou er een betere vergelijking gemaakt kunnen worden tussen de frequenties van de woordspelingen uit de twee onderzoeken.

Tabel 6. *Frequentie van type woordspeling uit dit onderzoek en het onderzoek van Otake en Cutler (2013)*

	Type woordspeling		
	homofonen	inbeddingen	mutaties
Dit onderzoek	284	195	133
Otake & Cutler (2013)	105	96	93