

WORD LID VAN DE
VERENIGING VOOR
FONETISCHE WETENSCHAPPEN

Vul het formulier in en stuur het naar het onderstaande adres of email de gegevens naar Mirjam.Ernestus@mpi.nl.

achternaam:
voorletter(s) evt. titel:
afdeling/vakgroep:
postadres
werk- of priveadres:
postcode en plaats:
emailadres:

De contributie is 7 Euro / jaar

Aanmelding als lid bij:

Mirjam Ernestus
Max Planck Institute for Psycholinguistics
Postbus 310
6500 AH Nijmegen
tel: +31-24-3612970
email: Mirjam.Ernestus@mpi.nl

Voor meer informatie over de Vereniging voor Fonetische Wetenschappen:

Rob van Son
Leerstoelgroep Fonetische Wetenschappen
Universiteit van Amsterdam
Herengracht 338
1016 CG Amsterdam
Tel.: 020-5252195/020-5252183
Fax: 020-5252197
Email: R.J.J.H.vanSon@uva.nl
URL: <http://www.fon.hum.uva.nl/FonetischeVereniging/>

De dag van de *Fonetiek* 2003

Over lopend onderzoek naar spraak en spraaktechnologie
(<http://www.fon.hum.uva.nl/FonetischeVereniging/>)

Donderdag 18 december 2003 in de Sweelinckzaal, Drift 21 te Utrecht
Georganiseerd door de *Nederlandse Vereniging voor Fonetische Wetenschappen*
;toegang gratis!



Nederlandse
Vereniging
Voor
Fonetische
Wetenschappen

15:45 De automatische generatie van foneemtranscripties en segmentaties voor het Nederlands

Kris Demuyne, Tom Laureys, Dirk Van Compernelle, & Patrick Wambacq, K.U.Leuven

In deze presentatie beschrijven we de automatische generatie van foneemtranscripties en de bijhorende foneem- en woordsegmentaties zoals die gemaakt worden voor het Vlaamse deel van het Corpus Gesproken Nederlands. Eerst beschrijven we de automatische generatie van een netwerk van alternatieve foneemtranscripties op basis van de orthografie. Uit dit netwerk selecteert de automatische spraakherkenner het akoestisch best passende pad. Vervolgens behandelen we de gebruikte algoritmes voor het maken van woord- en foneemsegmentaties. We besluiten met een gedetailleerde analyse van de verschillen tussen enerzijds de manueel geproduceerde foneemtranscripties en woordoplijningen en anderzijds de resultaten van het automatische proces. Deze evaluatie gebeurt op de uiteenlopende componenten die binnen het CGN aanwezig zijn: van voorgelezen spraak over spontane conversaties tot telefoonspraak.

16:05 Automatische fonetische transcripties: Wat is al mogelijk?

Helmer Strik, Universiteit Nijmegen

Fonetische transcripties zijn nodig voor vele applicaties. Het is bekend dat manuele fonetische transcripties nadelen hebben, onder andere omdat het maken ervan tijdrovend en duur is. Een mogelijk alternatief zijn automatische fonetische transcripties. Maar, in hoeverre is dat nu al mogelijk? Deze vraag krijgen we de laatste tijd steeds vaker te horen. In deze presentatie zal ik proberen een antwoord te geven op deze vraag door een overzicht te presenteren van onderzoek dat al uitgevoerd is en lopend onderzoek.

16:25 Hoe valideer ik een spraakdatabase?

Henk van den Heuvel, SPEX/CLST, Universiteit Nijmegen

Valideren van spraakdatabases komt neer op het controleren van de kwaliteit van grote spraakbestanden aan de hand van een aantal criteria. Die criteria bestaan uit de specificaties van de database aangevuld met een aantal tolerantie marges die aangeven hoever van het ideaal mag worden afgeweken.

SPEX heeft een langdurige ervaring met het valideren spraakdatabases. Sinds 1996 valideert SPEX wij de spraakdatabases die in vele vaak door de EC gesponsorde projecten, geproduceerd worden. Verder zijn wij het officiële validatiecentrum van ELRA (European Language Resources Association). In deze bijdrage ga ik in op de aspecten van een database die gevalideerd worden en de procedures die daarbij gehanteerd worden. Hierbij zal gekeken worden hoe validaties in het begin werden uitgevoerd en hoe procedures en controles zich in de loop der tijden hebben ontwikkeld tot de hedendaagse stand van zaken.

Het woordje 'ik' in de titel heeft in eerste instantie betrekking op mijzelf, maar na afloop van de voordracht hoop ik dat ook de toehoorder geïnteresseerd is geraakt in het onderwerp en er zijn/haar persoonlijke voordeel mee kan doen.

PROGRAMMA

9:00 Ontvangst met koffie

9:30 Welkom

9:35-10:55 Ochtendsessie I (Voorzitter Bert Schouten)

9:35 Pleidooi voor articulatorische of a-fonetiek

Luc van Buuren, Linguavox, Bloemendaal

9:55 De perceptieve ontwikkeling van een Brits-Engels foneemcontrast bij volwassen Nederlanders

Willemijn Heeren, Universiteit Utrecht

10:15 Mutual intelligibility of Chinese, Dutch, and American speakers of English

Wang Hongyan and Vincent J. van Heuven

10:35 U verstaat (een klein beetje) Hongaars

Holger Mitterer, Max-Planck-Institut für Psycholinguistik

10:55 Koffiepauze

11:15-12:35 Ochtendsessie II (Voorzitter Mirjam Ernestus)

11:15 Sprekerherkenning door getuigen

Tina Cambier-Langeveld & Jos Vermeulen, Nederlands Forensisch Instituut

11:35 De Gooise r: chique maar toch irritant?

Renee van Bezooijen, Universiteit Nijmegen

11:55 "Ik was toch wel blij dat ik van mijn hobby mijn beroep kon maken"

De uitspraak van het possessivum *mijn* in het Standaardnederlands

Hanne Kloots, Steven Gillis & Marc Swerts, Universiteit Antwerpen

12:15 Demonstraties en Lunch

12.15 NeXTeNS: een nieuw open source tekst-naar-spraak systeem voor het Nederlands

Erwin Marsi & Joop Kerkhoff, Universiteit van Tilburg, Universiteit Nijmegen

12.30 Nederlandse (LVCSR) spraakherkenning/Information Retrieval

Roeland J. F. Ordelman, Universiteit van Twente

14:00-15:20 Middagsessie I (Voorzitter Rob van Son)

14:00 Vocaalreductie in monomorfematische woorden

Evie Coussé & Hanne Kloots, Universiteit Antwerpen

14:20 Prevoicing in het Nederlands

Petra van Alphen, Max-Planck-Institut für Psycholinguistik

14:40 Restructuring Rhythm Patterns

Maartje Schreuder & Dicky Gilbers, Universiteit van Groningen

15:00 Over het perceptieve belang van ritme en metrum

Hugo Quené, Universiteit Utrecht

15:20 Thee

15:45-16:45 Middagsessie II (Voorzitter Erwin Marsi)

15:45 De automatische generatie van foneemtranscripties en segmentaties voor het Nederlands

Kris Demuyne, Tom Laureys, Dirk Van Compernelle, & Patrick Wambacq, K.U.Leuven

16:05 Automatische fonetische transcripties: Wat is al mogelijk?

Helmer Strik, Universiteit Nijmegen

16:25 Hoe valideer ik een spraakdatabase?

Henk van den Heuvel, SPEX/CLST, Universiteit Nijmegen

16:45 Afsluiting

9:35 Pleidooi voor articulatorische of a-fonetiek

Luc van Buuren, Linguavox, Bloemendaal

De a-fonetiek ging kopje onder na de eerste bloeiperiode ± 2500 jaar geleden. Na een tweede bloeiperiode (± 1870-1945), dreigt ze opnieuw kopje onder te gaan. De Indiase benadering (articulatorische 'yoga', introspectie, auditieve observatie) strookt niet met westerse (empirische, instrumentele, visuele) voorkeuren, c.q. de a-fonetiek. Een opleiding tot a-foneticus schijnt niet langer te bestaan. Linguïsten en fonetici beschouwen nu de a-fonetiek als marginaal of irrelevant.

Practisch argument voor a-fonetiek. De *fantastische* vooruitgang sinds 1870 in de beschrijving van *alle* vocalisatie stelt ons in staat de uitspraak van talen (Nederlands, Engels...) nauwkeurig te beschrijven en/of te doceren aan allochthonen. Maar willen we dat wel?

Theoretisch argument. Taal is een netwerk van sociaal bepaalde vorm-betekenis eenheden (Saussure). In uw menselijk brein zitten die betekenissen en (fonologische) vormen als met elkaar verbonden neuro-cognitieve patronen, zintuigelijk –vooral auditief– *geleerd*. Om uw taaltekens weer *sociaal* te gebruiken moeten zij uw motoriek (vocalisatie, gebaren) activeren. Ergo: het verschijnsel taal behelst een vicieuze cirkel: fysiologie(hersenactiviteit +motoriek) ↔ vorm ↔ betekenis ↔ fysiologie... Ergo: a-fonetiek is een onmisbare component van taalkunde.

9:55 De perceptieve ontwikkeling van een Brits-Engels foneemcontrast bij volwassen Nederlanders

Willemijn Heeren, Universiteit Utrecht

Mijn promotieonderzoek richt zich op de vraag hoe de perceptie van een nieuw foneemcontrast zich ontwikkelt bij verschillende leeftijdsgroepen, nl. volwassenen en kinderen in de basisschoolleeftijd. We gaan hierbij uit van twee hypothesen. De eerste hypothese, Acquired Distinctiveness, stelt dat luisteraars verschillen binnen of tussen nieuwe categorieën aanvankelijk slecht horen. Door training leert de luisteraar de verschillen tussen klanken die verschillend worden gecategoriseerd. De tweede hypothese, Acquired Similarity, stelt dat de luisteraar verschillen binnen en tussen categorieën voor het leren goed kan onderscheiden. Door training blijft enkel het verschil tussen klanken die verschillend worden gecategoriseerd overeind.

Het experiment, waarvan ik de opzet en de voorlopige resultaten zal bespreken, volgt het leren van het Brits-Engelse contrast, / -s/, door volwassen Nederlanders. In een pretest-posttest design wordt de foneemontwikkeling als gevolg van training bekeken. De pre- en posttest bevatten spraak van één spreker. Spraak van vijf andere sprekers vormt het trainingsmateriaal. Deze variatie dwingt de luisteraar te abstraheren over sprekerverschillen en robuuste categorieën te vormen.

14:40 Restructuring Rhythm Patterns

Maartje Schreuder & Dicky Gilbers, Universiteit van Groningen

For this experiment we wondered whether the influence of a higher speech rate leads to adjustment of the phonological structure, as it does in music, or just to 'phonetic compression', i.e. shortening and merging of vowels and consonants, with preservation of the phonological structure. If the rhythmic structure is adjusted, this implies that every speech rate has its own register, in terms of Optimality Theory (Prince & Smolensky, 1993) its own ranking of constraints.

The allegro (fast) data were obtained by means of a multiple-choice quiz in which two subjects competed each other in answering simple questions as quickly as possible. Afterwards the subjects were asked to read the words in a sentence, at a moderate speech rate. The data were judged by five trained listeners, and were phonetically analysed in PRAAT.

The results showed a preference for restructured rhythms in fast speech. Particularly for the fastest speakers correspondence constraints prevailed in their andante (moderate) speech, whereas in allegro tempo markedness constraints dominated the correspondence ones.

15:00 Over het perceptieve belang van ritme en metrum

Hugo Quené, Universiteit Utrecht

De seriële volgorde van sterke en zwakke lettergrepen is van belang voor de spraakperceptie, althans in het Engels en het Nederlands. Een woord met beklemtoonde eerste lettergreep wordt sneller herkend dan een woord beginnend met een onbeklemtoonde lettergreep. Maar waardoor wordt dit effect nu eigenlijk veroorzaakt? Door de isochronie van beklemtoonde lettergrepen (ritme) in de hoorbare spraak? Of door de alternantie van sterke en zwakke lettergrepen (metrum)?

In een foneem-detectie-experiment heb ik deze ritmische en metrische factoren proberen te ontwarren. Luisteraars hoorden woordenlijstjes, waarin het doelwoord hetzij metrisch voorspelbaar was (zelfde patroon als voorgangers in lijstje), hetzij metrisch onvoorspelbaar. Bovendien was de 'timing' tussen woorden of ritmisch (isochronie van klemtonen), of a-ritmisch. De resulterende reactietijden laten duidelijk effect zien van ritme, maar niet van metrum. Dit resultaat suggereert dat de ritmische 'timing' van belang is voor de herkenning van gesproken woorden, en dat luisteraars deze ritmes gebruiken bij de herkenning van gesproken woorden.

14:00 Vocaalreductie in monomorfematische woorden

Evie Coussé & Hanne Kloots, Universiteit Antwerpen

In deze bijdrage presenteren we een onderzoek naar vocaalreductie in het Standaardnederlands aan de hand van data uit het Corpus Gesproken Nederlands. Onder vocaalreductie verstaan we het verkorten van een fonologisch lange klinker tot zijn korte pendant (pr[o]bleem > pr[]leem), het verdoffen van een volle klinker tot een sjwa (m[i]nuut/ > m[]nuut) en de volledige deletie van een klinker (Int[e]resse > int[]resse). In de fonetische en fonologische vakliteratuur zijn een aantal hypothesen geponeerd over vocaalreductie. De ontwikkeling van het Corpus Gesproken Nederlands geeft ons de mogelijkheid om een aantal van die stellingen te testen op een grote dataset.

Concreet hebben we van een subcorpus van monomorfematische woorden de brede fonetische transcriptie gealigneerd met een referentietranscriptie. Door beide transcripties te vergelijken, kunnen we variatie in de uitspraak op het spoor komen. Uit het onderzoek blijkt dat vocaalreductie beïnvloed wordt door fonologische factoren als vocaalkwaliteit, aard van de omringende consonanten, aantal syllaben, syllabestructuur, de relatieve positie van de klemtoon en taalexterne factoren als regio, spreekstijl en woordfrequentie.

14:20 Prevoicing in het Nederlands

Petra van Alphen, Max-Planck-Institut für Psycholinguistik

In het algemeen wordt aangenomen dat de Nederlandse plofklanken [b] en [d] aan het begin van een woord met prevoicing worden geproduceerd. Het eerste experiment van dit project laat echter zien dat stemhebbende plosieven regelmatig zonder prevoicing worden gerealiseerd. Er zijn verschillende factoren die de productie van prevoicing beïnvloeden. In het tweede experiment wordt onderzocht welke andere potentiële akoestische cues aanwezig zijn, die gebruikt kunnen worden bij de perceptie van het fonologische onderscheid tussen stemhebbende en stemloze plosieven. Er is gebruik gemaakt van een CART-analyse om te voorspellen welke van deze akoestische maten het meest betrouwbaar zijn. Tenslotte laat een perceptie-experiment zien welke van deze potentiële cues daadwerkelijk worden gebruikt door luisteraars. Hieruit blijkt dat prevoicing de primaire cue is voor het onderscheid tussen stemhebbende en stemloze plofklanken, ondanks het feit dat prevoicing in het Nederlands regelmatig afwezig is.

10:15 Mutual intelligibility of Chinese, Dutch, and American speakers of English

Wang Hongyan and Vincent J. van Heuven

Very little is known about the loss of intelligibility that is incurred by L2 speakers when they communicate with native L1 listeners. Even less is known about the differences in intelligibility among L2 English speakers of diverse national backgrounds, such as Chinese-accented speakers of English versus Dutch-accented speakers. The first aim of our study is to test the hypotheses that (i) Dutch English is more intelligible to native listeners of English than Chinese English, and (ii) both foreign-accented varieties are less intelligible to L1 English listeners than native English. Hypothesis (i) follows from a contrastive analysis of the sound systems of the languages involved, showing that Dutch and English are much more similar in their sound structure than Chinese and English. Our third hypothesis relates to the relative intelligibility of the three types of English for non-native listeners. Two hypotheses are plausible here: (iii) L1 English is always more intelligible to listeners of any nationality, since it optimally conforms to the norm the foreign speaker/listeners were taught to adhere to, or (iv) Dutch English is more intelligible to Dutch listeners, and Chinese English to Chinese listeners, as these varieties embody precisely the interference phenomena that the L2 speakers are used to.

We recorded a male and a female speaker of (American) English, of Dutch English and of Chinese English. Speakers were young adults, studying at the university level with no specialisation in English. Five types of English materials were recorded for each speaker: (1) **vowel test**: a list of word containing the 20 vowels in identical /hVd/ contexts, (2) **consonant test**: a list of nonsense words /aCa/ containing 24 intervocalic single consonants, (3) **cluster test**: a list of 20 CC or CCC clusters in /aCC(C)a/ clusters, (4) **SUS-test**: 30 Semantically Unpredictable Sentences with high-frequency words occurring in syntactically correct but semantically nonsense sentences, and (5) **SPIN test**: 50 short sentences, with a contextually predictable or unpredictable target word in final position. The entire set of materials was then presented in perceptual identification and recognition tests three groups of listeners belonging to the same population as the speakers.

For each test, hypotheses (i), (ii) and (iv) but not (iii) were supported. In our talk we will present the confusion structure in the vowel, consonant, and cluster data, and show how intelligibility at the sentence level can be predicted through regression analysis from the phoneme-identification results.

10:35 U verstaat (een klein beetje) Hongaars

Holger Mitterer, Max-Planck-Institut für Psycholinguistik

In vloeiende spraak komen taalspecifieke aanpassingen, zoals assimilatie, voor. Zo kan het woordje tuin met een /m/ worden uitgesproken. In hoeverre tuim als een voorbeeld van tuin herkend wordt is afhankelijk van de fonologische context, met name in hoeverre die assimilatie toestaat. Zo wordt tuim door Nederlandse luisteraars als tuin herkend in tuimbank maar niet in *tuinstoel. De vraag is nu in hoeverre deze contextsensitiviteit een gevolg is van het leren van assimilatie regels, d.w.z. Nederlanders hebben geleerd dat een /m/ voor een /b/ een /n/ kan zijn? Om dit te onderzoeken hebben wij Portugese proefpersonen met Nederlandse assimilaties en Nederlandse proefpersonen met Hongaarse assimilaties geconfronteerd. Uit de resultaten blijkt dat de contextsensitiviteit bij het herkenning van assimilaties maar ten dele het gevolg is van taalspecifieke ervaring met assimilatieprocessen. Zou u dus Hongaars willen leren, zou de herkenning van geassimileerde vormen bij hoge uitzondering geen problemen opwerpen.

11:15 Sprekerherkenning door getuigen

Tina Cambier-Langeveld & Jos Vermeulen, Nederlands Forensisch Instituut

Naast het verrichten van vergelijkend spraakonderzoek, d.w.z. sprekerherkenning door deskundigen, wordt het Nederlands Forensisch Instituut (NFI) een enkele keer ook gevraagd een betrouwbare vorm van sprekerherkenning door getuigen in elkaar te zetten. Als de zaak voldoet aan een aantal criteria, wordt een ‘voice line-up’ geconstrueerd, waarbij de stem van een verdachte in een rijtje van soortgelijke stemmen wordt gezet en aan de getuige wordt gevraagd of hij/zij één van de stemmen herkent als de stem van de dader.

Om de waarde van de uitkomst van een dergelijke ‘voice line-up’ zo goed mogelijk in te kunnen schatten, is het belangrijk dat de line-up aan bepaalde eisen voldoet. Eén van de eisen die gesteld wordt is dat alle stemmen in de line-up moeten voldoen aan de beschrijving die de getuige geeft van de stem van de dader. Deze beschrijving is echter over het algemeen weinig specifiek. In voorkomende gevallen wordt altijd met een vragenlijst gewerkt, waarbij de getuige enige terminologie krijgt aangereikt. Een voorbeeld van zo’n vragenlijst zal tijdens de presentatie worden weergegeven. Deze is echter redelijk arbitrair tot stand gekomen. Het vinden van een eenduidige terminologie voor het beschrijven van stemmen blijft een zeer lastige zaak.

11:35 De Gooise r: chique maar toch irritant?

Renee van Bezooijen, Universiteit Nijmegen

De Gooise r - dat wil zeggen de approximantische realisatie van de /r/ in postvocale positie - lijkt zich in een rap tempo in het Nederlands te verspreiden. In Haarlem is hij onder kinderen nu de enige postvocale realisatie van de /r/, maar ook in Nijmegen heeft de helft van de kinderen hem al. Ook in de media tref je hem veelvuldig aan: tweederde van de televisiepresentatoren gebruikt hem op z’n minst af en toe. Vindt men de Gooise r dan zo veel aantrekkelijker dan de tongpunt-r en de huig-r? En wat straalt de Gooise r dan uit, met wat voor persoonlijkheid worden de verschillende varianten geassocieerd? Wat weten mensen van het voorkomen van verschillende r-varianten in Nederland? Deze vragen stonden centraal in een evaluatie-onderzoek dat ik heb uitgevoerd in vier regio’s: (de gebieden rondom) Hilversum, Nijmegen, Geleen en Leeuwarden. Het aantal luisteraars per plaats lag tussen de 30 en 40 personen, verdeeld over mannen en vrouwen en twee leeftijden. Er werd gebruik gemaakt van de matched guise techniek, waarbij dezelfde tekst door dezelfde spreker met verschillende combinaties van r-en werd ingesproken. Ik presenteer in mijn lezing de resultaten.

11:55 "Ik was toch wel blij dat ik van mijn hobby mijn beroep kon maken"

De uitspraak van het possessivum *mijn* in het Standaardnederlands

Hanne Kloots, Steven Gillis & Marc Swerts, Universiteit Antwerpen

In het kader van het VNC-project Variatie in de uitspraak van het Standaardnederlands werd een sociolinguistisch interview afgenomen van 80 Vlaamse en 80 Nederlandse leraren Nederlands. De steekproef was gestratificeerd naar regio (4 regio’s in Vlaanderen, 4 in Nederland), sekse (evenveel mannen als vrouwen) en leeftijd (de helft van de sprekers is geboren voor 1955, de andere helft na 1960). De spontane spraak die in het kader van dit project verzameld werd, vormt momenteel de basis voor onderzoek naar reductieverschijnselen in de standaardtaal. In deze presentatie brengen we verslag uit van een studie naar de uitspraak van het possessivum *mijn*. De 160 gesprekken bevatten in totaal 1253 realisaties van *mijn*. De stimuli werden gescoord door drie beoordelaars via een internetapplicatie. In de literatuur worden doorgaans twee uitspraakvarianten onderscheiden: de ‘volle’ vorm *mijn* en de ‘doffe’ vorm *m’n*. We gaan na of dit inderdaad de enige varianten zijn die in ons corpus voorkomen. Vervolgens onderzoeken we de invloed van de variabelen land, leeftijd en sekse, en we besteden daarbij ook aandacht aan factoren als aanwezigheid van klemtoon, toepassing van taalnormen en regionale herkomst van de spreker.

12.15 NeXTeNS: een nieuw open source tekst-naar-spraak systeem voor het Nederlands

Erwin Marsi & Joop Kerkhoff, Universiteit van Tilburg, Universiteit Nijmegen

NeXTeNS staat voor 'Nederlandse Extensie voor Tekst naar Spraak', en is een project dat tot doel heeft om een modern tekst-naar-spraak systeem te ontwikkelen voor onderwijs- en onderzoeksdoeleinden. Het systeem draait onder verschillende besturingssystemen (MS WIndows en Linux), de programmacode is vrij beschikbaar (open source), en het is gratis te verkrijgen en te gebruiken. Eerst beargumenteren we waarom er behoefte is zo’n nieuw systeem. Vervolgens presenteren we in het kort de doelstellingen, de deelnemers, de ontwikkelingsstrategie (nl. zoveel mogelijk gebruik maken van bestaande voorzieningen en programma's), en de architectuur. In de rest van dit praatje zullen we de nadruk leggen op het praktische perspectief: wat kunnen gebruikers met NeXTeNS doen? Tevens zullen we de grafische gebruikersinterface bespreken. Ter afsluiting zullen we een aantal voorbeelden van synthetische spraak laten horen.

12.30 Nederlandse (LVCSR) spraakherkenning/Information Retrieval

Roeland J. F. Ordelman, Universiteit van Twente

Aan de hand van twee demonstraties wil ik laten zien hoe Nederlandse spraakherkenning ingezet kan worden voor spraak-gebaseerde retrieval. De eerste demo laat zien hoe door middel van "alignment" optimaal gebruik kan worden gemaakt van al aanwezige, niet geheel overeenkomende transcripties van spraak (zoals notulen van vergaderingen) voor het zoeken in audio/multimedia bestanden. De tweede demo toont de op de Universiteit Twente voor het Nederlands ontwikkelde spraakherkenner in actie in het nieuws domein en geeft een idee van de huidige kwaliteit van de herkenning (met groot vocabulaire) en hoe de herkenningresultaten kunnen worden gebruikt voor het verkrijgen van aan het audiofragment gerelateerde informatie (koppeling nieuwsuitzending aan krantenmateriaal).