

DIXIT

TIJDSCHRIFT OVER TAAL- EN SPRAAKTECHNOLOGIE



TAAL- EN
SPRAAKTECHNOLOGIE

IN BEDRIJF



INCLUSIEF VERSLAG VAN HET CONGRES TAAL IN BEDRIJF 2011

INHOUD

- Algemeen	2
- Voorwoord	3
- Inleiding	6
- 10 jaar NOTaS	6
- De markt is klaar voor spraaktechnologie	8

Taal in Bedrijf

- Automatische spraakherkenning voor het taalonderwijs en de zorg	10
- AEGON Voiceportal	12
- Met spraaktechnologie telefonische contacten analyseren	14
- Spraakherkenning spoort criminelen op	16
- Contextgebaseerde spellingcorrectie met Valkuil.net	17
- Hoe kan de 'onderzoekende zoeker' beter ondersteund worden?	18
- In de frontlinie op Facebook en Twitter	20
- De digitale rechercheur	22
- Stemmen maken	24
- Watson voor het Nederlands	26
- Taaltechnologie elimineert obstakels voor recruitment van talent	28
- Een online reputatiemonitor	30
- Naam in Bedrijf: datakwaliteit van relatiegegevens	32
- Wordt vertaler een knelpuntberoep?	34
- Spraaktechnologie voor inclusive design	36
- Taal- en spraaktechnologie en navigatie: een gelukkig huwelijk	38
- Taal onderweg	40
- Onderzoekers die dit artikel lezen, lezen ook...	42
- Een spraakgestuurde sprekende applicatie voor materieelonderhoud	43
- "BATAVO was conservatief en braaf"	44

Congres Taal in Bedrijf 2011

- Taal in Bedrijf, een geslaagd evenement	46
- Technologie voorbij de menselijke intelligentie	47
- Verslag Sociale Media Taal in Bedrijf 2011	48
- Taal in Bedrijf	49

En verder

- Colofon	2
- Directory	51

COLOFON

DIXIT: Tijdschrift over toegepaste taal- en spraaktechnologie – 8e jaargang, editie TST in Bedrijf. DIXIT is een uitgave van Stichting NOTaS, Postbus 31070, 6503 CB Nijmegen. Tel. 024-352 88 88 – E-mail info@notas.nl – www.notas.nl

Redactieadres: Stichting NOTaS, Postbus 31070, 6503 CB Nijmegen **Redactie:** Arjan van Hessen: hessen@cs.utwente.nl, Henk van den Heuvel: h.vandenheuvel@let.ru.nl, Remco van Veenendaal: remco.vanveenendaal@inl.nl, Leonoor van der Beek: vdbeek@gmail.com, Anne Wijnen: info@notas.nl Hoofredactie themanummer: René Collier: rene.collier@chello.nl

Advertenties: Stichting NOTaS – Anne Wijnen: info@notas.nl, 024-352 88 88 **Abonnementen:** Voor een gratis abonnement dient u zich te wenden tot een van de NOTaS-deelnemers (zie www.notas.nl) **Druk:** Leonard Nijmegen bv **Verantwoording:** DIXIT is een uitgave van Stichting NOTaS. Overname van de artikelen is alleen toegestaan met bronvermelding en na toestemming van Stichting NOTaS. Stichting NOTaS en de bij deze uitgave betrokken redactie en medewerkers aanvaarden geen aansprakelijkheid voor mogelijke gevolgen die zouden kunnen voortvloeien uit het gebruik van de in deze uitgave opgenomen informatie.

Voorwoord

Zoekt en gij zult ... steeds beter vinden

Dankzij taal- en spraaktechnologie (TST)!

Weet u nog hoe het zoekproces vroeger ging? Wat vond u erger "No hits, check your spelling" of juist 55.000.000 hits op willekeurige volgorde met de impliciete boodschap "zoek maar zelf verder"? Zonder het te beseffen maken we allemaal in meer of mindere mate gebruik van TST. Niet alleen bij het zoeken op internet of spraakgestuurd nummerkiezen, maar wat dacht u van de spraaksynthese in uw navigatiesysteem of het omroepsysteem op het vliegveld of bij het chatten met een Avatar? En er zijn maar weinig mensen die géén gebruik maken van de spellingscontrole of die niet af en toe een online woordenboek raadplegen. TST is daarom belangrijk, ook voor het Nederlands. Dit was een van de drijfveren voor de inrichting van het STEVIN-programma, waarvoor de financiën door de Nederlandse en Vlaamse overheid bijeen werden gebracht. Gelukkig werd een significant deel van het geld geoormerkt voor samenwerking van academies en bedrijfsleven; iets waarvoor NOTaS de Nederlandse Taalunie en de financiers van STEVIN nog steeds zeer erkentelijk is. Dankzij deze financiële en organisatorische ondersteuning is er een zeer groot aantal zeer succesvolle projecten gerealiseerd waarmee de samenwerkende partijen zowel de kennis als het gebruik van TST sterk hebben vergroot.

Zoveel vooruitgang en toch nog een lange weg te gaan

Als vertaler en redacteur weet ik uit persoonlijke ervaring dat automatisch vertalen nog een lange weg heeft te gaan. Ook andere TST-toepassingen zijn hoewel al zeer bruikbaar, nog lang niet uitontwikkeld. De herkenning van gewone 'sloppy' spraak kan nog beter, zoeken in de enorme hoeveelheid informatie moet 'semantisch', automatische correcties van oude, ge-OCR-de documenten moet beter, net als de vindbaarheid van al die digitale en gedigitaliseerde informatie. Daarom steunt NOTaS van harte nieuwe initiatieven op het gebied van taal-, spraak- en ontsluitingstechnologie zoals CLARIN en CLARIAH.

10 jaar NOTaS

Dit jaar viert NOTaS haar 10-jarig jubileum. Sinds onze oprichting in 2001 is er veel gebeurd. Onze eerste voorzitter, Geert Kobus, helaas op een veel te jonge leeftijd overleden, zou onder de indruk zijn van de prestaties van onze deelnemers. Samen met het NOTaS-bestuur deel ik zijn visie dat samenwerking tussen kennisinstellingen, bedrijven en concullega's een enorme toegevoegde waarde creëert en dus een belangrijke bijdrage levert aan de verdere ontwikkeling en verbetering van de TST: de zo onzichtbare doch onmisbare technologie in onze informatie samenleving.

Als NOTaS kijken wij uit naar ons volgende jubileum en zullen u intussen via DIXIT, onze website www.notas.nl en onze LinkedIn-groep blijven informeren over de TST-ontwikkelingen en hoe zij uw werk- en privéleven makkelijker kunnen maken.

Een nieuwe DIXIT

Wij zijn de vele schrijvers van deze DIXIT dankbaar! De meest vooraanstaande onderzoekers en kennisdragers binnen onze industrie in Nederland en Vlaanderen hebben aan deze speciale editie een bijdrage geleverd. Onze gastredacteur, René Collier, heeft o.a. als voorzitter van het IOP Mens-Machine Interactie zijn sporen verdiend. We zijn hem dankbaar voor inspirerende inleiding en zijn kritisch oog op het geheel. Namens NOTaS en de DIXIT-redactie wens ik u een leerzame en inspirerende blik in de dynamische wereld van TST 'na STEVIN'.

Debbie Kenyon-Jackson, voorzitter stichting NOTaS

3000 jaar taal- en spraak-technologie

Taal en spraak worden meestal in één adem genoemd omdat ze zo nauw verwant zijn. Spraak is immers de natuurlijke manifestatie van taal. Spraak is echter een vluchtig medium met beperkte reikwijdte: het geluid sterft meteen uit en het draagt niet ver. Daarom worden al sinds mensenheugenis pogingen gedaan om die beperkingen op te heffen met de beschikbare technische middelen.

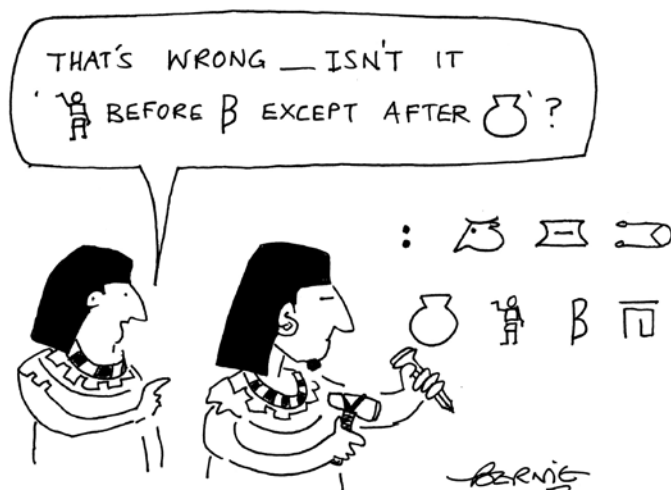
Van spijkerschrift naar spraakherkenning en FM radio

De uitvinding van het (spijker)schrift door de Sumeriërs, zo'n 3000 jaar geleden, is voor mij de belangrijkste spraaktechnologische uitvinding aller tijden: een vluchtige gesproken boodschap kon met een puntige stylus voor eeuwig worden gegraveerd in een kleitabletje! Je moest slechts voor ieder concept een teken afspreken. Een paar eeuwen later kwamen de Egyptenaren met een nog dieper en praktischer inzicht: de eenheid voor de registratie van gesproken taal is niet het concept, dus ruwweg het woord, maar de individuele klank. Spraak bestaat uit slechts een paar tientallen distinctieve geluiden. Als je voor elk daarvan een grafisch teken afspreekt, kun je gesproken boodschappen met een handvol symbolen veel compacter registreren: met k,a,s,t, schrijf je 'kast' enz. Geschreven spraak kun je bewaren, maar ook naar een verre bestemming laten transporteren. Als je echter simultaan diverse ontvangers wil bereiken, moet je het tabletje (of de papyrusrol) met de hand kopiëren. Die oplossing bleef gelden tot de late Middeleeuwen en werd pas opgeheven door de uitvinding van de drukpers. De schrijfmachine, als een soort persoonlijke

drukpers, was een volgende grote mijlpaal. En nog een paar eeuwen later maakten de PC en internet tekstverwerking en email mogelijk. Om niet zelf te hoeven typen werd in de jaren vijftig van de vorige eeuw gewerkt aan een 'voice-activated typewriter' als de gedroomde applicatie van automatische spraakherkenning. Met de rechtstreekse opslag van het spraakgeluid werden ook grote vorderingen gemaakt, van het kleitabletje naar de grammofoon, de magnetische en later de digitale recorder. En als opvolger van de tamtam zorgden de radio en de telefonie voor een wereldomspannende distributie van het geluid.

Uit de opsomming van een paar mijlpalen blijkt dat taal- en spraaktechnologie hun oorsprong niet hebben in de recente ontwikkeling van de elektronica of de informatica. Iedere generatie heeft er met de toen beschikbare technische middelen aan bijgedragen: het functionele verschil tussen het kleitabletje en het scherm van een PDA is slechts gradueel; ze hebben allebei een 'pen' nodig... Het meest opvallend is de snelheid waarmee innovatie toeneemt: de mijlpalen liggen geen eeuwen, maar slechts jaren of maanden uit elkaar.

René Collier
Voorzitter IOP
Mens-Machine
Interactie



Illustratie: www.cartoonstock.com

Kennis vs. kunde?

Het baanbrekende inzicht dat het aantal distinctieve klanken ('fonemen') per taal zeer beperkt is werd geformaliseerd in de structurele taalwetenschap, die ca. 1920 de 'fonologie' introduceerde. Latere ontwikkelingen in de linguïstiek en de fonetiek zijn maar voor een klein deel van invloed geweest op taal- en spraaktechnologie. Groeiend begrip van hoe het proces van spreken en verstaan in zijn werk gaat is nauwelijks terechtgekomen in moderne algoritmes voor automatische spraaksynthese, resp. spraakherkenning. Ook de linguïstische competentie van een professionele vertaler of een meertalige spreker worden niet nagebootst in een vertaalprogramma.

Er zijn in het spraakdomein vroege pogingen tot fonetisch gemotiveerde spraaksynthese. Met zijn mechanische spraakmachine bouwde Von Kempelen (1791) het productiesysteem na, van de longen tot de lippen. Het werd echter vooral gezien als een curiosum, ook omdat de uitvinder al eens gefraudeerd had met een mechanische schaa-machine waarin een dwerg verborgen zat... De eerste heuse elektrische synthesizer, de Voder -voice coder- van IBM (1939), was gebaseerd op akoestische principes. Het resultaat klonk niet overtuigend. De oplossing voor spraaksynthese werd rond 1970 uiteindelijk gezocht in het zorgvuldig aan elkaar plakken van stukjes van tevoren opgenomen spraak. Concatenatie of 'unit selection' is als synthesesetechniek minder triviaal dan het lijkt. Het is een echte uitdaging voor o.m. signaalbewerkers, informatici en andere technologiën. De aan elkaar gekoppelde eenheden moeten op zinsniveau door fonetici nog voorzien worden van accentuering, spraakmelodie en ritme. Om dat goed te doen is eerst nog een syntactische en semantische analyse nodig. Dus blijven er nog uitdagingen genoeg voor taalkundigen, in multidisciplinaire samenwerking met ingenieurs.

Voor spraakherkenning waren er geen vroege mechanische oplossingen. Een algoritmische benadering, gebaseerd op inzichten uit de spraakperceptie, kon alleen maar op een computer geïmplementeerd worden. Dat gebeurde voor het eerst in 1952, op de Bell Labs in de Verenigde Staten. Maar ook bij herkenning werden 'regels' al snel door 'waarschijnlijkheden' vervangen en won de ingenieur het van de linguïst. Berucht is de uitspraak (ca. 1970) van Fred Jelinek, toen toponderzoeker bij IBM: "Every time I fire a linguist, the performance of the speech recognizer goes up".

Automatisch vertalen als een gokje?

De samenwerking tussen Georgetown University en IBM leidde in 1954 tot het eerste automatische vertaalsysteem van 60 Russische naar Engelse zinnen. Het nieuws haalde de voorpagina van de New York Times. Het enthousiasme was groot: na 5 à 10 jaar zouden alle vertaalproblemen opgelost zijn. Ook in dit geval moest een taalkundige aanpak na een paar decennia wijken voor een statistische. De vertaalapplicatie van Google bijvoorbeeld, werkt nu met een database van 200 miljard woorden. De Nederlandse Google-vertaling van het bovenstaande citaat van Jelinek is: "Elke keer als ik brand een linguïst, de prestaties van de spraakherkenner gaat omhoog". Niet perfect (want de associatie 'fire' - 'ontslaan' komt niet vaak voor), maar qua woordkeuze net iets beter dan de huidige Systran-technologie waar Google tot oktober 2007 gebruik van maakte: "Elke keer steek ik een taalkundige in brand, gaan de prestaties van het toespraakwaarnemingssysteem uit".

TST en ICT

De 'data driven' aanpak van TST-problemen werd gestimuleerd en mogelijk gemaakt door aanhoudende en snelle ontwikkelingen in hardware en software.

Zoals 3000 jaar geleden maakt men ook nu gebruik van de beschikbare technologie, die een fundamenteel andere strategie ondersteunt. Internet is in ICT de meest recente grote doorbraak. Alle TST-applicaties kunnen nu 'slim' zijn, want ze kunnen online toegang krijgen tot gigantische databestanden, snelle zoekalgoritmen, krachtige processoren enz.

Spraakherkenning en -synthese of automatisch vertalen zijn niet langer puur 'academische' uitdagingen. Ze zijn nu 'enabling technology' voor een breed scala aan toepassingen. Het STEVIN-programma, waarvan een flink aantal artikelen in dit nummer een uitvloeisel zijn, heeft hier een belangrijke rol in gespeeld. Wim Luites van Telecats merkt terecht op dat de markt nu klaar is voor TST: de acceptatie van spraaktechnologie is de laatste jaren groter dan ooit, niet omdat de algoritmes spectaculair beter zijn, maar omdat de kwaliteit van applicaties door 'tuning' sterk is toegenomen.

Succes door maatwerk

Er zou een optimaal herkennings-, synthese- en vertaalsysteem moeten bestaan dat alle variëteiten van 'het Nederlands' dekt en inzetbaar is in iedere denkbare applicatie. Een dergelijk universeel gereedschap bestaat niet. Verschillende artikelen in deze DIXIT laten echter overtuigend zien dat heel behoorlijke resultaten haalbaar zijn door maatwerk: houd

rekening met het doel van de applicatie, met de beoogde gebruiker en zijn context, met zijn kennis- en ervaringsniveau, zijn motivatie om te oefenen, etc. Dan blijkt TST een prima hulp voor de onderhoudsmonteur van militaire voertuigen, bij het transcriberen van telefoontaps van criminelen, in taalonderwijs en zorg op afstand, in de interactie met navigatie en infotainment in de auto, bij spellingcontrole, bij het samenvoegen van klant- en relatiegegevens uit verschillende bronnen, etc.

Voor optimale prestaties is het ook van groot belang dat TST goed wordt ingebed in een systeem dat op vele manieren 'slim' kan zijn, omdat het niet alleen (redelijk) kan spreken en verstaan maar ook over veel soorten niet-linguïstische informatie beschikt. Zo kunnen we bijvoorbeeld in dit nummer lezen hoe het 'voice portal' van een bank- en verzekeringsmaatschappij een klantvraag zo analyseert dat die vanzelf bij de deskundige medewerker terecht komt of hoe spraaktechnologie met 'augmented reality' wordt gecombineerd bij servicewerkzaamheden. Met taaltechnologie

kunnen vacatureteksten van recruiters en cv's van sollicitanten (bijna) automatisch aan elkaar worden gekoppeld, terwijl het toch om ongestructureerde documenten gaat.

Een fascinerende ontwikkeling is de integratie van spraak-, taal- en vertaaltechnologie in een consumentenapparaat als de smartphone: je produceert een zin, je spraak wordt herkend en omgezet in taal (de zin verschijnt desgewenst op het scherm), dan wordt de boodschap vertaald naar een gekozen vreemde taal (de tekst is eveneens te zien op het scherm) en tenslotte geconverteerd naar spraak. Het is een aanrader om deze Google Translate app voor ruim 50 talen zelf eens uit te proberen¹. Voor de 'early adopter' is dit misschien slechts een gadget, maar voor de ingewijde TST-er is het een prestatie waar het vakgebied trots op mag zijn. Dergelijke nuttige toepassingen op de mobiele apparaten die straks iedereen op zak heeft, kunnen de doorbraak van TST alleen maar versnellen.

¹) www.google.com/mobile/translate

Deze DIXIT wordt u aangeboden door de deelnemers van NOTaS:

- CLST, Radboud Universiteit Nijmegen
- Data Archiving and Networked Services
- Dedicon
- Dutcheer
- Gridline
- Kecida, NFI
- Knowledge Concepts
- Linguistic Systems BV
- OSTT/Sint Maartenskliniek
- RightNow Technologies
- Instituut voor Nederlandse Lexicologie, TST-Centrale
- TeleCats
- Universiteit van Tilburg – CIW/TICC
- Universiteit Twente – HMI
- Universiteit Utrecht – Utrecht Institute of Linguistics OTS

En in het bijzonder door:
STEVIN en Nederlandse Taalunie



10 Jaar NOTaS

10 Jaar geleden werd de Stichting NOTaS opgericht, op initiatief van Geert Kobus. De visie en de passie van Geert Kobus en zijn collega's van het eerste uur, waarmee zij de belangen van TST in Nederland hebben behartigd, lagen aan de basis van NOTaS. Met een enorme gedrevenheid probeerde Geert de krachten te bundelen van bedrijven die taal- en spraaktechnologie en TST-toepassingen ontwikkelden. Dat waren er toen nog niet veel, en de meeste waren klein. Bovendien zaten er relatief veel op korte afstand van elkaar in Brabant.

**Harry Bunt en
Debbie Kenyon-
Jackson**
Bestuur NOTaS

Het begon in 1999, toen Geert een aantal concullega's uitnodigde om te komen brainstormen. Dat werd het begin van een serie bijeenkomsten, bij Geert thuis of in zijn tuin, op zijn kantoor, en ook bij andere bedrijven o.a. in Oss en Nuenen. Al snel haalde Geert er ook TST-onderzoekers bij van de universiteiten van Tilburg en Eindhoven, omdat hij samenwerking tussen bedrijven en kennisinstellingen essentieel vond. Zijn passionele gedrevenheid werkte enorm stimulerend, en er ontstond groot enthousiasme voor zijn plannen. In opdracht van de Brabantse Ontwikkelings Maatschappij (BOM) werd in 2000 een markverkenning Taal- en Spraaktechnologie uitgevoerd door het bedrijf AXYS, waarin o.a. de volgende conclusies getrokken werden:

1. Na decennia van research en ontwikkeling zijn voor het eerst in de recente periode 1997-2000 sterk verbeterde TST-applicaties *succesvol* op de markt gebracht.
2. Voor TST is sprake van een definitieve doorbraak; onomkeerbaar, structureel, internationaal en van voldoende schaal-

grootte. Deze positieve ontwikkeling geldt zowel voor taal- als voor spraaktechnologie.

3. Bij Nederlandse TST-bedrijven is behoefte aan *positionering* als sector met gemeenschappelijke aanpak van kennisuitwisseling, marktbenadering, standaarden, juridische en andere operationele randvoorwaarden.

Besloten werd tot het oprichten van het 'Samenwerkingsverband TST', waarvan de oprichtingsakte van 9 februari 2001 als doelstelling noemt "*het onder de aandacht brengen van de sector taal- en spraaktechnologie in de markt en het bewerkstelligen van een impuls voor marktontwikkeling door onder meer het organiseren van concrete samenwerkingsprojecten*" en verder vermeldt dat binnen een jaar beslist zal worden over een eventuele rechtsvorm zou krijgen. Daar waren maar drie maanden voor nodig: in mei 2001 werd de stichting NOTaS opgericht.

De pioniers zagen de enorme toegevoegde waarde in de samenwerking tussen bedrijven onderling en tussen bedrijven en kennisinstellingen. Veel van deze bedrijven waren immers ontstaan als spin-offs van universitaire onderzoeksgroepen. Van elkaar leren en kennis delen waren belangrijke speerpunten in de eerste jaren. Een andere belangrijke doelstelling was van begin af aan de lobby om TST op de kaart te krijgen bij EZ, OCW, NWO en andere instanties die technologisch onderzoek en ontwikkeling stimuleren.

Geert Kobus was behalve een gepassioneerd spreker ook een uitstekende schrijver en heeft diverse visiestukken geschreven, o.a. voor de Nederlandse Taalunie. Samen met medebestuurleden heeft hij ook een actieve rol gespeeld bij het tot stand komen van STEVIN. Dat was niet altijd gemakkelijk, en de belangen van NOTaS lagen niet altijd op één lijn met die van de overheid. Door de jaren heen is dit sterk veranderd, hebben de overheid (EZ, Senter, Agentschap NL) en



NOTaS elkaar gevonden, en er is bijvoorbeeld een nauwe samenwerking ontstaan tussen NOTaS en de Taalunie.

In de loop van de jaren zien we de twee pijlers, het delen van kennis door te netwerken en het lobbyen voor de TST, steeds weer terugkomen. Onderzoeksprojecten binnen STEVIN laten goed zien hoe successen geboekt kunnen worden door samenwerking, hoe nieuwe inzichten kunnen ontstaan, en hoe langdurige samenwerkingen tot stand komen.

Een derde pijler van NOTaS is markteducatie. Eindgebruikers schaffen nauwelijks taal- en spraaktechnologie als zodanig aan; meestal is men geïnteresseerd in een praktisch eindresultaat: een geautomatiseerd call-centre of self-help service desk; hulpmiddelen voor leerondersteuning; een zoekmachine die meedenkt als je niet precies weet hoe iets heet... In zulke systemen speelt TST een essentiële rol, dat is een rol achter de schermen. Markteducatie is erop gericht om het besef te vergroten dat TST een belangrijke 'enabling technology' is, de technologische basis van een scala van ICT-toepassingen.

Het tijdschrift DIXIT, door NOTaS vooral in het leven geroepen om het bereiken van netwerk- en markteducatiedoelstellingen te ondersteunen, is door de jaren heen een essentieel onderdeel geworden van het communicatiebeleid van NOTaS. Deze speciale editie van DIXIT informeert u in terugblik over het STEVIN-programma, aansluitend op de afsluitende STEVIN-dag op 29 november 2011.

Ter gelegenheid van het 10-jarig bestaan van NOTaS heeft het bestuur zich samen met de huidige deelnemers gebogen over de vraag welke rol NOTaS de komende jaren het beste kan spelen. Drie rollen kwamen daaruit naar voren als belangrijk en gewenst: NOTaS als netwerkorganisatie, als lobbyorganisatie, en als organisator van markteducatie-activiteiten. Instrumenten als DIXIT en Taal en Bedrijf worden daarbij hoog gewaardeerd. Opgemerkt is dat NOTaS als netwerkorganisatie een speciale rol zou kunnen spelen voor het behoud van de samenwerkingen die ontstaan zijn in het STEVIN-programma, nadat STEVIN als zodanig aan zijn eind gekomen is.

NOTaS was oorspronkelijk opgezet ten behoeve van de taal- en spraaktechnologie in Nederland, maar mede onder invloed van

de Vlaams-Nederlandse samenwerking in STEVIN zien we dat het ook zinvol is voor NOTaS om de belangen te behartigen van de Nederlandstalige taal- en spraaktechnologie. Daarom staat NOTaS sinds kort ook open voor Vlaamse deelnemers. Daarmee hopen we bredere samenwerking te stimuleren en de banden die zijn ontstaan tijdens STEVIN te behouden en te versterken.

We hebben de afgelopen 10 jaar al veel mooie dingen bereikt. Maar niet alleen de technologische ontwikkelingen gaan snel, ook het 'speelveld' verandert voortdurend. Er is nog veel te doen! De TST-pioniers en de oprichters van NOTaS hebben een fundament gelegd, waarop wij verder kunnen bouwen. De uitdagingen zijn groot en daarmee ook de kansen. De vertaling van de nieuwste onderzoeksresultaten van de kennisinstellingen naar de praktijk biedt steeds weer kansen voor nieuwe toepassingen. Door nieuwe initiatieven te steunen zoals Clarin en Clariah zal NOTaS een actieve rol blijven spelen in de voorhoede van de taal- en spraaktechnologie in Nederland – en Vlaanderen. Daarnaast willen we via markteducatie nieuwe klanten aanspreken. De overheid als launching customer biedt steeds meer nieuwe concrete commerciële mogelijkheden. NOTaS zal zich hierop blijven richten.

"De markt is klaar voor spraaktechnologie"

Directeur Wim Luimes van Telecats over de toekomst van taal- en spraaktechnologie

"Vijf jaar geleden vertelde ik aan klanten wat er allemaal mogelijk was, nu vertel ik hoe dat allemaal toegepast wordt. Maar we verkopen geen radicaal nieuwe toepassingen." Volgens directeur Wim Luimes van spraaktechnologiebedrijf Telecats zat de ontwikkeling in de markt voor taal- en spraaktechnologie (TST) de afgelopen jaren vooral in de marktacceptatie van de technologie. Voor de toekomst verwacht hij een voortzetting van dat patroon.



Leonoor
van der Beek

Kennis genereren

Dat wil niet zeggen dat er geen mooie dingen gebeuren. De politie draait een proef met het opnemen en onderzoekbaar maken van processen verbaal. In plaats van aantekeningen te maken en die later uit te werken, belt de agent een medewerker die de informatie gestructureerd invoert. "Stel, er is een gele scooter gejat", legt Luimes uit. "Dan spreekt die agent alle informatie in, en dat gesprek wordt met behulp van spraakherkenning omgezet naar tekst." De agent hoeft dus minder papierwerk te doen. Maar het mooiste moet nog komen: "Als nu een week later een gele scooter gebruikt wordt bij een overval, dan kan de compu-

ter die twee voorvallen aan elkaar linken." Dankzij het opnemen en omzetten van het gesprek naar tekst kan de computer dat zelfs wanneer de kleur niet in de database is ingevoerd. "Zo wordt kennis gegenereerd. Dat gebeurt nu nauwelijks."

Taaltechnologie

Met name door het optimaliseren van de bestaande spraakherkenningstechnologie voor specifieke toepassingen, en door het combineren van spraak- en taaltechnologie, wordt de functionaliteit en bruikbaarheid van spraakherkenningstoepassingen uitgebreid of verbeterd. "Tien jaar geleden wilden we het letterlijke woord *auto* herkennen. Nu willen we ook de juiste respons kunnen geven wanneer iemand *wagen* zegt, of het over een Volvo heeft. We moeten niet alleen herkennen wat er gezegd wordt, maar ook begrijpen wat er bedoeld wordt."

Tunen

De verbeteringen zitten hem volgens Luimes niet in betere technologie. "De spraakherkenning is de laatste jaren niet extreem verbeterd. Wat wel verbeterd is, is de manier waarop die wordt toegepast." De combinatie met taaltechnologie is daar een voorbeeld van. Maar er zijn er meer: "We verbeteren de herkenning door gebruik te maken van slimmigheidjes en statistiek. Wat zijn populaire onderwerpen? Dan is de kans dat een audiosignaal daarmee te maken heeft relatief groot." Zo bleek er voor het routeren van gesprekken naar het juiste politiecursus een lineair verband te zijn tussen het aantal inwoners en het aantal keer dat naar de politie in die plaats gevraagd werd. "De kans dat naar Haarlem of Arnhem gevraagd wordt, is gewoon groter dan de kans dat naar het dorp Haren in de provincie Groningen gevraagd wordt." Klinkt vanzelfsprekend, maar dat soort logica heeft de toepassingen enorm



Christian Wartena
Novay

Het wordt steeds moeilijker om relevante informatie te vinden. Het indexeren van het hele internet met de woorden van webpagina's voldoet niet meer: de systemen van de toekomst moeten niet alleen vast kunnen stellen welke woorden in een tekst genoemd worden, maar vooral wat de relevante concepten en onderwerpen zijn, hoe een stuk informatie samenhangt met andere vindplaatsen en hoe relevant het is in de context van de gebruiker. Hierbij zal taaltechnologie steeds belangrijker worden.

De grootste groei op het internet is te zien in multimedia en video. Video is niet meer beperkt tot amusement: er komen steeds meer video's beschikbaar van toespraken, lezingen, symposia en colleges. Deze wil je kunnen doorzoeken, en dat zal de komende jaren een grote impuls geven aan de verdere ontwikkeling van spraakherkenning. Taal- en zoektechnologie komen daarmee voor een dubbele uitdaging te staan: ze moeten kunnen omgaan met de fouten en onzekerheid in de spraakherkenning én ze moeten in staat zijn de inhoud van (gesproken) documenten te interpreteren. Bestaande methodes en technieken zullen verder ontwikkeld moeten worden. Onderzoek, en dus investering in taal- en spraaktechnologie, blijft daarom ook na afloop van het STEVIN-programma nodig!

verbeterd. "Je kunt de technologie niet out of the box inzetten. Je zult altijd moeten tunen. Dat realiseren klanten zich ook steeds beter."

Nuance

Verbeterd de taal- en spraaktechnologie dan niet meer? Luimes gelooft van wel. "De computer gaat steeds meer 'begrijpen'. Die gaat de logica die wij er nu in stoppen zelf toepassen. Bijvoorbeeld context of non-verbale signalen zoals wijzen." Hij verwacht dat die ontwikkelingen voornamelijk aangestuurd zullen door de academische wereld en niet door bijvoorbeeld Nuance, verreweg de grootste leverancier van spraakherkenningstechnologie. "Want", zo stelt Luimes, "Nuance moet toch allereerst geld verdienen." En bovendien: "Het belang voor Nuance om dat extra procentje herkenning te krijgen is wellicht ook niet zo groot."

Duurzaamheid

Luimes ziet de aandacht voor duurzaamheid als een stimulans voor spraaktechnologie. "Ik hoorde gisteren dat we ondanks alle digitalisatie vorig jaar toch weer meer geprint hebben. Organisaties die veel vergaderen, zoals bijvoorbeeld Provinciale Staten of de gemeenteraad, produceren enorme pakketten papier over wat er verteld is. Ik zie voor mij dat er in de toekomst nog slechts 2 A4tjes worden uitgedeeld met de hoofd- en actiepunten, en als je precies wilt weten wat er gezegd is, dan kun je hier de opnames vinden en dan kun je zelf zoeken."

Privacy

Aan de andere kant worden ontwikkelingen in de spraaktechnologie beperkt door privacyreglementen. "Persoonlijk vind ik dat vreemd", aldus Luimes. "Iedereen vertelt alles maar op Twitter, maar aan de andere kant wordt persoonlijke data zeer streng beschermd door juridische restricties op dit gebied. Dat kan een struikelblok zijn voor het testen van nieuwe technologie. Dat geldt niet voor openbare in-



Erwin Schaapman
directeur Data Direction

Content is een zwaar onderschatte bedrijfsasset en bepalend voor de prestaties van taal- en spraaktechnologie. Het is hét belangrijkste ingrediënt voor elke dialoog die een organisatie met haar doelgroepen heeft en die medewerkers nodig hebben voor hun contact met de klant. Tot op heden is de verantwoordelijkheid voor content versnipperd binnen organisaties en veelal alleen op operationeel niveau belegd. De verantwoordelijkheid voor content én contentapplicaties moet strategisch worden gepositioneerd.

formatie, zoals vergaderverslagen, maar wel data uit contactcenters, en zeker ook bij de pilot bij de politie."

Visionairs

Maar de toekomst voorspellen, dat kan ook Luimes niet. "Tien jaar geleden werd er voorspeld dat er niet meer gebeld zou worden, maar de jeugd belt steeds meer. Er werd voorspeld dat er niet meer geprint zou worden, en dat is ook niet waar gebleken." Bovendien besteedt Luimes zijn aandacht liever anders dan aan het volgen van grootse toekomstvisioenen: "Uiteindelijk moeten we gewoon in gesprek blijven met onze klanten. Dat is het belangrijkste, we moeten die markt verder ontwikkelen. Waar het naar toe gaat, ik weet het niet. Bij ons in bedrijf werken veel visionairs. Zij vinden mijn ideeën soms nogal achterlopen. Maar als ik die ideeën bij klanten vertel, dan loop ik op hen voor." Nou voorruit, één voorspelling wil hij wel doen. "Het spraakgestuurd aansturen van applicaties, dat gaat wel gebeuren. Daar zie ik veel beweging. Maar dat wil niet zeggen dat Telecats dat gaat doen."

Marcel Smit

oud-CEO Q-go, nu VP Natural Language Solutions,
RightNow Technologies



Ik verwacht dat taal- en spraak technologie veel meer geïntegreerd gaan worden in bedrijfsapplicaties. Daar is het nog steeds lastig om in grote hoeveelheden data precies datgene te vinden wat je nodig hebt, bijvoorbeeld de tien meest winstgevendende klanten rond Amsterdam in de maand juni. Er zullen veel meer dialoog-systemen gaan komen waarin je in normaal Nederlands dit soort gegevens kunt opvragen in plaats van met complexe analyses. Een vergelijkbare ontwikkeling ga je - veel meer dan nu het geval is - ook zien in consumententoepassingen, zoals spraaksturing om een tv-programma op te nemen of een favoriete song op de iPod te vinden.

Automatische spraakherkenning voor het taalonderwijs en de zorg

Automatische spraakherkenning (ASH) wordt al ingezet voor tal van toepassingen, variërend van treinreisinformatiesystemen, het melden van klachten bij krantenbezorging, het verschaffen van informatie over flitspalen, en 'Google Search'. Al lang wordt gedroomd van het inzetten van ASH ten behoeve van de zorg en het taalonderwijs. Dit heeft vooral te maken met het enorme potentieel dat deze technologie kan bieden voor spraaktherapie en het trainen van spreekvaardigheid, waarvoor nu doorgaans niet genoeg tijd beschikbaar is. Door gebruik te maken van ASH zou het mogelijk moeten zijn om programma's te ontwikkelen waarmee patiënten en cursisten in hun eigen omgeving, op hun eigen tempo eindeloos kunnen spreken, en het liefst ook feedback krijgen van de computer. Voor taalleren zijn er intussen al wel commerciële producten op de markt die dit beweren te doen. Cursisten zijn daar niet altijd tevreden over, omdat die producten niet goed genoeg werken, hun beloftes niet altijd waarmaken, maar ook omdat de cursisten vaak te hoge verwachtingen hebben.

**Helmer Strik,
Catia Cucchiari
CLST Radboud
Universiteit
Nijmegen**

Om te begrijpen in hoeverre ASH nuttig kan zijn voor taalleren en de zorg moet eerst duidelijk worden gemaakt dat deze technologie niet feilloos is, en nog wel degelijk beperkingen heeft. Dit houdt bijvoorbeeld in dat het op dit moment nog niet mogelijk is om 'spontaan' tegen een computer te praten, in de hoop dat die alles zal verstaan wat tegen hem gezegd wordt. Ondanks deze beperkingen is het toch mogelijk om nuttige oefeningen te ontwerpen waarbij cursisten verschillende aspecten van spreekvaardigheid kunnen oefenen door tegen de computer te spreken, die vervolgens de spraak analyseert en daar feedback over geeft.



Automatische spraakherkenning ten behoeve van het taalonderwijs

Op het Centre for Language and Speech Technology (CLST) van de Radboud Universiteit Nijmegen wordt al jaren onderzoek hiernaar gedaan. In het kader van het NWO-project Dutch-CAPT¹ werd onderzocht in hoeverre ASH kan worden gebruikt om feedback te geven op de uitspraak van Nederlandse klanken die voor anderstaligen moeilijk zijn. Het systeem is getest en bleek effectief te zijn: voor de cursisten die ons

systeem gebruikten was de reductie in het aantal uitspraakfouten significant groter dan bij een controlegroep.

Een vervolg op het succesvolle Dutch-CAPT project is het STEVIN-project DISCO². Hierin wordt een geheel nieuw systeem ontwikkeld dat taalleerders automatisch feedback geeft over niet alleen uitspraak, maar ook morfologie en syntaxis. In het kader van dit project is ook onderzoek verricht om het opsporen van fouten steeds nauwkeuriger te doen. Voor de cursist is het immers belangrijk dat de feedback die de computer geeft zo goed mogelijk is.

In het NWO-project FASOP³, wordt een systeem met ASH gebruikt om onderzoek te doen naar het effect van correctieve feedback bij het leren van een tweede taal. Dit maakt het mogelijk om verschillende vormen van feedback op een systematische manier te onderzoeken. Ook kunnen alle interacties tussen leerling en computer worden opgeslagen. Deze kunnen later worden gebruikt om te onderzoeken hoe en hoeveel leerlingen oefenen, welke fouten ze maken, hoe ze reageren op de feedback die ze krijgen, en in hoeverre ze vooruitgaan, maar ook om technologie en systeem te verbeteren.

Bovenstaande projecten hebben allemaal als onderwerp het leren van Nederlands als tweede taal (NT2). Recent zijn we ook begonnen met onderzoek naar het leren van het Engels in het project MPC, dat mogelijk is door een STW 'valorisation grant'. In MPC⁴ wordt een systeem ontwikkeld dat onderwijsinstellingen kunnen gebruiken bij het ondersteunen van het onderwijs Engels. In eerste instantie gaat het om een systeem

dat feedback geeft op de uitspraak van het Engels door Nederlandstaligen. Het ligt in de bedoeling om het systeem later uit te breiden naar andere aspecten van spreekvaardigheid en naar andere talencombinaties.



Automatische spraakherkenning voor de zorg

Spraak van taalleerders wijkt op vele manieren af van de spraak van moedertaalsprekers. Dat is ook het geval bij spraak van mensen met communicatieve beperkingen. Automatische herkenning van moedertaalspraak is al lastig, maar herkenning van deze zogenaamd atypische spraak is nog veel complexer, juist omdat die spraak op allerlei manier afwijkend kan zijn, qua uitspraak, morfologie, en syntaxis. Daarom is het nodig om taal- en spraaktechnologie te ontwikkelen die specifiek geoptimaliseerd wordt voor de afzonderlijke doelgroepen. Naast onderzoek over de inzet van ASH bij taalleren, doen wij ook onderzoek naar de inzet van ASH bij ondersteuning van mensen met communicatieve beperkingen.

In het ZonMw-project PEDDS⁵ wordt technologie ontwikkeld die later toegepast kan worden in therapie voor dysartriepatiënten. De technologie maakt het mogelijk om fouten in de uitspraak te detecteren en daarover feedback te geven. Ten slotte is er het ZonMw-project CRDP⁶. In steeds meer e-Health-applicaties moeten patiënten communiceren met computerprogramma's, bijvoorbeeld op patiëntenwebsites, de zogenaamde DigiPoli's. Voor mensen met communicatieve beperkingen is dit vaak problematisch. Technologieën zoals 'chat-by-click', woordpredictie, spraaksynthese, en spraakherkenning kunnen ingezet worden om het voor deze mensen makkelijker te maken om te communiceren 'met de computer'.

Meer applicaties in de toekomst

In de toekomst zullen dit soort applicaties nog belangrijker worden. Door de toegenomen mobiliteit is er een groeiende vraag

naar taalleren. De vergrijzing leidt tot meer vraag naar zorg. De inzet van docenten en therapeuten is duur, en er is een tekort aan gekwalificeerde mensen. Ten slotte, op vele terreinen moet bezuinigd worden; e-Learning en e-Health bieden nieuwe mogelijkheden om goed onderwijs, ondersteuning en zorg te blijven bieden. Interessante ontwikkelingen zijn ook te verwachten door de combinatie van ASH met 'virtual reality', 'serious gaming', en 'mobile learning'. Daarvoor is het dan wel nodig dat er goed onderzoek verricht wordt naar deze nieuwe methodes door multidisciplinaire teams waarin experts uit het veld samenwerken met taal- en spraaktechnologen.

- ¹⁾ lands.let.ru.nl/~strik/research/Dutch-CAPT
- ²⁾ *Development and Integration of Speech technology into Courseware for language learning*: lands.let.ru.nl/~strik/research/DISCO
- ³⁾ *Feedback and the Acquisition of Syntax in Oral Proficiency*: lands.let.ru.nl/~strik/research/FASOP.html
- ⁴⁾ *My Pronunciation Coach*: lands.let.ru.nl/~strik/research/MPC-STW-VG1 en lands.let.ru.nl/~strik/research/MPC-STW-VG2
- ⁵⁾ *Pronunciation Error Detection for Dysarthric Speech*: lands.let.ru.nl/~strik/research/PEDDS
- ⁶⁾ *Communication & Revalidation DigiPoli*: lands.let.ru.nl/~strik/research/CRDP

AEGON Voiceportal

Innovatieve oplossingen met open vraag spraakherkenning bij AEGON

Bij AEGON werden we geconfronteerd met een interessante puzzel. AEGON heeft de verschillende labels zoals bijvoorbeeld Spaarbeleg samengevoegd en onder één noemer AEGON gebracht. Daar hoort één telefoonnummer voor alle klanten bij. AEGON heeft drie vestigingen en meer dan 1200 producten waarvan de kennis verspreid is over deze drie locaties. Met klassieke middelen is dit eigenlijk niet meer te behappen en AEGON wilde daarom een voiceportal waarmee alle klanten in één keer gerouteerd konden worden naar die medewerker die de specifieke vraag het best kon beantwoorden. Open vraag spraakherkenning vormt de basis van de oplossing.

Martijn van de Runstraat
Telecats BV

Dialoog

Klanten die AEGON bellen, worden eerst door het systeem geïdentificeerd op basis van de postcode en huisnummer. Deze hoeft u overigens niet telkens opnieuw in te spreken bij elk gesprek; het systeem herkent herhaalbellers aan de hand van het telefoonnummer. Het systeem zoekt na identificatie de klantgegevens erbij en weet precies wat voor producten de klant heeft. Vervolgens kan de klant de vraag inspreken.



Voorbeeld van klantvraag classificatie met productbezit filter

Open Vraag SpraakHerkenning (OVSH)

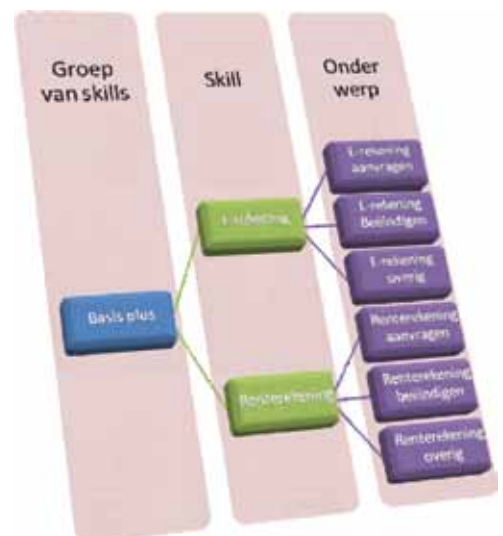
wordt gebruikt om antwoorden op een open vraag te herkennen, bijvoorbeeld: "Waarmee kan ik u van dienst zijn?". De beller kan vervolgens op zijn/haar eigen wijze de aanleiding van het gesprek vertellen. Bij AEGON hebben we er voor gekozen om iets meer richting te geven aan de beller met de vraagstelling: "spreek uw vraag in, en noemt u daarbij zo mogelijk het product". We hebben vastgesteld dat dit helpt om de beller de vraag zo te laten formuleren dat we goed kunnen vaststellen waar het om gaat en wie ze daarmee het beste kan helpen. De vraag wordt door de spraakherkenner zo goed mogelijk herkend en daarna door een classifier voorzien van een aantal mogelijke categorieën waar deze vraag bij hoort. Nu kan het systeem de optimale bestemming voor het

gesprek bepalen. Als iemand bijvoorbeeld iets over zijn "polis" wil weten, is het heel handig dat we weten dat deze persoon een schadeverzekering heeft.

Elke skill in het contactcenter bevat één of meerdere categorieën of onderwerpen. De skills zijn weer gebundeld in groepen, waarin agenten zijn vertegenwoordigd die al deze skills bezitten. De voice portal maakt van deze gelaagde structuur gebruik om slim te routeren: Als de voiceportal bijvoorbeeld wijfelt tussen het onderwerp "renterekening" en "e-rekening", dan kan deze gewoon naar de skillgroep routeren waarin deze beide skills vallen. De klant merkt zo niets van deze onzekerheid en kan toch direct geholpen worden zonder extra vragen of doorschakelingen. Sommige vragen kunnen het beste door de computer zelf worden afgehandeld. Bijvoorbeeld overboekingen kunnen geheel automatisch door het systeem worden afgehandeld met een selfservice toepassing.

Analyse van de klantvraag

Hoe zet je zo'n applicatie in elkaar? We zijn begonnen met de analysefase waarin we



Gelaagde classificatiestructuur

een Wizard-of-Oz experiment deden. Hierbij hebben we een klein gedeelte van de klanten door de datacollector geleid. In plaats van het gebruikelijke keuzemenu, werd deze mensen gevraagd om de vraag in te spreken, precies alsof er al een spraakherkenning systeem achter zat. We kregen hiermee een precies beeld van wat mensen op welke manier inspreken. De meeste mensen spreken de vraag in alsof ze een bericht inspreken op het antwoordapparaat: "Met meneer Jansen, mijn wasmachine is kapot gegaan en nu heb ik waterschade.", sommigen houden het wat korter: "Ik heb een vraag over de lijfrente", of "Renterekening". De vraag die de klant ingesproken had, werd vervolgens door een medewerker van een speciaal hiervoor ingericht team beluisterd, waarna de medewerker direct kon beginnen door de vraag te beantwoorden, of het gesprek ongemerkt door te verbinden naar een medewerker met de juiste kennis.

Training

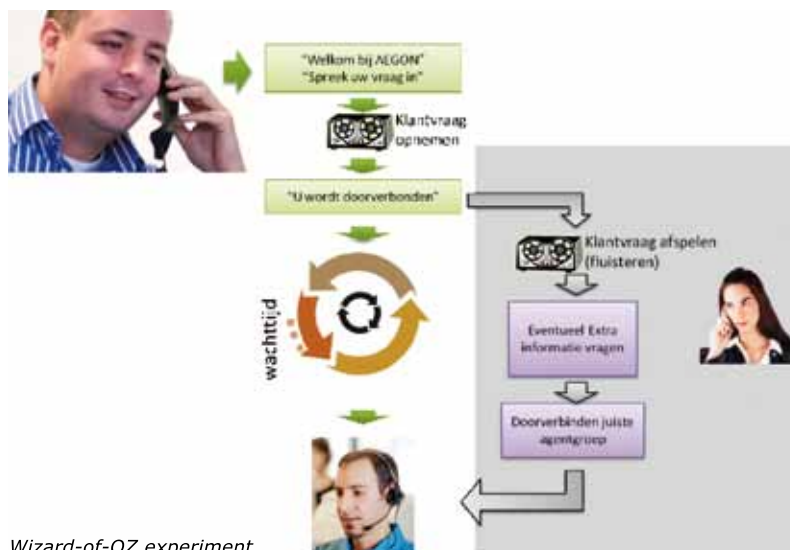
De verzamelde vragen zijn allemaal beluisterd en uitgeschreven. Met deze vragen, aangevuld met documenten van AEGON waarin veel van het gebruikte jargon voorkomt, hebben we de spraakherkenner getraind. Vervolgens hebben we de geregistreerde vragen op onderwerp ingedeeld (geclassificeerd). Het onderwerp heeft vaak betrekking op een product, maar het kan ook iets anders zijn als bijvoorbeeld een verhuizing. In overleg met AEGON hebben we de vragen in de gewenste categorieën ondergebracht en daarmee het classificatiesysteem getraind. Met de simulatie-opstelling konden we al voordat we live gingen met de applicatie precies voorspellen hoe het systeem zich zou gedragen.

Succes

Waarom zijn klanten en medewerkers van AEGON nu zo blij met deze voiceportal? Natuurlijk is het fijn dat het systeem de klanten bij de juiste medewerker aflevert. Maar door een koppeling met de Agent DeskTop applicatie die de medewerkers gebruiken, zetten we de klantgegevens en de vraag van de klant direct op het scherm. De medewerkers kunnen het gesprek daarom meteen heel prettig beginnen: "U heeft een vraag over uw inboedelverzekering?". De klant voelt zich begrepen en merkt dat het nuttig was om de gegevens en de vraag in te spreken. Het gesprek gaat direct over de inhoud waardoor de klant zo snel en goed mogelijk te woord te wordt gestaan.

Resultaat

Van het eerste idee tot en met het live gaan van het systeem, zijn we ruim een half jaar bezig geweest. Omdat we in het project ook

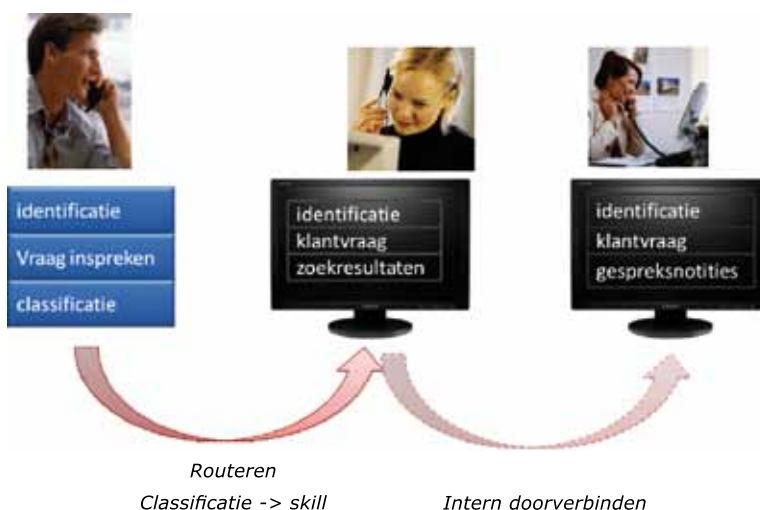


Wizard-of-Oz experiment

te maken hadden met andere ontwikkelingen en afhankelijkheden bij AEGON, heeft het uiteindelijk wat langer geduurd, maar altijd hebben we gezamenlijk het einddoel voor ogen gehouden, wat tot een prachtig resultaat heeft geleid.

In de periode na de livegang van de voiceportal hebben we het systeem verder geoptimaliseerd. Zo hebben we analyses gemaakt van de gesprekken die zijn doorverbonden. Zijn de gesprekken door de agent, die het gesprek kreeg, afgehandeld, of is er al snel doorverbonden? Daar waar er snel werd doorverbonden is er vaak ruimte voor verdere finetuning van de voiceportal. Momenteel worden meer dan 90% van de gesprekken die door de voiceportal zijn gerouteerd, afgehandeld door de agent die het gesprek aangeboden krijgt. Daarmee is de belangrijkste doelstelling van AEGON behaald.

De voiceportal draait nu sinds 2009 naar grote tevredenheid van AEGON en haar klanten. Maar we zitten niet stil. We blijven met AEGON kijken naar mogelijkheden om het systeem verder te optimaliseren en bij te stellen aan de hand van de veranderende business.



Met spraaktechnologie telefonische contacten analyseren

De telefoon blijft een belangrijk medium om in contact te treden met bedrijven en instellingen. Veel algemene vragen kunnen ook via internet afgehandeld worden, maar voor persoonlijke vragen en voor zaken waarbij emotie een rol speelt, zal de klant vaak naar de telefoon willen grijpen. Telefonische contacten gelden echter als kostbaar, dus als het mogelijk is wil je als bedrijf of organisatie de onnodige telefoontjes zoveel mogelijk voorkomen. En voor de gesprekken die je wel afhandelt, wil je natuurlijk dat dat op een professionele en klantvriendelijke manier gebeurt. Om hier greep op te krijgen en te houden is spraakanalyse een krachtige oplossing. Met taal- en spraaktechnologie worden telefonische contacten analyseerbaar gemaakt, zonder dat alle gesprekken moeten worden beluisterd.

Martijn van de Runstraat
Telecats BV

Stuurinformatie voor het contactcenter

In het contactcenter worstelt men met vragen als: "Waarover bellen onze klanten en wanneer doen ze dat?", "Wat zijn de belangrijkste redenen om ons te bellen en hoe gaan we daar vervolgens mee om?", "Zijn sommige gesprekken wellicht gemakkelijk te voorkomen?", "Handelen we gesprekken wel efficiënt en klantvriendelijk af?". Om goede sturing te geven aan een contactcenter is inzicht hierin onontbeerlijk. Op de één of andere manier moet men dus zorgen dat informatie, die aanwezig is in het contactcenter, beschikbaar gemaakt wordt voor de rest van de organisatie. Een bekende mogelijkheid hiervoor is gesprekslogging, waarbij na elk gesprek door de agenten wat gegevens worden vastgelegd over het gesprek. Dit proces is tijdrovend, maar ook redelijk onbetrouwbaar. Hoe zorg je ervoor dat de onderwerpen waaruit de agent kan kiezen volledig zijn en toch snel vindbaar? Hoe voorkom je dat onderwerpen waarop je niet geanticipeerd had, terecht komen onder 'overig' en dus onopgemerkt blijven? Hoeveel tijd en energie wil je van je agenten vragen voor het invoeren van deze informatie en hoe zorg je dat ze dat zorgvuldig doen en niet uit gemakzucht veel onder 'overig' classificeren?

betekenis van het gesprek vaak beter herkend worden. Deze woorden zijn vaak langer en worden duidelijker in het gesprek uitgesproken. Bijvoorbeeld de zin:

"Ik zat te kijken in dat *overzicht* van jullie waarop het *termijnbedrag* staat..."

De twee woorden 'overzicht' en 'termijnbedrag' zijn hier al voldoende om het onderwerp van de zin te bepalen. Dit zijn gemakkelijk te herkennen woorden voor de spraakherkenner. Om de spraakherkenning verder te verbeteren, gebruiken we verschillende technieken. Gesprekken moeten bij voorkeur worden opgenomen op twee sporen (stereo). Zo kunnen we de beller en de medewerker gemakkelijk uit elkaar houden. We kunnen beter zien welke teksten bij elkaar horen en we kunnen specifiek analyseren op basis van wie iets gezegd heeft. Voordat spraakherkenning wordt losgelaten op de opnames, maakt het systeem eerst een audio-analyse. Hierbij wordt bepaald welke onderdelen van het gesprek bestaan uit wachttijd, stiltes, interactie met een IVR, etc. Ook wordt vastgesteld of gesprekken worden doorverbonden, doordat het stemprofiel van de kant van de medewerker opeens verandert na een doorschakeling.

Verdere analyse

Na het uitvoeren van de spraakherkenning hebben we de gesprekken in tekstvorm en kunnen we daarin gaan zoeken naar woorden en woordcombinaties. Maar we gaan veel verder. De computer kan het gesprek automatisch analyseren op basis van onderwerp, afhandel-tijd, interactie met de beller, woordgebruik, etc. Ook wordt het mogelijk om het gesprek te segmenteren. Zo is er bij een typisch gesprek sprake van een vaste opbouw:

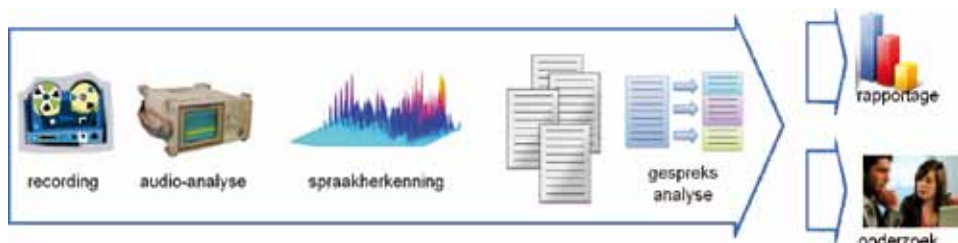
- introductie, eerste probleemstelling
- identificatie
- analyse
- oplossing
- afsluiting

Spraakanalyse

Er is nu een mooi alternatief: Spraakanalyse. De computer zet de gesprekken om in tekst, waarna we in de gesprekken kunnen zoeken naar woorden of onderwerpen en allerlei analyses kunnen uitvoeren op basis van de inhoud van de gesprekken.

Spraakherkenning van telefonische gesprekken werkt nog niet foutloos, maar dit is meestal ook niet nodig om het gespreksonderwerp eruit te halen. Ongeveer de helft van de woorden wordt doorgaans correct herkend en dat is in de praktijk voldoende voor een goede analyse. Maar effectief is de herkenning veel hoger, aangezien woorden die belangrijk zijn voor de





Het spraakanalyseproces

De kwaliteit van een gesprek is op elk van de segmenten meetbaar te maken. Hiermee kunnen we de kwaliteit beoordelen en gesprekken lokaliseren die beter bekeken of beluisterd moeten worden.

Gemeten karakteristieken kunnen automatisch in beeld worden gebracht en gehouden. Zo kan je bijvoorbeeld van dag tot dag in de gaten houden of bepaalde onderwerpen meer dan anders voorkomen, in welke mate gesprekken over een bepaald onderwerp snel en in één keer kunnen worden afgehandeld en of er veel herhalingen in gesprekken voorkomen.

Naast de dagelijkse 'vinger aan de pols' door het contactcenter management, kan een analist uit de gesprekken natuurlijk nog veel meer te weten komen over het klantcontact. Hoe vaak komt het niet voor dat er rondom een bepaald onderwerp (zoals bijvoorbeeld het overstappen van of naar een bepaalde concurrent of vragen over een nieuw product) door het management gevraagd wordt om te onderzoeken hoeveel er over gebeld wordt en of dat goed loopt? Op basis van enkele woorden kan de analist gemakkelijk wat gesprekken opzoeken en vaststellen of de juiste informatie wel altijd wordt gegeven. De automatische transcripten van de gevonden gesprekken kunnen eventueel worden bekeken en de gesprekken kunnen worden beluisterd om een goed beeld te krijgen van wat er aan de hand is. Hiermee heb je ook meteen een aantal goede voorbeeldgesprekken.

Een krachtige techniek die hier gebruikt kan worden, is clustering. Hierbij verdeelt de computer de gesprekken geheel automatisch op 'onderwerp' in een door de gebruiker gevraagd aantal groepen. Vergelijk het met een handvol punaises die je op tafel uitspreidt. Op basis van kleur en positie kan een clusteralgoritme uit de schijnbaar ongeordende zaken groepen maken, die elk 'redelijk' homogeen zijn. Bij gesprekken kijkt de computer naar woorden die bij bepaalde gesprekken voorkomen en bij andere gesprekken juist niet. Elk cluster wordt daarmee beschreven door een aantal beschrijvende woorden en een aantal discrimine-

rende woorden (woorden waarop het juist van andere clusters afwijkt).

Door de gespreksgegevens te verrijken met gegevens die met spraakanalyse zijn bepaald, krijg je een compleet beeld van de aanleiding en afhandeling van gesprekken. Door technieken en gegevens te combineren, kun je onderwerpen ontdekken die wellicht niet efficiënt of correct worden afgehandeld. Zo is het mogelijk om onderwerpen te vinden waarbij veel herhaalverkeer voorkomt, of onderwerpen met een hogere gemiddelde gesprekstijd. Vaak kunnen ook gespreksclusters worden gelokaliseerd die met enkele eenvoudige verbeteringen van de processen en/of website kunnen worden gereduceerd.

What's next?

De kwaliteit van Spraakanalyse zal blijven verbeteren. De spraakherkenning zal nauwkeuriger worden, de snelheid zal toenemen en de dialoog-analyse zal krachtiger worden. Semantiek en sentiment mining krijgen een grotere rol. Daarmee zal er een verschuiving plaatsvinden van analyse op basis van woorden naar analyse op basis van begrip en betekenis. Zo kunnen we in de komende jaren toepassingen verwachten waarbij gesprekken worden beoordeeld op de wijze waarop de agent reageert op de klant en vice-versa. Begrijpen ze elkaar? Spreken ze elkaar tegen of zijn ze het eens? Gaat dat op een prettige manier of is er sprake van onbegrip? In welke mate is de agent in staat de beller tevreden te stellen? Praten beide partijen op hetzelfde niveau? En dat alles kan op veel grotere schaal worden uitgevoerd omdat de capaciteit van de systemen die de gesprekken moeten verwerken alleen maar toeneemt.



Clustering



Voorbeeld van een analyse van gesprekken van klanten die vaker dan 1 keer bellen

Spraakherkenning spoort criminelen op

In Nederland bestaat de mogelijkheid om toestemming te krijgen voor het afluisteren van telefoongesprekken (telefoontaps). Het is goed te weten dat er niet zomaar opnames gemaakt mogen worden, bijvoorbeeld door je werkgever. Dit heeft alles te maken met de privacyrechten die je als burger hebt. Toestemming wordt gegeven door de rechter als de noodzaak tot afluisteren is aangetoond.

Hans Jongebloed
Dutchear

Huidige werkwijze

Voor een gemiddelde strafzaak en/of veiligheidsonderzoek worden al gauw honderden uren aan telefoongesprekken opgenomen. Dit betreft dan alle telefoongesprekken van en naar een aantal specifieke telefoonnummers. Veel van deze gesprekken hebben niets met het onderzoek te maken, zoals pizza bestellen, korte kletspraatjes, enquêtes, enzovoort. Om toch belangrijke informatie boven water te krijgen worden alle gesprekken afgeluisterd en (samengevat) in tekst omgezet door medewerkers met speciale bevoegdheden. In sommige gevallen moet zelfs een tolk ingeschakeld

worden. De tekstuele versies van de gesprekken worden vervolgens doorzocht op het voorkomen van steekwoorden, zoals namen, plaatsen of voorwerpen. Gesprekken waarin deze woorden voorkomen worden uiteindelijk aan een nadere analyse onderworpen. Twee nadelen van deze methode zijn: het kost veel tijd (en dus geld) om alle gesprekken af te luisteren en bovendien is het voor mensen lastig om alle gesprekken letterlijk weer te geven; als later in het onderzoek de zoektermen veranderen, dan ontbreken deze in de tekstuele uitwerkingen (vooral als alleen een samenvatting is gemaakt).

telefoontaps heeft Dutchear de spraakherkenning zodanig aangepast, dat de kans op het herkennen van de steekwoorden bijna 100% is geworden. De rest van de gesproken woorden mocht dan verkeerd herkend worden. Hierbij bleken er ook steekwoorden opgeleverd te worden die in werkelijkheid niet in de audio te horen waren (de zogeheten false positives). Aangezien dit ook verwacht werd, was al bepaald dat een gespecialiseerde onderzoeker van Fox-IT verifieert welke gevonden steekwoorden terecht of onterecht zijn door het specifieke gedeelte van de geluidsopnames te beluisteren. De uiteindelijke gevolgde stappen zijn weergegeven in het stroomschema. Aangezien de steekwoorden maar een klein deel van alle opnames betrof hoefde er nog maar heel weinig audio afgeluisterd te worden, zelfs als ook de onterecht gevonden steekwoorden meegeteld werden. Het is bovendien goed mogelijk om bij wijzigingen in de gezochte steekwoorden opnieuw alle geluidsopnames opnieuw te laten herkennen.

Evaluatie door experts

In een onafhankelijke evaluatie van de spraakherkenning door Fox-IT bleek dat het systeem heel goed werkte. Het bleek mogelijk om de hoeveelheid af te luisteren materiaal met 90% te reduceren, terwijl nog steeds bijna alle steekwoorden gevonden werden. Hierbij waren ook opnames met een andere taal; de steekwoorden waren zodanig geconfigureerd dat ook met de uitspraak in een andere taal het steekwoord gevonden werd. In deze evaluatie werd onderstreept dat mensen deze taak ook niet perfect kunnen uitvoeren: een aantal door de machine gevonden steekwoorden waren door de mens gemist!

Deze methode biedt hiermee de mogelijkheid om meer audiomateriaal doorzoekbaar te maken, zonder dat er veel mensinzet nodig is. Denk hierbij vooral ook aan bedrijven die grote call centers hebben waar dagelijks honderden gesprekken binnenkomen met waardevolle informatie over producten en diensten voor een bedrijf.

Links: www.dutchear.nl en www.fox-it.com



Toepassing van spraakherkenning

Dutchear is de uitdaging aangegaan om voor Fox-IT – een bedrijf gespecialiseerd in digitale recherche – het proces met spraakherkenning te versnellen. Er zit een sterk raakvlak tussen het opsporen van de steekwoorden en het achteraf zoeken à la Google in grote hoeveelheden audio, zoals Dutchear al heeft toegepast bij de Rechtbank van Amsterdam. Het bijzondere van deze uitdaging was dat - indien mogelijk - ALLE voorkomens van de steekwoorden moesten worden opgespoord.

Tot op de dag van vandaag werkt spraakherkenning helaas niet 100% perfect. Om het toch te kunnen inzetten bij het afluisteren van de



Contextgebaseerde spellingcorrectie met Valkuil.net

Sinds 31 mei 2011 is het Nederlandse taalgebied een spellingcorrector rijker: Valkuil.net. De online spellingcorrector is contextgebaseerd, wat betekent dat niet alleen typfouten gedetecteerd worden, maar ook verwarringen tussen bestaande woorden. Waar de spellingcorrectie van Word geen fout ziet in de zin 'hij zij dat hij naar huis ging', lukt dat Valkuil.net wel. Valkuil.net let ook op spatiefouten ('verpleeg kundige', 'mooieboel'), en kan d/t-fouten detecteren.

Valkuil.net signaleert spelfouten aan de hand van de context waarin ze voorkomen. De corrector is getraind op ongeveer anderhalf miljard woorden aan Nederlandse tekst, en vergelijkt iedere nieuwe tekst met dit achtergrondgeheugen. Fouten vallen op doordat ze in contexten voorkomen waarin Valkuil.net een ander woord zou verwachten. Het woord 'geld' is een prima Nederlands woord, maar in de zin 'Dat geld ook voor jou' ziet Valkuil.net dat er een sterk gelijkend woord is, 'geldt', dat veel beter past. Bij deze statistische vergelijking wordt geen expliciete taalkundige kennis gebruikt, en ook geen vaste woordenlijst. Net als Google Translate kan Valkuil.net gezien worden als een voorbeeld van taaltechnologie die gebaseerd is op een grote hoeveelheid basismateriaal. Valkuil.net kent het fragment 'Dat geldt ook' simpelweg omdat het vaak voorkomt in teksten.

Het systeem is niet perfect: het mist nog wel eens echte fouten en corrigeert soms onterecht een prima woord. Momenteel is de corrector zo ingesteld dat hij zo min mogelijk vals alarm slaat (in tegenstelling tot bijvoorbeeld de spellingcorrector in Word) en tegelijkertijd zo weinig mogelijk fouten over het hoofd ziet; een delicate balans. Een voorbeeld. In de zin 'Het licht niet aan het duobestuur' ziet Word geen probleem in 'licht', maar wel in het woord 'duobestuur' dat niet in de woordenlijst van Word staat. Valkuil.net verbetert 'licht' naar 'ligt' en laat 'duobestuur' ongemoeid.

Valkuil.net komt voort uit het onderzoeksproject 'Implicit Linguistics', uitgevoerd aan de Universiteit van Tilburg bij het Tilburg center for Cognition and Communication. Het project, geleid door Antal van den Bosch, werkte ook aan een automatisch vertaalsysteem gebaseerd op dezelfde basistechnologie.

Valkuil.net is bedoeld voor iedereen die Nederlandse teksten schrijft, voor studie, werk of persoonlijke doeleinden, en die snel een externe 'lezer' over de tekst wil laten lopen om wellicht nog enkele fouten af te vangen. Door psycholinguistisch onderzoek weten we

dat zelfs de meest geoefende schrijvers nog fouten maken, vaak in de haast, en dat deze fouten vaak verwarringen zijn tussen gelijkklinkende woorden zoals 'gebeurt' en 'gebeurd', of 'nog' en 'noch'. Waar de meeste andere spellingcheckers deze fouten ongemoeid laten, is Valkuil.net in staat om dit soort verwarringen op te sporen. De tool is daarom een nuttige uitbreiding voor iedere professionele schrijver, maar het is zeker ook een tool om te gebruiken voordat de sollicitatiebrief de deur uit gaat, of het opstel, samenvatting of werkstuk wordt ingeleverd. De spellingcorrector geeft geen garanties, maar iedere verbeterde fout kan een belangrijke zijn.

Valkuil.net is een voorbeeld van 'cloud computing': een dienst die ergens op internet beschikbaar is en toegankelijk is voor mensen (via de webpagina) of software (via de webservice). De webapplicatie die op <http://valkuil.net> te gebruiken is, is gebaseerd op de webservice-software CLAM (Computational Linguistics Application Mediator), die ook in de Tilburgse groep werd ontwikkeld.

De onafhankelijkheid van menselijke expertise of taal maakt Valkuil.net een uiterst snel inzetbare technologie; een spellingcorrector voor een andere taal kan in het tijdsbestek van enkele dagen ontwikkeld worden. De praktische uitdaging van het team achter Valkuil.net, dat ook software voor automatisch vertalen en parafraseren ontwikkelt, is om de technologie zo bruikbaar en beschikbaar mogelijk te maken voor iedereen. De wetenschappelijke uitdaging is om taalkundigen te overtuigen dat modellen als Valkuil.net echte taalkundige modellen zijn, die net als klassieke taalkundige modellen taalkundige oordelen, voorspellingen en inzichten kunnen genereren.

Links:

- valkuil.net
- ilk.uvt.nl/clam - webservice-software
- ilk.uvt.nl/il - het onderzoeksproject
- www.tilburguniversity.edu/nl/thema/innovatie/taalsoftware/

Antal van den Bosch
CLST Radboud
Universiteit
Nijmegen

Hoe kan de 'onderzoekende zoeker' beter ondersteund worden?

Zoekmachines op internet zijn heel geschikt om allerlei soorten informatie te vinden over een onderwerp. Als ik in Google de zoekopdracht 'Multatuli' geef dan krijg ik 811.000 verwijzingen naar websites met informatie over de 19e-eeuwse Nederlandse schrijver. Maar als ik specifieke informatie over Multatuli zoek, zoals wat voor verbanden er zijn tussen zijn werk en politieke gebeurtenissen uit die tijd, of hoe andere schrijvers op zijn boeken reageerden, dan is de Google-strategie niet meer zo behulpzaam. Google geeft mij namelijk altijd verwijzingen naar complete webpagina's, hoe specifiek mijn zoekvraag ook is.

Suzan Verberne
CLST Radboud
Universiteit
Nijmegen

Complexe zoekvragen door 'onderzoekende zoekers'

Complexe zoekvragen, die bijvoorbeeld in gaan op verbanden tussen entiteiten, zijn voorbeelden van vragen die een historicus of literatuurwetenschapper zou kunnen hebben. Als deze onderzoeker geïnteresseerd is in literaire bronnen, dan kan hij of zij terecht bij gespecialiseerde tekstcollecties, bijvoorbeeld de Digitale Bibliotheek voor de Nederlandse Letteren [1]. DBNL heeft, net als de meeste andere online tekstcollecties, een zoekinterface die is geïnspireerd op Google: als ik de zoekvraag 'Multatuli' ingeef in het DBNL-zoek-systeem [2] dan krijg ik 127 pagina's met resultaten. Als ik op één van de gevonden fragmenten klik, dan krijg ik een document te zien waarin alleen mijn zoekterm Multatuli is gemarkeerd. De tekst die ik op mijn scherm heb, is lang en ik zal hem moeten lezen om erachter te komen waar de informatie staat waar ik naar op zoek ben. 'Onderzoekende zoekers' zoals de literatuurwetenschapper in dit voorbeeld zijn bereid tijd te besteden aan het bestuderen van tekstuele

bronnen. Daarbij zouden ze echter veel beter geholpen kunnen worden dan met alleen het tonen van volledige documenten. Een eerste stap zou hierbij zijn als het zoekstelsel namen en gebeurtenissen zou markeren in het document dat de onderzoeker bekijkt. Als ik bijvoorbeeld een recensie van de Max Havelaar lees in DBNL zouden titels van boeken, schrijvers en tijdschriften (Vaderlandsche Letteroefeningen, bijvoorbeeld) gemarkeerd kunnen worden.

Een volgende stap zou zijn om op verzoek van de gebruiker informatie te tonen over de gemarkeerde termen. Die informatie zou bijvoorbeeld uit een encyclopedie kunnen komen. Welke informatie precies van belang is, hangt af van de doelgroep van het zoekstelsel. Als ik op 'Vaderlandsche Letteroefeningen' klik, zou ik bijvoorbeeld informatie uit een dossier van de Koninklijke Bibliotheek kunnen krijgen: "Vaderlandsche Letteroefeningen was meer dan een eeuw lang een van de toonaangevende literair-culturele tijdschriften van Nederland.", met een link naar externe bronnen en andere plaatsen in

DBNL waar het tijdschrift genoemd wordt. Als documenten in tekstcollecties zoals DBNL verrijkt worden met dit soort informatie, kunnen 'onderzoekende zoekers' zoals literatuurwetenschappers en historici in de toekomst beter geholpen worden bij hun literatuurstudie.

I. *Max Havelaar, of de koffij-veilingen der Nederlandsche Handel-Maatschappij, door Multatuli. Twee deelen. Tweede druk. (Te) Amsterdam, (bij) J. de Ruijter, 1860. Prijs f 4,00.*
II. *Waarheen? Een woord aan de lezers van Max Havelaar. Tweede druk. (Te) Leiden, (bij) P. Engels, 1860. Prijs f 0,40.*
In een tijdschrift dat den naam van Vaderlandsche Letteroefeningen draagt, mag een werk dat zooveel sensatie gemaakt heeft in den lande als *Max Havelaar*, niet onvermeld blijven, vooral wanneer zulk een werk een literarisch meesterstuk is. Dat eerst nu eene aankondiging van *Multatuli's* geschrift plaats heeft, wijte men den uitgever de Ruijter, die geen exemplaar van het werk *ter recensie* zond aan de redactie van een tijdschrift, waarin alleen bij zeldzame uitzondering niet ingezonden boekwerken worden aangekondigd. Onze uitgever, de heer van der Mast, gaf een blijk van zijnen ijver door het kwaad te verhelpen, en zond mij eerst kort geleden *Max Havelaar* ter aankondiging. Mijne taak is door deze min wenschelijke vertraging zeer moeilijk geworden. Reeds is er zoo menig woord over het bijna alom gelezen boek gevallen, zoo menig verschillend oordeel daarover geveld, dat ik bijna schroom onzen lezers mijnen schralen mosterd na den maaltijd voor te zetten. Vooral gevoel ik dien schroom, wanneer ik mij herinner, welk eene meesterlijke beschouwing van *Max Havelaar* Prof. P.J. Veth

Taaltechnologie die onderzoekend zoeken mogelijk maakt

Het markeren van belangrijke termen en namen in tekst is een taak die al goed mogelijk is met bestaande taaltechnologie. Zo worden in PoliticalMashup [3], een samenwerkingsproject van NRC Handelsblad en het Informatica Instituut van de Universiteit van Amsterdam, politieke debatten automatisch geanalyseerd op

onderwerpen die aan bod komen per partij en politicus, en welke namen en termen daarbij het vaakst gebruikt worden.

Voor het vergaren van feitelijke informatie die gekoppeld kan worden aan termen t.b.v. het ondersteunen van zoekers is nog meer onderzoek nodig. Daarvoor moet namelijk niet alleen een naam herkend worden als een entiteit (Multatuli is Eduard Douwes Dekker), maar ook de relatie met andere entiteiten gevonden worden ("Multatuli schreef de Max Havelaar"). Die relaties kunnen op veel verschillende manieren worden uitgedrukt: "de Max Havelaar is door Multatuli geschreven", "de auteur van de Max Havelaar is Multatuli", etc.

Initiatieven wereldwijd

In het onderzoeksproject Extracting factoids from Dutch texts, dat in 2011 wordt uitgevoerd aan de Radboud Universiteit Nijmegen, wordt hieraan gewerkt. Met financiering van Google in het kader van het Europese digital humanities-programma wordt in dit project software ontwikkeld die feitelijke informatie vergaart uit Nederlandse Wikipedia-teksten.

zelfs een artikel in Nature, geschreven door de Oren Etzioni, de directeur van het Turing Center van de universiteit van Washington in Seattle [6]. Hij noemt onder andere de doorbraak van de IBM-computer Watson, die een kennisquiz won van de beste menselijke deelnemers. Op <http://tiny.cc/jeopardy> kun je zien hoe Watson de quiz 'Jeopardy' won. IBM onderzoekt nu de mogelijkheden om Watson in te zetten voor zoekvragen in plaats van trivia-vragen.

- ¹⁾ DBNL: www.dbnl.org
- ²⁾ www.dbnl.org/zoeken/zoekeninteksten
- ³⁾ www.nrc.nl/denhaag
- ⁴⁾ Alle projecten die dit jaar worden gefinancierd staan op tiny.cc/googleawards
- ⁵⁾ rtw.ml.cmu.edu/rtw
- ⁶⁾ www.nature.com/nature/journal/v476/n7358/full/476025a.html



Met behulp van die software kunnen in de toekomst tekstuele bronnen verrijkt worden met achtergrondinformatie. Google ondersteunt met de digital humanities awards de ontwikkeling van technologie ten behoeve van onderzoekers in de geesteswetenschappen [4]. In de Verenigde Staten worden momenteel een aantal grote projecten uitgevoerd die zich bezighouden met het automatisch leren van feiten uit teksten [5]. Sommige wetenschappers vinden dat zoeksystemen op internet veel preciezer te werk zouden moeten gaan in het beantwoorden van vragen. In augustus van dit jaar verscheen daarover

In de frontlinie op Facebook en Twitter

In het najaar van 2010 lanceerde Mozilla Firefox een nieuw concept in klantenservice: The Army of Awesome¹. Het bedrijf verzamelde een 'legertje' vrijwilligers dat reageert op Tweets van anderen. Zodra een tweet het woord 'firefox' of 'mozilla' bevat, wordt deze opgenomen in de tijdlijn. De vrijwilligers selecteren hier een tweet die ze willen beantwoorden. Vervolgens presenteert Firefox hen een lijstje met de meest gestelde vragen en een 'Twitter-proof' (lees: maximaal 140 tekens lang) antwoord hierop, dat ze nog kunnen personaliseren alvorens te verzenden². Maar niet ieder bedrijf heeft zoveel deskundige fans, dat het in goed vertrouwen zijn klantenservice kan crowdsourcen. Hoe dan toch de enorme hoeveelheden vragen in de verschillende kanalen te verwerken? Taaltechnologie helpt.

Kim Sterenborg
Ordina

Kanalen

Consumenten verwachten van bedrijven dat die altijd en overal service kunnen verlenen. Daarnaast verwachten zij dat de geboden service via alle kanalen van dezelfde kwaliteit is. En het aantal kanalen groeit gestaag. Mensen pakken weliswaar nog steeds de telefoon om een vraag te stellen, maar doen dat lang niet altijd meer door te bellen. Ze gebruiken een bedrijfsapp of gaan via Twitter of Facebook op jacht naar informatie. Ruim zeventig procent van de Nederlanders is actief op tenminste één sociaal netwerk en ruim veertig procent maakt gebruik van mobiel internet³. Ook daar verwachten zij klantenservice te vinden.



Goedkoop

Niet alleen consumenten willen graag online geholpen worden. Ook voor bedrijven is het belangrijk goede webservice te verlenen. Het is namelijk vele malen goedkoper om een klant online te helpen dan face-to-face of per telefoon: als een klant online alle informatie vindt, zal hij geen kostbaar telefoontje meer plegen. Dat geldt niet alleen voor de website van het bedrijf zelf, maar ook voor de Social Media kanalen van het bedrijf.

Onderzoek

Uit onderzoek onder bedrijven die RightNow Intent Guide gebruiken, blijkt dat bijna 65% van de bedrijven verwacht dat eindgebruikers self-service tools in Social Media zullen gebruiken⁴. Iets minder, ongeveer de helft

van deze bedrijven, gebruikt op dit moment al Social Media om vragen te beantwoorden. Maar slechts een klein gedeelte hiervan is ervan overtuigd dat de antwoorden op vragen in alle kanalen uniform zijn - hier is dus winst te behalen.

KLM

Een goed voorbeeld van een groot Nederlands bedrijf actief in Social Media is KLM. Zo verraste de luchtvaartmaatschappij passagiers eind 2010 met kleine, gepersonaliseerde cadeautjes op basis van hun Social Media activiteiten⁵. Op de Facebookpagina en het Twitter-account van KLM komen vooral vragen van passagiers voorbij: vragen over bagage, inchecken en boekingen. Wat opvalt, is dat het veelal vragen zijn van dezelfde strekking. Vragen die nu nog één voor één door agents beantwoord worden, maar die een bedrijf ook gestandaardiseerd kan oplossen, mits het zijn online content goed op orde heeft en die content ook goed te vinden is.

Taaltechnologie

Natuurlijke taaltechnologie kan een grote bijdrage leveren aan de vindbaarheid van informatie: als de zoekfunctie de intentie van een gebruikersvraag weet te achterhalen, kan de klant beter automatisch naar relevante content geleid worden. Wanneer de technologie spelfouten corrigeert en syntactische of lexicale varianten van de vraag kan herkennen, zal de klant snel zijn antwoord vinden, ook als hij niet de officiële termen gebruikt. Dat is in Social Media nog belangrijker dan elders op het web, omdat het taalgebruik er nog informeler is, en er heel waarschijnlijk een groot verschil bestaat tussen hoe gebruikers over producten of processen praten, en de manier waarop bedrijven dat op hun officiële website doen. Een bedrijf als KLM zal de workload van de Social Media agents kunnen verminderen, door de natuurlijke taaltechno-

logie die zij op hun website gebruiken, ook op Facebook in te zetten – uiteraard uitgebreid met Facebookspecifieke content. Op de Britse Facebookpagina doen zij dit al [6].

Twitter

Klantenservice via Twitter heeft nog een extra uitdaging. Niet alleen zijn er heel veel service gerelateerde vragen met het woord of de hashtag KLM (of enig ander bedrijf), er zijn daarnaast nog vele malen meer niet-servicegerichte tweets met datzelfde woord of hashtag. Agents zijn veel tijd kwijt met het selecteren van de relevante tweets, die zij kunnen en willen beantwoorden. Ook hier kan taaltechnologie helpen. RightNow onderzoekt op dit moment de mogelijkheden om de tweets op basis van intentie te classifi-

ceren. In een dashboard worden alle tweets verzameld op onderwerp, bijvoorbeeld annuleringen of aswolk. De agents kunnen deze vervolgens gericht beantwoorden met de informatie uit de database – met natuurlijk de mogelijkheid het antwoord te personaliseren. Zo vraagt het Social Media leger slechts een minimum aan manschappen [7].

¹⁾ support.mozilla.com/nl/army-of-awesome

²⁾ www.youtube.com/watch?v=P970H6dS88I

³⁾ discoverdigitallife.com/

⁴⁾ www.rightnow.com/cx-suite-intent-guide.php

⁵⁾ www.youtube.com/watch?v=pqHWA8GDEk

⁶⁾ http://www.facebook.com/klmuk?sk=app_232425713461135

⁷⁾ www.youtube.com/watch?v=3SuNx0UrnEo

- advertentie -



CLST

fundamenteel en toepassingsgericht onderzoek op het gebied van TST

- Data mining & knowledge discovery (tekst en audio)
- Taalleren: evaluatie en training
- Zorg: diagnostiek, therapie en ondersteuning
- Forensisch onderzoek en toepassingen daarvan
- Resources

CLST zoekt samenwerking met kennisinstellingen en bedrijven.
Tevens biedt **CLST** consultancy.

Bezoek ons op: www.ru.nl/CLST

CLST

Centre for Language and Speech Technology

Radboud Universiteit Nijmegen




De digitale rechercheur

"Entiteit Extractie als middel in de strijd tegen illegale activiteiten"

In de wereld van toezichhouders, fraudeonderzoekers en ordehandhavers draait het groten-deels om het tijdig en correct herkennen van 'named entities' (zoals personen, organisaties, datums, telefoonnummers en kentekens) in combinatie met de 'gedragingen' van deze entiteiten. Het in de tijd kunnen ordenen van deze gegevens is essentieel bij de strijd tegen velerlei illegale activiteiten.

Omdat de wereld nu eenmaal in snel tempo aan het digitaliseren is, richt het een deel van het onderzoek zich vooral op het omgaan met de enorme hoeveelheid beschikbare digitale informatie. Vroeger volstond het wellicht voor instanties die zich met de controle en handhaving op illegale praktijken bezig houden, om wat te googelen op internet. Nu gaat dat niet meer en moet men gebruik maken van specialistische statistische en taalkundige software om die enorme bulk aan informatie te kunnen processen.

Jan Paul Raven
Knowledge
Concepts

Uitdaging

Enigszins versimpelend, kan men zeggen dat het bij de te nemen horden bij opsporing en handhaving vooral gaat om:

- De hoeveelheid informatie die de verschillende diensten dagelijks te verwerken krijgen
- De verscheidenheid in opslagmethoden en documenttypes waarmee gewerkt moet worden
- De variabelen in taal- en woordgebruik in de te verwerken informatie
- De tijdspanne waarbinnen de onderzoeken dienen plaats te vinden
- De beschikbaarheid van adequate en getrainde medewerkers om de onderzoeken te doen

Om adequaat te kunnen werken dient een opsporingsinstantie het vermogen te hebben om de beschikbare informatie snel en grondig te kunnen analyseren, maar ook om aan de hand van nieuw gevonden indicaties, retrospectief te kunnen zoeken in al eerder verwerkte informatie. Immers niet alle illegale activiteiten worden de eerste maal als zodanig herkend en geen enkele instantie heeft het vermogen qua mankracht, tijd of budget om alles uit het verleden nogmaals handmatig door te spitten. De 'kennis van nu' moet dus gebruikt worden om met terugwerkende kracht in de eerder ontsloten data te gaan zoeken. Gedacht kan worden aan een overval gepleegd met een rode Vespa. Dan moet automatisch in de 'oude data' gezocht worden naar mogelijke vermeldingen van een rode Vespa (bv een gestolen rode scooter).

Cold case

Bij een inval vorig jaar in het kader van Antiterrorisme Bestrijding werden twee dozen met papieren ordners, 18 cd's, 4 dvd's, 1 usb-stick en 2 computers met elk een 2TB harddisk in beslag genomen. Deze data wilde men snel analyseren alvorens de sporen 'koud' zouden worden.

Bij een quick scan bleek een groot gedeelte van de data in een taal te zijn waarvan het aantal 'native speakers' zeldzaam was binnen de opsporing-instantie. Tegen de tijd dat men alles vertaald, bestudeerd en geanalyseerd had was elke kans op een snelle arrestatie verlopen. De combinatie van een wereld zonder grenzen en moderne ICT vereist duidelijk een andere aanpak dan we van de meeste detectiveseries gewend zijn!



Content Enabler Discovery

De afgelopen jaren werd hard gewerkt aan software waarmee in ieder geval een deel van de hierboven geschetste problemen opgelost zou kunnen worden: de 'Content Enabler Discovery' (CED).

Een onderzoeksproces zoals gebruikelijk bij veiligheidsdiensten is op te delen in drie logische blokken:

1. Verkregen data wordt omgezet in leesbare tekst of documenten via oplossingen en tools als bijv. Forensic Tool Kit, EnCase, RIMAGE, scan- en OCR-processen e.d.
2. Deze stroom wordt door CED geanalyseerd op taal, grammaticale inhoud en indeling, de verschijning van entiteiten en waar mogelijk worden al relaties tussen de entiteiten bepaald aan de hand van de tekst en de zinsbouw. Een eerste snelle documentanalyse en verrijking kan vervolgens plaatsvinden op documentniveau binnen CED.
3. De originele data samen met de entiteiten

en verrijkingen kunnen vervolgens ingelezen worden in analysetools om document overschrijdende analyses te doen of om deze data te linken aan andere al bekende informatie. Deze tools (Analyst's Notebook en I-Base) visualiseren de relaties tussen de verschillende entiteiten en kunnen hiermee veel inzicht brengen in complexe netwerken.

Hoe werkt het?

De data wordt allereerst van de verschillende datadragers afgenomen en via een aantal methodes gestandaardiseerd (leesbaar gemaakt). Deze gestandaardiseerde datastromen worden vervolgens langs een aantal processen geleid welke de inhoudelijke gegevens herkent en indien gewenst extraheert en/of relateert. De stappen zijn te simplificeren tot de volgende hoofdelementen:

- taalherkenning
- de tekst opbreken in logische elementen
- zinsontleding
- algemene herkenning van termen of samengestelde termen aan de hand van verschillende methoden (regels – lijsten – logica)
- herkennen van entiteiten (gebruikmakend van de eerdere verfijning en analyse)
- en als laatste stap de koppeling van entiteiten en gebeurtenissen (denk hier aan zaken als: de heer Marcel Jansen, geboren te Den Haag op ..-.-..... wonende te).

Deze geëxtraheerde informatie kan vervolgens gegroepeerd worden over grotere datahoeveelheden heen waardoor verbanden en samenhangen kunnen worden vastgesteld die boven de individuele dataelementen uitstijgen.



Uitsnede van de basisinterface van CED

Door deze integrale inzet van forensische middelen, taal- en zoektechnologie en analysesoftware worden medewerkers van opsporingsinstanties in staat gesteld snel, flexibel en vooral nauwgezet grote datastromen te analyseren op verdachte activiteiten.

Conclusie

Met het (tijdig) herkennen van patronen en het attenderen op 'vreemd' gedrag in de informa-

tiestroom is de opsporingsinstantie in staat snel en onderbouwd triage te verrichten. Men besteedt mankracht en middelen daar waar men denkt dat er iets te halen valt! Dit is goedkoper, sneller en uiteindelijk effectiever. Bij nieuwe kennis kan het systeem snel aangepast worden aan een nieuwe situatie.

CED is een in Nederland, door Knowledge Concepts, ontwikkeld product dat het mogelijk maakt om snel en in vele talen verdachte activiteiten en personen te herkennen en deze naar boven te brengen vanuit een grote informatiestroom. Het past naadloos in het werkproces van de opsporingsambtenaren in de eerste fase van een onderzoek en is door de modulaire opzet eenvoudig in te bouwen in bestaande werkprocessen en oplossingen.

Stemmen maken

*In 2005 publiceerde de Taalunie het rapport *Taal- en spraaktechnologie en communicatieve beperkingen*, een inventarisatie van wat er op TST-gebied zou moeten gebeuren om gebruikers met visuele of spraakhandicaps beter van dienst te zijn. Een van de conclusies was dat er op het gebied van synthetische spraak behoefte bestond aan een ruimer aanbod aan stemmen.*

Arthur Dirksen
Fluency

Ons kleine taalgebied was, ook in 2005, niet slecht bedeed met computerstemmen. Er waren verschillende aanbieders actief op de Nederlandse markt, en elk bood minstens twee stemmen aan. Maar het helpt een spraakgehandicapt meisje van, zeg, twaalf jaar, niet zoveel als ze kan kiezen uit drie of vier volwassen vrouwenstemmen.

Idealiter kan elke gebruiker van spraaksynthese een stem kiezen die bij hem of haar past. Wat zou het mooi zijn als je de computer een paar voorbeelden kon laten horen van de stem die je wilt, en hopla, geheel automatisch rolt zo'n stem eruit. Daar wordt wel onderzoek naar gedaan, maar voor praktische toepassing is de kwaliteit nog onvoldoende.

Ook binnen de bestaande technieken zijn er echter mogelijkheden om met beperkte inspanning en tegen relatief lage kosten, stemmen te maken voor spraaksynthese. In 2008 bracht Fluency een vernieuwde spraaksynthesizer op de markt, waarvoor snel een nieuwe stem gemaakt kan worden. Bij de release waren zeven stemmen beschikbaar (drie mannen-, twee vrouwen-, en twee tienerstemmen), en recent zijn nog vier stemmen toegevoegd. Daarnaast produceert Fluency ook speciale stemmen, onder meer voor patiënten met ALS, een spier- en zenuwaandoening.

Kleiner corpus

De nieuwe synthesizer maakt gebruik van de populaire unit-selectietechniek. Hierbij zoek je in een geannoteerde database met spraak eenheden bij elkaar waaruit je een nieuwe zin kunt samenstellen. De eenheden kunnen woorden zijn of lettergrepen, maar in de praktijk wordt meestal gebruik gemaakt van difonen, de overgangen tussen spraakklanken.

Normaliter bevat de spraakdatabase minstens drie uur spraak, die in een professionele stu-

dio wordt ingesproken door een professionele spreker. Deze opnames worden automatisch geannoteerd, en vaak handmatig verder bijgewerkt.

De Fluency-synthesizer maakt het makkelijker om een stem te creëren. We werken met een relatief klein corpus, bestaande uit 387 woorden en 387 zinnen, bij elkaar slechts drie kwartier spraak. Voor het inspreken maken we gebruik van low-budget apparatuur (een notebook met een goede usb-microfoon), en het idee is dat iedereen die in staat is de woorden en zinnen voor te lezen een geschikte spreker zou kunnen zijn.



Het opnameprogramma is vrij beschikbaar, en een spreker kan zelf, in een rustige ruimte zonder al te veel echo, de opnames maken. Elk item (woord/zin) wordt door de spraaksynthese voorgezegd, en de spreker zegt dit zo precies mogelijk na, met name wat betreft uitspraak en pauzes, maar wel met zijn eigen intonatie en manier van spreken. De spreker luistert of alles goed is opgenomen, eventueel samen met een begeleider.

De opnames worden door Fluency grotendeels automatisch verwerkt tot een spraakdatabase. Met de stemmen die via deze procedure tot stand komen kan goed verstaanbare,

levendige spraak gegenereerd worden, die een goede gelijkenis vertoont met de spraak van de inspreker.

Tienerstemmen

Met name voor jonge spraakgehandicapten is het belangrijk dat er synthetische stemmen zijn in hun leeftijdscategorie. Maar ook voor dyslectische kinderen is het prettig als zij in hun voorleessoftware een leeftijdgenoot kunnen kiezen. Om deze reden, maar ook om te zien of het überhaupt kon, hebben we enkele tieners achter de microfoon gezet.

De jongste spreker, David, was 13 jaar. Hij had wat moeite om de vaak moeilijke woorden en soms lange zinnen zonder haperen in te spreken. Maar hij was enthousiast genoeg om een opname desnoods tien keer over te doen. Het resultaat, de synthetische variant van een Amsterdams straatschoffie, is een goed voorbeeld van wat je met deze techniek kunt bereiken.



Mijn nichtje Fiona (16 jaar) had geen enkele moeite met het materiaal, zodat we in twee zondagmiddagen een stem hebben opgenomen die zelfs geschikt is voor professionele toepassingen.

Een leuke verrassing was de bijdrage van Koen, een 14-jarige jongen, die geheel zelfstandig met zijn eigen DJ-apparatuur het materiaal heeft ingesproken, met als resultaat een zeer bruikbare jonge mannenstem.

Het is een open vraag of het mogelijk is om ook kinderstemmen op deze manier te realiseren. Hier is zeker belangstelling voor, maar voor zulke jonge sprekers blijkt het in te spreken materiaal toch iets te volwassen. Het zou mogelijk moeten zijn om een aangepast inspreekcorpus samen te stellen, dat beter past bij de woordenschat en leesvaardigheid van kinderen. De uitdaging daarbij is echter om toch de volledige rijkdom en complexiteit van de Nederlandse uitspraak in zo'n corpus te representeren.

Hagenezen, Surinamers en Vlamingen

In opdracht van de Gemeente Den Haag hebben we voor de lees voor-functie van hun website een Haagse stem gemaakt. Het moest een soort Haagse Harry zou worden, en de in het Haagse zeer populaire entertainer Jeffrey Huf werd gekozen als spreker.

De schorre schreeuwstem waarmee Jeffrey de 774 woorden en zinnen in de microfoon slingerde, bleek een zware test voor onze analysesoftware, en er was nog flink wat handwerk nodig om hier een goede spraakdatabase van te maken. Het resultaat is echter erg vermakelijk, en de Jeffrey-demo is dan ook een populair onderdeel van onze website. Dat we onze software zonder aanpassingen plat Haags kunnen laten spreken is een indicatie dat varianten van het Nederlands geen bijzondere problemen zouden moeten opleveren. Om dit terrein wat verder te verkennen hebben we ook een jonge Surinaamse vrouwenstem en een Vlaamse mannenstem opgenomen. Kanttekening bij de Vlaamse stem: strikt genomen is alleen een Vlaamse tongval niet voldoende. Ook de uitspraak van veel woorden is in het Belgisch-Nederlands anders, en daar is de software nog niet op aangepast.

De hoeveelheid combinaties van geslacht, leeftijd en tongval is schier oneindig. De tot nu toe geproduceerde reeks stemmen is dan ook nog maar een druppel op een gloeiende plaat.

Meer info en audio-demonstraties: www.fluency.nl

Watson voor het Nederlands

Op dinsdag 15 februari 2011 werd in de Verenigde Staten de finale uitgezonden van een reeks afleveringen van de quiz show Jeopardy, waarin IBM's computersysteem Watson het opnam tegen de twee beste Jeopardy-kandidaten uit de geschiedenis. Het evenement trok veel media-aandacht, waarbij de vergelijking met de historische overwinning van schaakcomputer Deep Blue op wereldkampioen Kasparov niet werd geschuwd. Watson won glansrijk.

Gosse Bouma
Informatiekunde
Rijksuniversiteit
Groningen

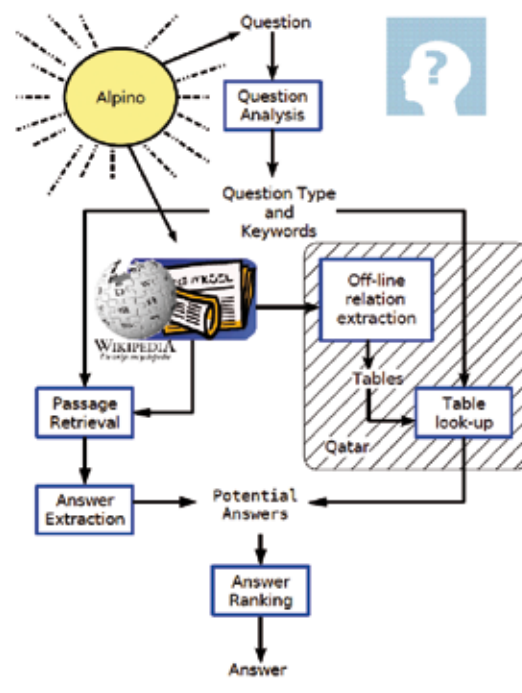
Vraag-antwoordsystemen zoals Watson, die hun kennis halen uit encyclopedieën, waaronder natuurlijk Wikipedia, andere boeken en tijdschriften, en talloze andere meestal weinig gestructureerde bronnen op het web, zijn gebouwd rond een slimme combinatie van information retrieval, information extraction, en natuurlijke taalverwerking. Omdat grote hoeveelheden informatie moeten worden doorzocht en de inhoud van deze informatie geanalyseerd moet worden, is veel rekenkracht nodig. In de publiciteit van IBM wordt op het laatste, om begrijpelijke reden, nogal de nadruk op gelegd. Hoewel de overwinning van Watson ongetwijfeld ook een overwinning is voor natuurlijketaaltechnologie, is het nogal lastig te zien hoeveel taaltechnologie nu echt nodig is voor een systeem als Watson. Een goede lakmoestest in dit geval is de vraag te stellen wat het zou betekenen een Watson voor het Nederlands te maken.

Joost

In figuur 1 wordt de architectuur geschetst van Joost, een vraag-antwoordsysteem voor het Nederlands. Joost kan vragen beantwoorden door te zoeken in tekstbestanden, zoals (de Nederlandstalige) Wikipedia of krantentekst. Joost is online beschikbaar via www.vraaghetaanjoost.nl.

De modules van het systeem zijn redelijk generiek. In vrijwel alle delen van het systeem speelt op het Nederlands toegesneden taaltechnologie een rol. In de analyse van de vraag bijvoorbeeld, wordt een vraagzin taalkundig geanalyseerd en wordt bepaald welke inhoudswoorden in de vraag voorkomen, welke namen, en welke grammaticale relaties tussen woorden bestaan.

Op basis hiervan bepaalt Joost het vraagtype (definitie, geboortedatum, hoofdstad, etc.) en het antwoordtype (datum, locatie, hoeveelheid, etc.). Voor de taalkundige analyse maken we gebruik van Alpino, een robuuste taalkundige ontleder voor het Nederlands, waar in verschillende STEVIN-projecten gebruik van wordt gemaakt.



Figuur 1: Architectuur van Joost, een Nederlandstalig vraag-antwoordsysteem

Het belang van taalkundige informatie kan worden geïllustreerd door de volgende twee vragen:

1. Waar leven cicaden?
2. Waar leven cicaden van?

In het eerste geval wordt gevraagd naar een locatie, terwijl in het tweede geval een zelfstandig naamwoord als antwoord wordt verwacht. Door te kijken naar meer dan het vraagwoord alleen kan dit verschil worden geïdentificeerd.

Alle tekstcollecties waarin kan worden gezocht naar een antwoord zijn ook volledig door Alpino geanalyseerd. Dit maakt het o.a. mogelijk in passage retrieval (dat is information retrieval op het niveau van alinea's of zelfs individuele zinnen) gebruik te maken van taalkundige informatie. Het antwoord op vraag 1 is bijvoorbeeld een locatie en dus zijn alinea's die geen locatie bevatten minder relevant. Voor het antwoord op vraag 2 zijn

zinnen die niet alleen het werkwoord 'leven' en het voorzetsel 'van' bevatten, maar waar 'van' bovendien een voorzetselvoorwerp inleidt bij 'leven' extra relevant. De module voor antwoord-extractie zoekt naar een daadwerkelijk antwoord binnen relevante alinea's. Hierbij scoren antwoorden gevonden in zinnen die dezelfde grammaticale structuur hebben als de vraag extra hoog.

De module voor offline relation extraction, tenslotte, wordt gebruikt om antwoord te geven op veelvoorkomende vraagtypes. De tekstcollectie wordt vooraf doorzocht op land-hoofdstad of persoon-geboortedatum paren, die worden opgeslagen in een tabel. Bij een vraag naar een hoofdstad of geboortedatum kan vervolgens het antwoord uit de tabel worden gegeven, en kunnen de stappen voor passage retrieval en answer extraction worden overgeslagen. Doordat de tekstcollecties volledig taalkundig geanalyseerd zijn kunnen zoekpatronen geformuleerd worden met behulp van grammaticale relaties.

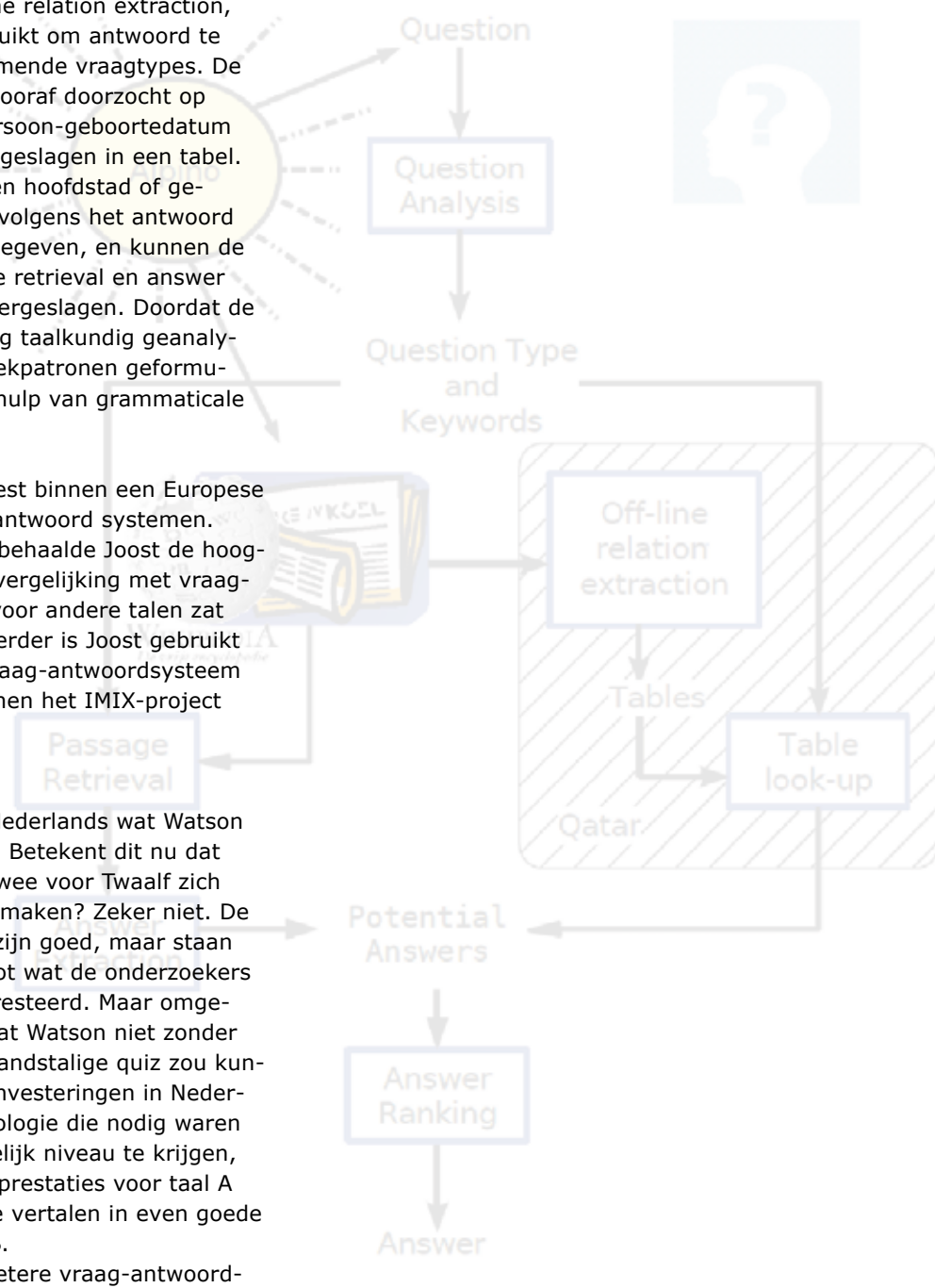
Joost is driemaal getest binnen een Europese evaluatie van vraag-antwoord systemen. Voor het Nederlands behaalde Joost de hoogste scores en ook in vergelijking met vraag-antwoord systemen voor andere talen zat Joost in de top-10. Verder is Joost gebruikt voor het medische vraag-antwoordsysteem dat is ontwikkeld binnen het IMIX-project (NWO).

Twee voor Twaalf

Joost doet voor het Nederlands wat Watson voor het Engels doet. Betekent dit nu dat de deelnemers aan Twee voor Twaalf zich zorgen moeten gaan maken? Zeker niet. De prestaties van Joost zijn goed, maar staan in geen verhouding tot wat de onderzoekers van IBM hebben gepresteerd. Maar omgekeerd is het ook zo dat Watson niet zonder meer aan een Nederlandstalige quiz zou kunnen deelnemen. De investeringen in Nederlandstalige taaltechnologie die nodig waren om Joost op een redelijk niveau te krijgen, laten zien dat goede prestaties voor taal A niet eenvoudig zijn te vertalen in even goede prestaties voor taal B.

Wat is er nodig om betere vraag-antwoord-systemen voor het Nederlands te ontwikkelen? Meer onderzoek natuurlijk. Aan de Radboud Universiteit werkt Suzan Verberne aan betere informatie-extractie voor vraag-antwoord systemen. (Zie haar artikel hierover elders in deze DIXIT, red.) Naast zoeken in tekst lijkt ook het aloude bevragen in natuurlijke taal van databases weer aan

belang te winnen. De informatie uit Wikipedia is bijvoorbeeld ook beschikbaar als online bevroagbare semantic web database (dbpedia). Hiermee kunnen zeer veel vragen beantwoord worden, mits je weet wat de juiste database query is. Om de toegankelijkheid voor grotere groepen gebruikers te verbeteren werken we momenteel aan een versie van Joost die Nederlandstalige vragen omzet in dbpedia (sparql) queries.



Taaltechnologie elimineert obstakels voor recruitment van talent

Taaltechnologie maakt het mogelijk om vraag en aanbod op de arbeidsmarkt bij elkaar te brengen op basis van ongestructureerde vacatures en cv's. Textkernel ontwikkelt basistechnologie voor het begrijpen van tekst en semantisch zoeken. De productfocus is: het leven van recruiters en sollicitanten eenvoudiger maken. In de komende 'war on talent' zal dit soort technologie voor werkgevers steeds belangrijker worden.

Jakub Zavrel
Textkernel

Om effectief te zoeken moet je 'begrijpen'. Anno 2011 zijn er nog steeds veel situaties waar informatie uit ongestructureerde documenten pas echt bruikbaar wordt na het invoeren ervan in een formulier. Op dat moment wordt het pas, doordat de velden van het formulier een betekenis en structuur geven aan de informatie, echt mogelijk om effectief te zoeken, te filteren en te analyseren. Je kunt nu eenmaal op combinaties van domeinspecifieke criteria veel krachtiger zoeken en matchen dan wanneer je puur afgaat op het wel of niet voorkomen van losse woorden, aangezien deze hetzij te ambigu of te gevarieerd kunnen zijn.

Met het overtikken van informatie in formulieren wordt echter ook heel erg veel tijd verspild. De informatie is op zich in het originele document al 'beschikbaar' en het overtikken is, gezien de huidige stand van de techniek, ook steeds vaker overbodig. Door gebruik van taaltechnologie kunnen voor vele types ongestructureerde documenten nauwkeurige automatische herkenners gebouwd worden die het merendeel van de informatie zelf kunnen 'begrijpen'.

Van machine learning R&D naar oplossen voor recruiters

Het gebruik van machine learning maakt de kennisacquisitie voor het bouwen van goed werkende systemen veel goedkoper en maakt het ook mogelijk om eenvoudig met meertaligheid om te gaan. De kennis in het systeem wordt immers met statistische modellen geleerd uit door native speakers gelabelde voorbeelden. Het maken van een systeem voor een nieuwe taal kost zo in de praktijk slechts weken, en dat is werkend vanuit een kleine thuismarkt als Nederland een belangrijk voordeel.

Textkernel heeft in de tien jaar van haar bestaan een scala aan interessante taaltechnologieprojecten opgeleverd, onder andere voor Unilever, Funda, TNT Post, Buzzcapture, Kadaster en KVSa. Deze projecten kenmerken zich door grote volumes ongestructureerde

tekst en de noodzaak om, op maat, een zeer accuraat werkende domeinspecifieke semantische herkenningseengine te leveren.

Curriculum vitae's en vacatures zijn typische voorbeelden van semantic search. Wanneer recruiters de beste kandidaat willen vinden voor een vacature hebben zij aan de ene kant een functiebeschrijving met voorgaande studie- en werkervaring, vaardigheden, talen, werklocatie, arbeidsvoorwaarden, carrièrekansen en andere criteria. Aan de andere kant een grote set documenten waarin kandidaten met motivatiebrieven, cv's en sociale netwerkprofielen uitleggen hoe ze aan de criteria voor de baan voldoen. De bemiddelaar die een werkzoekende probeert te matchen aan alle vacatures die openstaan zal ook op vele criteria filteren, wegen en zoeken. Het echte matchen, het selecteren van de juiste kandidaat voor zijn of haar ideale baan, zal altijd mensenwerk blijven, maar technologie kan wel de meest tijdrovende obstakels in het proces wegnemen.

Het elimineren van barrières in recruitment

Wanneer mensen solliciteren, dan hebben ze over het algemeen hun cv als digitaal bestand bij de hand. Dit is geen gestandaardiseerd bestand, maar een free format document. Cv-extractie van Textkernel maakt het mogelijk om dit document te gebruiken om het online sollicitatieformulier volautomatisch te vullen met de gegevens uit het cv. Dit bespaart op het 20 tot 30 minuten overtikken van informatie tot 90% van de tijd. Werkgevers die de technologie inzetten lukt het om een aanzienlijk betere ervaring voor de kandidaat realiseren, met als gevolg dat de conversie van bezoekers naar sollicitanten sterk verhoogd wordt. Daarnaast is de resulterende informatie vollediger en consistentier ingevuld en de recruiter kan daardoor effectiever zoeken. Door de basis in machine learning is de oplossing al in een dozijn verschillende talen beschikbaar.

Nederland heeft een sterk ontwikkelde uit-

Een online reputatiemonitor

Iedereen wordt beïnvloed door zijn omgeving. Aanbevelingen van vrienden, familie of collega's bepalen voor een belangrijk deel welke films we zien, welke boeken we lezen en in welke restaurants we eten. Onderzoek wijst zelfs uit dat consumenten aanbevelingen van anderen als de meest betrouwbare vorm van reclame ervaren. Tot voor kort beperkte mond-tot-mond reclame zich tot directe, gesproken communicatie op straat, op het werk of in de kroeg, maar tegenwoordig zijn online platforms minstens zo belangrijk geworden. Blogs, forums en social media sites als Twitter, Facebook en Hyves spelen een steeds belangrijker rol in het bepalen van het koopgedrag van de consument. Deze platforms stimuleren het verspreiden en delen van informatie, waarmee de potentiële impact van een individuele mening sterk vergroot wordt: een mening kan zich razendsnel verspreiden en daarmee direct invloed hebben op de reputatie van bedrijven, overheden en individuen.

**Thijs Westerveld,
Stefan de Bruijn,
Arthur van
Bunningen,
Rolf Schellenberger en
Sylvain Perenes**
Teezir

Wel of niet reageren?

Voor bedrijven kan het lastig zijn om te bepalen of, wanneer en op welke manier te reageren op de continue stroom van meningen die online gedeeld worden. Inzicht in de meest in het oog springende sentimenten, de belangrijkste bronnen en de meest invloedrijke auteurs kan helpen bij het bepalen van de juiste reactie. Teezir's dashboards eCare en Webcare kunnen daarbij helpen. Deze tonen niet alleen de individuele berichten, maar ook samenvattingsstatistieken zoals bijvoorbeeld hoeveel er in totaal gesproken is over een onderwerp, wat het gemiddelde sentiment is en in welke bronnen en bij welke auteurs de discussie plaats vindt.

eCare

Met eCare kunnen gebruikers leren wat er online over bepaalde onderwerpen – de zoekvraag - gezegd wordt. Zo'n zoekvraag kan bestaan uit losse termen, frasen en Boolean operatoren. Met behulp van filters kan ingezoomd worden op een bepaalde periode, een specifieke bron, of een individuele auteur, en met de verschillende widgets, grafieken en overzichten maken gebruikers vervolgens hun eigen dashboard en delen deze met collega's. Het dashboard biedt zowel overzichtsstatistieken ("Wat is de verhouding tussen positief en negatief sentiment rondom ons product?") als gedetailleerde informatie ("Welke auteurs zijn het meest negatief en wat zeggen ze precies?"). Enkele voorbeelden:

- Documentresultaten: een lijst met de voor de zoekvraag gevonden documenten, gesorteerd op datum, relevantie of sentiment
- Sentimentoverzicht: een weergave van de sentimentverdeling in de gevonden documenten
- Tijdslijn: een overzicht van de ontwikkeling van volume en sentiment over de tijd (Figuur 1)
- Termanalyse: de meest onderscheidende

termen in de gevonden documenten, plus of ze vooral in positieve of negatieve documenten voorkomen (Figuur 2)

- De verdeling van volume of sentiment over verschillende brontypen (nieuws, blog, forum, etc.), over bronnen of over auteurs

eCare stuurt emailalerts als het sentiment onder een door de gebruiker gedefinieerde grens zakt. De resultaten kunnen geëxporteerd worden om zelf verder te analyseren of te combineren met externe data. De eCare API maakt het mogelijk om de data volledig te integreren in eigen analysesystemen.

Webcare

Webcare (Figuur 3) is erop gericht om snel en effectief te kunnen reageren op individuele binnenkomende berichten. In een dashboard



Figuur 1: Tijdslijnvergelijking van resultaten voor Ziggo en KPN tussen januari en november 2010

configureren gebruikers zoekprofielen en monitoren ze verschillende datastromen (Twitter, Facebook, forums, blogs). De gevonden documenten worden gesorteerd op tijdstip en de met de beschikbare metadata getoond. Het is mogelijk om direct te reageren, maar ook om



Figuur 2: Woordenwolk met de meest onderscheidende termen rondom de zoekvraag Ziggo, inclusief de context

berichten toe te kennen aan specifieke gebruikers voor een reactie op een later tijdstip. De belangrijkste social media zijn in het dashboard geïntegreerd: reacties op Twitter en Facebook kunnen direct vanuit de dashboards worden geplaatst. Statistieken over reactietijden, de verdeling van taken over gebruikers en over categorieën worden verzameld en zijn beschikbaar in heldere overzichten. Met de webcare dashboards monitor je een brede set bronnen en handel je de volledige engagement workflow af in één en dezelfde tool. Er zijn twee demonstratiesystemen beschikbaar: Whorules (www.whorules.nl) en Influencer Insights (<http://influencerinsights.teezir.com>).

Crawling

De berichten in eCare en Webcare worden verzameld door zogeheten crawlers. Deze monitoren continue de belangrijkste Nederlandstalige blogs, forums, nieuwsbronnen en social media platforms. De crawlers zijn specifiek getraind voor het doorzoeken van de verschillende brontypen (blogs, forums, nieuwssites). Op basis van voorbeelden hebben ze de karakteristieken geleerd van de links die gevolgd moeten worden en van de elementen die opgeslagen dienen te worden. Zo wordt ervoor gezorgd dat de data schoon blijft: de berichten met de bijbehorende metadata worden opgeslagen, maar menustructuren, advertenties en andere verstoorde onderdelen van een webpagina worden genegeerd. Ook bij veranderingen van een site blijven de crawlers in staat om de juiste data te verzamelen. Naast crawlers worden ook RSS-feeds en social media APIs ingezet om relevante berichten te verzamelen.

Sentimentanalyse

Alle verzamelde berichten worden onderworpen aan automatische sentimentanalyse. Dit

gebeurt op basis van lijsten van termen waarvan bekend is dat ze een duidelijk positief of negatief sentiment dragen. Eerst bepaalt een Part-of-Speech tagger wat de rol is van een woord in de zin. Op basis van de gevonden woordsoort wordt vervolgens vastgesteld of een woord in de bewuste context daadwerkelijk sentiment draagt. Zo is de term naar bijvoorbeeld negatief als bijwoord ("Mijn bezoek aan de tandarts was naar"), maar neutraal als voorzetsel ("Ik ga naar de tandarts"). Verder wordt onderzocht of er modifiers aanwezig zijn die het sentiment versterken, afzwakken of omkeren. Deze analyse is nodig voor de juiste verwerking van frasen zoals "... tamelijk slechte service..." en "...niet zo goed geholpen...". De sentimentsscore voor het gehele document wordt vervolgens bepaald op basis van de scores van de gevonden sentimentdragende fragmenten.

De oorspronkelijke lijst met positieve en negatieve termen komt van het Instituut voor de Nederlandse Lexicografie. In de afgelopen jaren is deze lijst echter constant aangepast en uitgebreid, zodat Teezir overweg kan met het constant veranderende taalgebruik in social media en in andere online bronnen.

Index

Voor het berekenen van samenvattingsstatistieken over volume en sentiment van de relevante documenten, is het nodig om relevantiescores en sentimentsscores te combineren. Om deze combinaties efficiënt te kunnen berekenen heeft Teezir eigen indexstructuren ontwikkeld, waarin relevantiescores en sentimentsscores van documenten snel toegankelijk zijn. Bij het toevoegen van nieuwe documenten aan de collectie en bij het zoeken binnen de bestaande collectie, wordt gebruik gemaakt van dezelfde indexbestanden. Daardoor is het niet nodig het zoeken te onderbreken of indexen te vervangen om toegang tot nieuwe berichten te verkrijgen. De nieuw toegevoegde data wordt eerst in het geheugen opgeslagen en pas later naar disk weggeschreven. Een zorgvuldig ontworpen recovery-proces zorgt ervoor dat de nog niet weggeschreven indexdata hersteld kan worden wanneer er iets misgaat tijdens het indexeerproces.



Figuur 3: Teezir Webcare

Naam in Bedrijf: datakwaliteit van relatiegegevens

Zeg me wie je vrienden zijn en ik zal zeggen wie je bent

Een bekend gezegde, waar een hoop over te zeggen valt – maar in ieder geval veronderstelt het dat je kunt vertellen wie je vrienden zijn. Ok, daarvoor hebben we natuurlijk een adressenboek. Daar staan de meeste vrienden wel in. Behalve dan de meest recente, die staan alleen in de smartphone. Sommige staan in beide en als je dan het huisadres wilt hebben, dan moet je het boekje hebben, maar het e-mailadres, dat staat alleen in de telefoon. O ja, en die ene vriendin die net verhuisd is, die staat nog met haar oude adres in het boekje. Maar in de telefoon staat wel het nieuwe adres. En dan hebben we ook nog een stel LinkedIn-connecties; daar hebben we geen adressen bij, maar dat is daar ook niet zo belangrijk. Trouwens, over adressen gesproken – die ene vriend had het toch pas over een gemeentelijke herindeling, waardoor hij ineens een andere postcode had?

Emil van den Berg
Human
Inference BV

En dan gebeurt het onvermijdelijke: een bruiloft, een geboorte, hoe dan ook: we willen iedereen een kaartje sturen. Dat wordt een avondje buffelen.

Vermenigvuldig dit scenario met een factor 10.000 of 100.000, en dan kom je in de buurt bij wat er speelt wanneer een groot bedrijf z'n klanten een kerstkaart wil sturen. Of wanneer twee banken fuseren. Of wanneer een verzekeringsbedrijf wil controleren of iemand die een nieuwe verzekering wil afsluiten, niet in een zwarte lijst van fraudeurs voorkomt.

Dit is het probleemgebied waarin Human Inference werkt. Het levert software om dubbele voorkomens van dezelfde persoon (of hetzelfde bedrijf) op te sporen, om adressen te verbeteren aan de hand van postcodetabellen, om hoofdletters goed te zetten en dergelijke [1].

Bij het realiseren van deze functionaliteit speelt taaltechnologie een belangrijke rol. Daarbij gaat het niet om het interpreteren van zinnen, maar om het interpreteren van de contactgegevens, zoals (bedrijfs)namen, adressen, telefoonnummers en e-mailadressen. Je zou er natuurlijk over kunnen twisten in hoeverre we het bij dit soort dingen over menselijke taal hebben, maar in ieder geval: het is mogelijk grammatica's te schrijven en die door de computer te laten gebruiken om die gegevens te ontleiden. En net als bij gewone zinnen is de grammatica met name van belang bij woorden die meerdere betekenissen kunnen hebben. Bijvoorbeeld, in 'Jan Advocaat Jr.' betekent 'Advocaat' wat anders dan in 'Jan Jansen, Advocaat'.

Bij de producten van Human Inference speelt die taaltechnologie verschillende rollen: er is functionaliteit die eigenlijk niet goed denkbaar is zonder die technologie; en daarnaast is er functionaliteit die ook zonder taaltechnologie kan bestaan, maar waarbij die technologie wordt gebruikt om de kwaliteit te verhogen. Van beide varianten zal hier een beeld geschetst worden.

Taaltechnologie als fundament

Bij het opslaan van klantgegevens zijn er verschillende mogelijkheden om structuur in de gegevens aan te brengen: je zou ervoor kunnen kiezen om de complete naam in één veld op te slaan, en het hele adres in een ander veld. Maar meestal wordt een wat meer gedetailleerde structuur gebruikt: de voornamen worden gescheiden van de achternamen opgeslagen en de eerste regel van het adres wordt gescheiden van de tweede adresregel. Of nog gedetailleerder: de voegsels worden van de achternaam gescheiden en het huisnummer wordt apart opgeslagen.



Als er nu data vanuit een minder gedetailleerde bron naar een meer gedetailleerde bron moet worden overgebracht, moeten de gegevens opgesplitst worden. Om dit te kunnen doen, moeten alle onderdelen van de naam gelabeld worden, zodat uitgezocht kan worden wat in welk veld terecht moet komen.

Een ander probleem is het woekeren met beschikbare ruimte. Als een naam te lang is om in het naamveld te passen, moet hij ingekort worden. Dat kan door simpelweg het overschietende stuk weg te laten, maar daarbij bestaat het risico dat er teveel informatie verloren gaat. Daarom is het belangrijk om goed uit te zoeken welke onderdelen het meest relevant zijn en welke het minst; die laatste kunnen dan als eerste weggelaten kunnen worden. Daarnaast kan er ook ruimte gewonnen worden door woorden af te korten, liefst op zo'n manier dat ze begrijpelijk blijven.

Taaltechnologie als component

Een ander belangrijk onderdeel van datakwaliteit is het bij elkaar vinden van relaties waarvan de gegevens niet helemaal met elkaar overeenstemmen [2]. Normaal gesproken zijn computers wel goed in het bij elkaar vinden van dingen die exact gelijk zijn, maar niet zo goed in het geval van dingen die slechts ongeveer gelijk zijn. Maar mensen maken nu eenmaal fouten bij het invoeren van gegevens - en niet alleen mensen: bij het scannen van al dan niet handgeschreven gegevens ontstaan ook fouten.

Voor het bij elkaar vinden en vergelijken van ongeveer gelijke teksten zijn algoritmes beschikbaar. Die kunnen het aantal verschillen tussen twee teksten vaststellen en zo de mate van gelijkenis uitrekenen. Met behulp van deze algoritmes kun je bijvoorbeeld alle records bij elkaar vinden die voor minstens 80% aan elkaar gelijk zijn, zoals bijvoorbeeld 'kruyf' en 'kruif', omdat op 5 letters maar 1 verschil voorkomt.

Als deze paartjes van records door mensen bekeken worden, valt iets op: er komen paartjes voor die heel hoog scoren, maar toch door mensen als verschillend beschouwd worden, en andersom: sommige paartjes die duidelijk dezelfde persoon betreffen, hebben een lage score gekregen. Bij de menselijke beoordeling spelen meer aspecten een rol dan alleen het aantal verschillende letters.

Als deze aspecten meegenomen kunnen worden in de beoordeling door de computer, wordt het verschil tussen de door de computer gegenereerde beoordeling en de menselijke beoordeling kleiner. Dit heeft voordelen

aan twee kanten: het risico dat je goede paartjes mist, wordt kleiner en het aantal slechte paartjes dat door mensen bekeken moet worden, wordt ook kleiner.

Een duidelijk voorbeeld van zo'n aspect is synoniemie: vervoer is transport en Scheveningen is - binnen een bepaald postcodegebied - uitwisselbaar met Den Haag. Het aantal overeenkomende letters is erg laag, maar mensen hebben - afhankelijk van hoe de rest van de records eruit ziet - erg weinig moeite met synoniemen.

Iets vergelijkbaars geldt voor de uitspraak van woorden: cadeau kun je ook als kado schrijven, en shop kom je ook wel tegen als sjop.

Toekomstige productontwikkelingen

Traditioneel wordt de software van Human Inference bij de klant geïnstalleerd en geïntegreerd in de software van de klant. Voor kleinere bedrijven is deze benadering over het algemeen te kostbaar; om deze bedrijven ook te kunnen bedienen is Human Inference zich nu ook aan het richten op in-the-cloud oplossingen, waarbij de implementatiekosten veel lager zijn, en waarbij betaald kan worden naar rato van het gebruik.

Op taalkundig vlak is er een ontwikkeling gaande waarbij het werk dat voorheen met de hand gedaan werd, steeds meer wordt geautomatiseerd. Dit geldt voor dingen zoals het verzamelen van kennis, het schrijven van grammatica's en regelsets en het correct afstellen van de wegingen van de verschillende aspecten die een rol spelen bij het vergelijken.

Onder andere over dit soort dingen wordt van tijd tot tijd geblogd op www.datavaluetalk.com.

Onder het motto 'voorkomen is beter dan genezen' is Human Inference bezig met het realiseren van nieuwe functionaliteit die helpt om vervuiling te voorkomen tijdens het toevoegen van nieuwe relatiegegevens [3].

¹⁾ Zie: www.humaninference.com/floating-pages/the-value-of-data-quality

²⁾ Voor meer informatie: www.humaninference.nl/producten/data-matching

³⁾ Een demo hiervan is uit te proberen op: demo.easydq.com/ftr/crm/edit

Wordt vertaler een knelpuntberoep?

Door het internet en de globalisering van de maatschappij is de vraag naar vertaling nog nooit zo groot geweest. Een bedrijf maakt een website in één taal en deze website moet beschikbaar gemaakt worden in de talen van al de potentiële klanten. Bepaalde secties die vaak bijgewerkt worden moeten telkens weer zo snel mogelijk vertaald worden. Dit leidt tot een tekort aan vertalers die binnen een beperkte tijdsperiode hoge volumes tekst kwalitatief kunnen vertalen. Op dit vlak kan automatische vertaling een meerwaarde leveren.

**Vincent
Vandeghinste**
KU Leuven

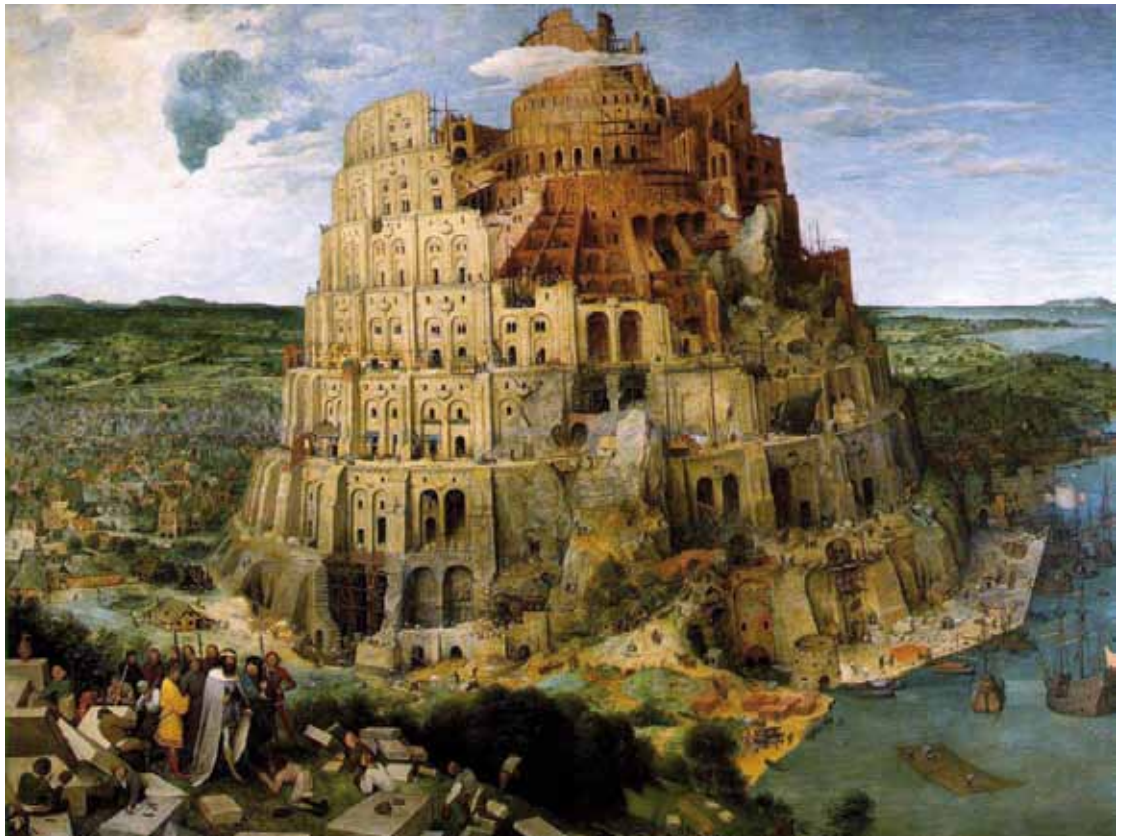
Versnellen

De meeste vertalers gebruiken momenteel tools zoals terminologiedatabanken (gespecialiseerde woordenboeken voor domeinafhankelijke woorden) en vertaalgeheugens (databanken van reeds vertaalde stukken tekst). Beide leiden tot een behoorlijke versnelling van het vertaalproces en een consistentere vertaling dan wanneer men zonder deze hulpmiddelen werkt.

Door gebruik te maken van automatische vertaling kan men de vertaalsnelheid nog verder opdrijven, afhankelijk van onder meer het specifieke taalpaar en het domein van de te vertalen teksten. De beschikbaarheid van een uitgebreide collectie vertaalde teksten binnen het domein speelt een cruciale rol in de nieuwste lichting vertaalsystemen, omdat op basis van deze gegevens geleerd wordt hoe te vertalen.

Integratie

Automatische vertalingen moeten gecorrigeerd worden door een vertaler. Hoewel nog onderzocht wordt hoe dit optimaal gebeurt, is al wel duidelijk dat het gebruik van automatische vertaling pur sang op de nodige weerstand stuit bij veel vertalers. Deze kan gereduceerd worden door automatische vertaling optimaal te integreren in de workflow. Hierdoor vervaagt het onderscheid tussen een fuzzy match uit het vertaalgeheugen (een zin uit het geheugen die niet voor 100% overeenkomt met de te vertalen zin) en een automatische vertaling. Als we bijvoorbeeld een zin in het vertaalgeheugen vinden die enkel verschilt in persoon van het onderwerp (zij vs. wij), en we dit verschil automatisch aanpassen in de gesuggereerde vertaling, gebruiken we dan het vertaalgeheugen of automatische vertaling? Zulke aanpassingen aan het vertaalgeheugen kunnen natuurlijk



nog veel verder gaan. Het komt erop aan te weten wanneer de gegenereerde vertaling niet meer aangeboden moet worden omdat de kwaliteit te laag is, of omdat de vertaler te veel tijd verliest door de gegeven vertaling aan te passen.

Online systemen

De meesten onder ons kennen automatische vertaling via online-vertaalsystemen zoals Google of Bing. Deze systemen zijn gemaakt voor algemene taal en geven slechts een beperkte indicatie van kwaliteit van de technologie indien geoptimaliseerd voor een bepaald taalpaar en tekst domein. Het zijn frasegebaseerde statistische vertaalsystemen en ze maken zo goed als geen gebruik van linguïstische kennis. Ze focussen op het omzetten van groepen van opeenvolgende woorden (frases) van de ene taal in de andere. En net omdat ze geen benul hebben van de grammaticale structuur van de zinnen bestaat de mogelijkheid dat ze zinnen uitvoeren waarin belangrijke delen missen, zoals bijvoorbeeld een werkwoord. Hierdoor kan de hele zin onbegrijpelijk worden.

De nieuwste algoritmes en systemen introduceren hiërarchie en structuur in de analyse van de te vertalen zinnen. Deze analyses zijn al dan niet taalkundig, maar staan in alle geval toe om in hogere mate te abstraheren over de reeds vertaalde teksten en hier meer uit te leren. Zulke systemen zijn nog in ontwikkeling en worden nog nauwelijks industrieel gebruikt.

PaCo-MT

Voor vertaling van en naar het Nederlands hebben de universiteiten van Leuven en Groningen samen met OneLiner bvba het PaCo-MT systeem ontwikkeld dat gebruikmaakt van automatische taalkundige analyses van reeds vertaalde teksten, al dan niet binnen een bepaald domein. Door deze analyses te combineren met modellen voor alignment (welk stuk tekst in de ene taal is een vertaling van welk stuk tekst in de andere taal) kunnen we automatisch linguïstische vertaalregels afleiden. Aan de hand van deze regels kunnen we vertalingen genereren die beter zijn dan vertalingen die gegenereerd worden door frasegebaseerde statistische vertaalsystemen. Hoewel PaCo-MT nog niet kan wedijveren met systemen à la Google of Bing genereert het in sommige gevallen vertalingen die correcter zijn dan wat door deze industriële vertaalsystemen gegenereerd wordt. Zo wordt bv. de zin "vanaf 2010 zamelen we 75 procent van het huishoudelijk afval selectief in" door PaCo-MT naar het Engels vertaald

als "from 2010 we selectively collect 75 % of the domestic waste". Als we dezelfde zin laten vertalen door Bing of Google, dan zien we dat 'selectively' helemaal achteraan eindigt: "from 2010 we collect 75 % of household waste selectively". Voor deze systemen is het theoretisch onmogelijk om het woord 'selectively' te plaatsen bij het werkwoord waar het thuishoort. Hiërarchische systemen zoals PaCo-MT hebben hier geen probleem mee. Als dergelijke hiërarchische systemen voldoende data krijgen in het domein dat vertaald moet worden en op een weloverwogen manier geïntegreerd worden in de workflow van vertalers dan kan het vertaalproces in die mate versneld worden dat vertaler geen knelpuntberoep wordt.

Spraaktechnologie voor inclusive design

TNO gebruikt Informatie en Communicatie Technologie (ICT) om producten en diensten toegankelijk te maken voor iedereen, dus ook voor speciale doelgroepen zoals gebruikers met verminderde cognitieve vaardigheden. Hierbij zetten we vaak spraaktechnologie in. Een voorbeeld hiervan is de elektronische reisassistent voor mensen met een verstandelijke beperking. Deze reisassistent kan relevante informatie voorlezen door gebruik te maken van spraaksynthese. Een ander voorbeeld is een elektronische coach voor kinderen en ouderen. Door de inzet van spraakherkenning en spraaksynthese kan op een natuurlijke manier met de e-coach gecommuniceerd worden.

**Anita Cremers,
Judith Kessens
en Mark
Neerincx**
TNO

Inclusive design

TNO gebruikt een ontwerpmethodiek die inclusive design wordt genoemd. Bij deze methode wordt de gebruiker vanaf het begin in het ontwerpproces betrokken. De methodiek bestaat uit een aantal onderdelen: een analyse van de beschikbare technologieën (waaronder spraaktechnologie), een inventarisatie van de problemen waar gebruikers tegenaan lopen, de gebruikerseisen die ze stellen en de context waarin het product of de dienst wordt gebruikt. Het resultaat is een lijst van ontwerpspecificaties die in een iteratief proces worden verfijnd. Dit wordt gedaan door gebruikers te laten werken met een prototype van het systeem en hen om feedback te vragen. Dit eerste prototype hoeft niet volledig te werken; technologie kan gesimuleerd worden of papieren versies van het ontwerp worden aan gebruikers voorgelegd. Op basis van de feedback van gebruikers wordt het ontwerp aangepast en het aangepaste ontwerp wordt opnieuw aan gebruikers voorgelegd. Dit proces wordt een aantal keer herhaald.

Spraaktechnologie

Interactie tussen de gebruiker en een product of dienst zal bijna altijd plaatsvinden via meerdere modaliteiten. Zo kan de gebruiker een sy-

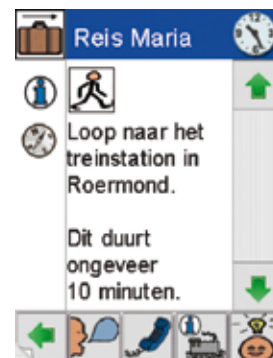
steem bedienen met gebaren, via knoppen of aanraakschermen en kan het systeem informatie teruggeven in de vorm van tekst, iconen, vibraties en geluid. Juist voor gebruikers met beperkingen in één modaliteit kan het aanbieden van informatie via een andere modaliteit meerwaarde geven. Het voorlezen van teksten kan bijvoorbeeld ondersteuning bieden aan gebruikers die slechtziend zijn of die moeite hebben met het lezen van teksten. Ook kan op basis van spraak informatie afgeleid worden over de gebruiker (zoals emotie, taal, taalvaardigheidsniveau). Deze informatie maakt het mogelijk om het systeem beter aan te passen aan zijn/haar specifieke behoeftes en wensen. Spraak is een natuurlijke en intuïtieve manier om met een systeem te interacteren. Onderstaande voorbeelden laten zien hoe spraak als extra modaliteit ingezet kan worden voor het toegankelijk maken van producten en diensten.

De elektronische reispartner

De elektronische reispartner is ontwikkeld voor mensen met een verstandelijke beperking en maakt het voor deze mensen mogelijk om zelfstandig met het openbaar vervoer te reizen. Een eerste versie van de elektronische reispartner is gebaseerd op bestaande ontwerprichtlijnen en op criteria die door represen-



Figuur 1: Computerprogramma gebruikt door begeleider.



Figuur 2: Voorbeeld van het scherm van de PDA die tijdens de reis gebruikt wordt

tanten van de doelgroep en hun begeleiders zijn opgesteld. Het concept bestaat uit twee onderdelen: een computerprogramma dat de begeleider gebruikt om de reis te plannen en een gebruikersprofiel op te stellen (figuur 1) en een draagbare PDA voor de verstandelijk beperkte reiziger die in simpele bewoordingen informatie en instructies geeft over de reis (figuur 2). Voor iedere stap in het reisschema wordt een apart scherm gebruikt en een audiosignaal geeft aan dat er nieuwe informatie op het scherm verschijnt. De PDA kan bediend worden door de iconen op het scherm aan te raken. Tekstuele instructies worden op het scherm getoond, maar kunnen ook voorgelezen worden door de icoon met de tekstballon aan te raken.

De elektronische reispartner is geëvalueerd door drie gebruikers uit de doelgroep tijdens een reis met het openbaar vervoer. Hieruit blijkt dat de alle proefpersonen in staat waren om de reis zelfstandig te maken, maar dat het systeem op een aantal punten verbeterd kan worden. De spraaksynthese als alternatief voor het lezen van teksten werd als positief ervaren, maar verbeterpunten zijn: een betere kwaliteit van de spraaksynthese, de mogelijkheid het volume aan te passen en het aanbieden van de spraak via een koptelefoon voor een betere verstaanbaarheid en meer privacy.

De elektronische coach voor chronisch zieken In veel westerse landen neemt het aantal

wordt (met een Wizard-of-Oz methode). De iCat kan emoties tonen via gezichtsuitdrukking en spraakgebruik.

In het prototype van de e-coach worden patiënten geleerd en gemotiveerd om een gezonde leefwijze te volgen, door het afnemen van interviews en door het bijhouden van een elektronisch dagboek dat gezondheidsgesegunde informatie bevat. Alle dialogen zijn toegespitst op een specifieke gebruikscontext (bijv. simpele vragen), zodat we verwachten dat de prestaties van automatische spraakherkenning voldoende zijn om praktisch toe te passen. Daarnaast verwachten we ook dat het mogelijk is om belangrijke aspecten van de communicatie automatisch af te leiden (zoals wie wanneer spreekt, en wat de emoties zijn). Gebaseerd op deze informatie kan de robot op een gepaste manier reageren. Door de e-coach als een soort "buddy" voor zowel serieuze als ontspannende activiteiten in te zetten, wordt deze aantrekkelijk gemaakt voor kinderen.

Uit evaluaties blijkt dat ouderen de e-coach als empathisch ervaren. De kinderen zijn meer geneigd om met de e-coach te interacteren dan met standaard desktop-applicaties. De communicatie is efficiënter en wordt als plezieriger ervaren, onder andere omdat het geen computervaardigheden vereist zoals lezen en typen. Het blijkt belangrijk te zijn dat de iCat zich in alle modaliteiten op een consistente manier uitdrukt en dat deze uitdrukkingsvorm past bij de situatie/context. Over het algemeen worden gelijke resultaten gevonden voor een virtuele iCat getoond op een computerscherm.

Conclusie

Onze praktijkvoorbeelden laten zien dat ICT hulpmiddelen (zoals de elektronische reispartner en de e-coach) de toegankelijkheid vergroten van producten en diensten voor speciale doelgroepen. Spraaktechnologie is hierbij een belangrijke ondersteunende technologie omdat het specifieke cognitieve vaardigheden kan ondersteunen, maar ook omdat het een natuurlijke en intuïtieve manier van communiceren mogelijk maakt.



Figuur 3: De iCat als virtuele assistent voor ouderen.

patiënten met chronische ziekten zoals obesitas en diabetes toe. Het doel van een elektronische coach (e-coach) is om patiënten thuis op te leiden, te informeren, te motiveren en te instrueren. Daarnaast dient het als communicatielaag tussen de patiënt en de specialist op afstand. Dergelijke ondersteuning is vooral aantrekkelijk voor ouderen en kinderen, omdat deze doelgroepen minder geneigd zijn om desktopachtige hulpmiddelen te gebruiken. Om deze reden is een spraakgestuurde robot ontwikkeld, gebaseerd op de iCat van Philips, zie Figuur 3 en 4. In de robot is automatische spraaksynthese geïmplementeerd, terwijl de spraakherkenning momenteel nog gesimuleerd



Figuur 4: De iCat als virtuele assistent voor kinderen.

Taal- en spraaktechnologie en navigatie: een gelukkig huwelijk

De toepassing van taal- en spraaktechnologie is een belangrijk onderdeel in het ontwerp en de ontwikkeling van navigatieproducten. Een nuance in betekenis van een label kan het verschil zijn tussen een gebruiksvriendelijke of onbruikbare interface. Een stem die niet alleen foutloos spreekt, maar ook aangenaam is om naar te luisteren, zorgt ervoor dat route instructies begrepen en als plezierig ervaren worden. Verkeersveiligheid neemt toe doordat bepaalde taken door middel van spraaktechnologie sneller en makkelijker kunnen worden uitgevoerd en de bestuurder zijn aandacht bij het verkeer kan houden.

Ingrid van de Bovenkamp,
Ingrid Halters
TomTom

Taal- en spraaktechnologie binnen TomTom

TomTom maakt gebruik van bestaande commerciële spraaksynthese en -herkenningssoftware voor handsfree bediening en flexibele audio output. De afdeling 'User Experience Design' werkt aan de toepassing van taal- en spraaktechnologieën in user interfaces. Taalkundigen testen en verbeteren de spraaksynthese en interaction designers maken het spraak-interactieontwerp.

Een andere R&D afdeling binnen TomTom houdt zich bezig met de creatie van 'voice maps'. Dit zijn fonetische databases die digitale kaarten aanvullen. Met deze fonetische transcripties is het mogelijk om spraaksynthese systemen steden, straatnamen en Points-of-Interest (POI's) correct uit te laten spreken, terwijl spraakherkenningssystemen gesproken locaties kunnen ontcijferen met behulp van deze data.

In 2006 kwam het eerste TomTom navigatieproduct met spraaksynthese (Text-To-Speech, TTS) op de markt en een jaar later werd

spraakherkenning (Automatic Speech Recognition, ASR) geïntroduceerd.

Spraaksynthese – "Volg de A4 gedurende 3,8 kilometer"

De hoogstaande spraaksynthese software waar TomTom gebruik van maakt, is gebaseerd op de zogenaamde concatenative synthesis. De geluidssegmenten van de uiteindelijke audio output worden via 'unit selection' uit een database van menselijke geluidssegmenten gehaald, en zorgt zo voor een natuurlijk geluid. Dit kan worden gecombineerd met opgenomen instructies: dezelfde persoon (stem) die als basis dient voor de TTS-stem, neemt voor TomTom bepaalde zinnen op, zodat de uiteindelijke audio output een combinatie is van opnames en synthese.

Gesproken route-instructies zijn een belangrijk onderdeel van de navigatie, dus de kwaliteit van de TTS-stemmen is erg belangrijk. Dit begint al bij selectie van de leverancier van het TTS-systeem en daarna wordt elke taal en stem afzonderlijk onderworpen aan een keuring.



Figuur 1: USA vs. Europa, er is een duidelijk verschil in zichtbaarheid van straatnaamborden.

Eerst wordt de kwaliteit van de computerstem gecontroleerd. Stemmen die als te hoog (onprettig), te laag (niet verstaanbaar in de auto) of mechanisch (vermoeiend om naar te luisteren) worden ervaren, komen niet door de test. Deze testen zijn uiteraard subjectief; sommige gebruikers worden meer door de verstaanbaarheid beïnvloed (begrijp je wat er wordt gezegd?), terwijl andere de natuurlijkheid (hoe menselijk klinkt de stem?) belangrijker vinden.

Naast deze eerste kwaliteitstest wordt de 'front-end processing' geëvalueerd: Worden nummers goed uitgesproken, hoe behandelt het systeem afkortingen, wat doet het met veelvoorkomende woorden van buitenlandse origine, etc.

In 2006 werd TTS op de TomTom-modellen geïntroduceerd: aan de bestaande auditieve route instructies - opgenomen door de persoon achter de TTS-stem - werden straatnamen en wegnummers toegevoegd. Wegwijzer informatie kon worden gebruikt in constructies als 'richting Den Haag'.

Alhoewel het toevoegen van straatnamen in sommige Europese landen vooral een leuke feature is, is het in Noord-Amerika bijna een must: Amerikanen vinden de weg aan de hand van straatnamen, die over het algemeen ook duidelijk staan aangegeven. Deze culturele verschillen zijn bij niet-westerse culturen nog groter: In India hebben zelfs 70% van de straten geen naam en veel huizen geen nummer. Points-of-Interest en het uitspreken daarvan zal belangrijker worden en met de flexibiliteit van TTS-systemen kunnen route instructies als "Ga links bij de tempel" elk moment aan de uitgesproken instructies toegevoegd worden.

Spraaksynthese wordt ook toegepast bij het voorlezen van nuttige informatie tijdens het rijden, zoals het weer op plaats van bestemming, de verkeerssituatie op de route en de afstand tot het dichtstbijzijnde benzinestation.

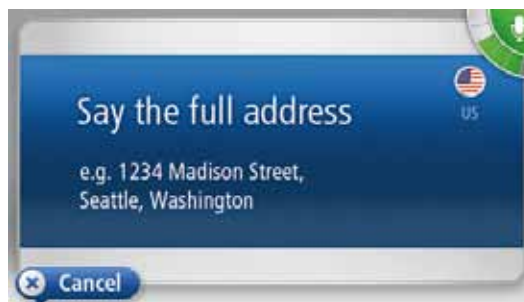
Spraakherkenning – "Breng me naar het dichtstbijzijnde benzinestation, alsjeblieft"

Bij de introductie van spraakherkenning een jaar later, was de 'multi-step address entry' een belangrijke feature: in plaats van adressen intypen kon je eerst stad, dan straat en vervolgens huisnummer inspreken. De herkenningsgraad was hoog en de invoer sneller dan manueel.

Begin 2011 introduceerden we de zogenaamde 'one-shot-address entry': Het hele adres kan in één keer worden gesproken en als dusdanig herkend worden. De spraakinterface volgt niet

langer de grafische user interface als alternatieve invoermethode, maar spraak krijgt z'n eigen ontwerp.

Door de jaren heen is er steeds meer flexibiliteit gekomen in wat een gebruiker kan zeggen. "Navigeer naar" en "Rijd naar" worden beide als commando herkend en zullen hetzelfde resultaat opleveren. De gebruikersvriendelijkheid neemt daardoor toe: de dialoog is intuïtiever en gebruikers hoeven geen lijst van commando's uit het hoofd te leren. Ook kunnen variaties van Points-of-Interest en straatnamen herkend worden. "Navigeer naar John Fitzgerald Kennedy Boulevard" moet hetzelfde resultaat opleveren als "Navigeer naar JFK Boulevard".



Figuur 2: Bij de 'one-shot address entry' kan de gebruiker het hele adres in één keer uitspreken.

Meer factoren spelen een rol bij een goede gebruikservaring, zoals een intelligent helpstelsysteem. En van 100 commando's die de speech engine in 2007 herkende, zijn we naar bijna 1000 gegaan, inclusief de variaties. Het systeem moet het in één keer werken voor gebruikers met verschillende accenten en zonder ingewikkelde trainingsprogramma's. Het onderliggende akoestisch model moet geoptimaliseerd zijn voor auto omgevingsgeluiden.

Bestemming bereikt?

Verschillende onderzoeken wijzen uit dat spraakinvoer veiliger is en minder belastend voor de 'mental workload' dan handmatige invoer. Aangezien op dit moment spraak de enige alternatieve manier is voor manuele invoer op een navigatiesysteem (wie weet wat de toekomst brengt?) en veiligheid een grote rol speelt bij TomTom's technologieën, verwachten we dat het een belangrijke rol blijft spelen in navigatiesoftware. De use cases in de auto lijken oneindig. Van het af laten spelen van je favoriete muziek tot het vinden van een Chinees restaurant in de buurt: de onderliggende technologie ontwikkelt zich in rap tempo en een volledig natuurlijke gesproken dialoog tussen mens en navigatieapparaat is geen science fiction meer.

Taal onderweg

De tijd dat auto's voornamelijk mechanische machines waren die ons van punt A naar punt B hoorden te brengen ligt al lang achter ons. Elektronische snufjes die het comfort en de veiligheid van bestuurder en passagier verbeteren zijn niet meer weg te denken in het huidige aanbod van elke constructeur. Om deze nieuwe functies te bedienen zetten heel wat automerken volop in op automatische spraakherkenning en spraaksynthese.

De iCar

Het is dinsdagochtend. Na een vlug ontbijt stap je in je wagen, klaar voor een nieuwe werkweek, mobieltje op zak. Vlug even de bestemming ingeven van de klant die je wil bezoeken. Het adres had je gelukkig al op een papiertje geschreven: menu kiezen, met de draaiknop stad, straat en nummer ingeven. Je zal er zijn in 53 minuten, zo staat althans op het schermje aangeduid. Je wilde graag even de binnengekomen e-mailberichten lezen, maar dat zal voor de eerstvolgende tankbeurt zijn. Hopelijk is die levering gisteren nog op tijd vertrokken. Na een halfuur verschijnt een gevarendriehoek op het scherm. Even klikken, en lezen dat er een ongeval is gebeurd een eindje verderop. Bijna zelf op de voorligger ingereden. Geen probleem, er wordt een alternatieve route berekend. Vlug even de klant bellen, denk je. Even aan de kant gaan staan om André te bellen. Hij is er niet, maar geen nood. Vlug een SMS'je intypen. Verdorie, denk je, onze jongste typt die dingen toch sneller in. Terug op weg. Je hebt nog een eindje rijden voor de boeg, dus kies je een leuke CD van je CD-lader om de tijd te doden. Toch fijn dat je uit 6 schijfjes tegelijk kan kiezen, denk je.

Of liever dit:

Het is dinsdagochtend. Na een vlug ontbijt stap je in je wagen, klaar voor een nieuwe werkweek, iPhone op zak. Vlug even de bestemming inspreken van de klant die je wil bezoeken: "InfoTec in Hasselt". Je zal er zijn in 53 minuten, zo zegt althans een vriendelijke stem. Onderweg beluister je even de binnengekomen e-mailberichten. Mooi, de levering van gisteren is net op tijd vertrokken. Na een halfuur komt de vriendelijke stem terug: er is een eindje verderop een ongeval gebeurd. Geen probleem, ze stuurt je wel langs een alternatieve route. Vlug even de klant bellen, denk je. Gelukkig heeft je auto het adresboek van je iPhone reeds netjes opgeslagen.

"Bel André van InfoTec", zeg je. Hij is er niet, maar geen nood. Vlug een SMS'je inspreken: "Stuur bericht: ik zal een halfuurtje later zijn. Tot straks. Bart." Even wat muziek kiezen: "Speel Mad About You van Hooverphonic". Toch fijn dat je hele muziekcollectie op je nieuwe iPhone past, denk je.

Nieuwe toepassingen

Allerlei computersystemen proberen ons een zo comfortabel mogelijke rit te bezorgen. Het navigatiesysteem laat ons weten welke weg we best nemen naar onze bestemming, rekening houdend met eventuele ongevallen of files. Ondertussen luisteren we naar onze favoriete muziek, die we meebrachten met onze MP3 speler. Als we dan toch in een file terechtkomen, willen we de tijd nuttig gebruiken en naar klanten bellen, onze e-mail



berichten alvast beluisteren of een SMS-je sturen om te laten weten dat we iets later zullen zijn. Het is zelfs mogelijk geworden om onderweg even op het net te surfen om te zoeken naar een leuk restaurant.

Nieuwe mogelijkheden

Tot nu toe is spraaktechnologie nooit de motivator geweest om specifieke voorzieningen in te bouwen in wagens. De opties die voor bepaalde functies nodig zijn, zoals een microfoon om te kunnen bellen, een audio-installatie om de radio te horen, en recent een

Bart D'hoore
Automotive ASR
Research,
Nuance Communications

internetverbinding om restaurants of hotels te zoeken, hebben ook de weg geëffend voor TST om al deze toepassingen veilig te bedienen. Nu de microfoon er toch is om handenvrij te telefoneren, kunnen we hem evengoed gaan gebruiken om met spraakherkenning ook het nummer te kiezen, of gewoon de naam uit je contactlijst in te spreken. En als dat dan kan, waarom niet een adres inspreken in het navigatiesysteem? Dat is nu immers mogelijk geworden omdat de nodige rekenkracht beschikbaar is, nu die ingebouwd wordt om mooie driedimensionale plaatjes te tekenen op het centrale scherm.

Productdifferentiatie

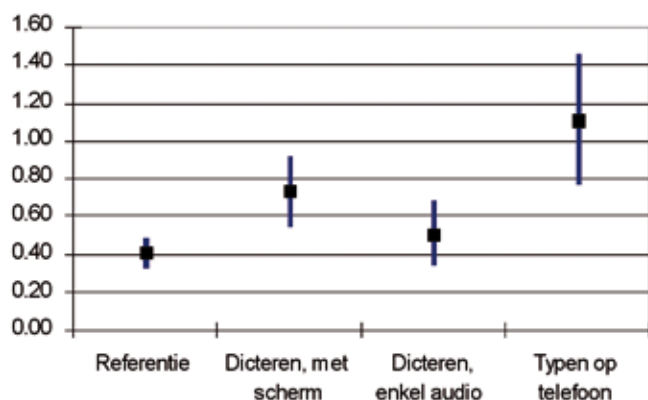
Met draaiknoppen en aanraakscherm alleen is het vaak moeilijk zich een weg te banen doorheen de menubalkjes en opties. Om het gebruiksgemak te verhogen, trekken autobouwers volop de TST kaart. In de Verenigde Staten gebruiken vandaag 6 miljoen mensen het met de stem bestuurd OnStar systeem van General Motors. Concurrent Ford rustte reeds 3 miljoen auto's uit met Ford Sync, een volledig geïntegreerd communicatie- en entertainmentsysteem, waar de stembesturing als essentieel element in advertenties wordt gepromoot. Volgens deze fabrikant kiest 80% van de kopers van een nieuwe wagen ervoor het Sync-systeem te laten installeren. Europese constructeurs gaan volop mee in deze trend. Bovendien wordt de spraakinteractie gezien als een essentieel element dat mee het gezicht van een merk bepaalt. De stem voor de gesproken boodschappen moet in lijn liggen met het imago van de fabrikant. Een Mercedes-stem of BMW-stem moet uniek en herkenbaar zijn, en de degelijkheid of dynamiek van het merk uitstralen.

van de bestuurder afleiden van de weg. Niet handenvrij bellen of SMS'en is reeds in vele landen verboden. Studies tonen aan dat het gebruik van TST in deze context een belangrijke bijdrage kan leveren om de veiligheid te verhogen. Een veel gebruikte methode hierbij is in een simulator nagaan hoever bestuurders van een opgelegd pad afwijken bij het uitvoeren van bepaalde taken. Een recente studie over het sturen van SMS berichten toont aan dat deze taak uitvoeren met spraaktechnologie een stuk veiliger kan. De gemeten afwijking is significant lager bij het dicteren dan bij het intypen van een bericht op de telefoon zelf.

De toekomst

We verwachten dat TST zich verder zal doorzetten in de auto. Dergelijke technologie biedt in deze omgeving een enorme meerwaarde op vlak van gebruiksgemak en veiligheid. De intrede van netwerkverbindingen met hoge snelheid in auto's opent de deur naar toepassingen die ergens in een computercentrum draaien, met heel wat minder rekenkrachtbeperkingen. Gebruik door de grote massa zal ook nieuwe eisen stellen. Een hoge gemiddelde herkenningsgraad zal plaats ruimen voor de eis naar een hoge prestatie voor iedereen. De verwachting dat burgers in de mobiele wereld van vandaag meertalig worden, zal zich ook vertalen in eenzelfde vraag voor spraakgestuurde systemen. Misschien zijn we binnen tien jaar allemaal een beetje Knight Rider.

Gemiddelde trajectafwijking bij SMS'en (m)



Is het veilig?

Verkeersdeskundigen maken zich zorgen over de groei aan toepassingen die de aandacht

Onderzoekers die dit artikel lezen, lezen ook...

Samen met CLST van de Radboud Universiteit Nijmegen voerde de ILK onderzoeksgroep van het Tilburg center for Cognition and Communication (Tilburg University) van 2004 tot en met 2008 het À Propos-project uit. À Propos had tot doel kenniswerkers te ondersteunen in het professionele schrijfproces door relevante leesstof te suggereren tijdens het schrijfproces. Een specifiek voorbeeld is het vinden en suggereren van relevante artikelen voor wetenschappers die bezig zijn met het beschrijven van hun eigen onderzoek.

Toine Bogers
Royal School of
Library and
Information
Science,
Copenhagen,
Denmark
Antal van den
Bosch
CLST Radboud
Universiteit
Nijmegen

Deze pro-actieve automatische suggestie-generator is een voorbeeld van een zogenaamd recommender system. Een recommender system probeert de meest interessante en relevante objecten (zoals wetenschappelijke artikelen) in een collectie te vinden voor een specifieke gebruiker, of een groep gebruikers. Voor het genereren van zulke suggesties kunnen verschillende informatiebronnen gebruikt worden: demografische informatie, meta-data behorende bij de objecten, informatie over gebruikersvoorkeuren in het verleden, etc. Binnen het À Propos-project is gewerkt aan verschillende recommender systems: voor wetenschappelijke literatuur, voor nieuwsartikelen en voor het suggereren van experts.

Als onderdeel van het IOP-MMI-programma (een onderzoekssubsidieprogramma van Agent-schap NL, toen SenterNovem) was het opzetten van samenwerkingen met het bedrijfsleven voor À Propos van groot belang. Een voorbeeld van de praktische toepassing van het onderzoek verricht binnen À Propos, was de samenwerking met CiteULike, een Britse start-up. CiteULike is een zgn. social reference manager, een website die het mogelijk maakt voor haar gebruikers om hun favoriete wetenschappelijke artikelen online op te slaan en te beheren middels een web-interface, in plaats van deze artikelen lokaal op de computer op te slaan. Het 'sociale' aan het onderliggende systeem is dat het alle opgeslagen artikelen toegankelijk maakt voor alle gebruikers. CiteULike-gebruikers hebben de mogelijkheid om de artikelen die ze aan hun profiel hebben toegevoegd te beschrijven met steekwoorden. Deze steekwoorden worden ook wel tags genoemd en vormen een aanvulling bovenop de gebruikelijke metadata zoals titel en samenvatting. Tags, zoals bijvoorbeeld text mining of factorization, kunnen gebruikt worden om de artikelen van een gebruiker te categoriseren en verhogen de toegang tot en de hervindbaarheid van de artikelen. Tags worden beschikbaar gemaakt voor alle gebruikers. Veel gebruikers voegen deels dezelfde artikelen toe, met deels overlap-pende tags. Dit leidt tot het ontstaan van een

rijk netwerk van gebruikers, artikelen en tags, ook wel een folksonomy genoemd.

Voor het vinden van nieuwe, interessante artikelen moet een gebruiker in de klassieke zoekmachine-wereld zelf actie ondernemen: zoektermen intypen, browsen. Om dit proces makkelijker te maken ontwikkelde À Propos-medewerker Toine Bogers in samenwerking met CiteULike een recommender systeem voor wetenschappelijke artikelen in CiteULike. Bogers vergeleek het nut van de verschillende informatiebronnen, zoals de artikelen die door haar gebruikers toegevoegd zijn, metadata en tags. Hieruit bleek dat tags en metadata de beste mogelijkheden bieden om gelijksoortige, interessante artikelen te vinden voor een gebruiker. Artikelen die overeenkomen in de tags en metadata die eraan zijn toegekend leiden tot betere suggesties dan het gebruiken van informatie over gebruikersvoorkeuren in het verleden. De beste suggesties komen voort uit de combinatie van de verschillende algoritmes, zodat de zwaktes van de afzonderlijke algoritmes door elkaar worden opgevangen.

Dit recommender system voor wetenschappelijke artikelen werd in september 2010 toegevoegd op de CiteULike website. De initiële respons op de artikelsuggesties was erg positief, met een gemiddeld acceptatiepercentage van boven de 30%. Dit betekent dat van elke 10 suggesties die het recommender system maakt, er gemiddeld meer dan 3 toe worden gevoegd aan de collectie met artikelen die gebruikers bijhouden op CiteULike. In de toekomst zouden combinaties van meerdere recommendation-algoritmes dit acceptatiepercentage nog verder omhoog moeten kunnen stuwen, zoals al is aangetoond in experimenten die Bogers beschreef in zijn proefschrift.

Links

- <http://ilk.uvt.nl/apropos>
- www.citeulike.org
- <http://ilk.uvt.nl/~toine/phd-thesis>
- www.iva.dk



Figuur 1: Het À Propos recommender system in actie binnen de social bookmarking website CiteULike. De gebruiker wordt gewezen op een aantal artikelen die hij zelf nog niet gebookmarkt heeft en die goed aansluiten bij wat hij al wel in CiteULike heeft ingebracht.

Een spraakgestuurde sprekende applicatie voor materieelonderhoud

Het onderhouden van auto's en vrachtwagens is voor medewerkers van een garage al behoorlijk complex. Bij defensiematerieel wordt deze complexiteit nog eens vergroot doordat de voertuigen vaak speciaal zijn aangepast voor de inzet in zeer divers terrein en omstandigheden (denk ook aan toevoeging van extra bepantsering). Om de onderhoudstaak voor monteurs makkelijker te maken presenteert dit artikel een zeer geavanceerd systeem dat gebruik maakt van virtual reality, spraakherkenning en spraaksynthese.

Project bij Defensie

In november 2009 heeft het consortium bestaande uit BlueTea, Tedopres en Dutcheer de Defensie Innovatie Competitie gewonnen met het idee om een speciaal onderhoudssysteem te ontwikkelen (project SAMM – System for Asset Maintenance Management)¹. Naast spraakherkenning en spraaksynthese is een andere belangrijke component het gebruik van augmented reality². Door augmented reality wordt het met behulp van een speciale bril mogelijk om extra visuele informatie te geven aan de monteur. Neem als voorbeeld het onderhouden van een gepantserde vrachtwagen. Zodra de monteur door de bril naar de vrachtwagen kijkt, verschijnen er in beeld aanwijzingen waar op de vrachtwagen de handelingen moeten worden verricht. Er verschijnt bijvoorbeeld een pijl in beeld waar de motorkap moet worden geopend. Ook is het mogelijk om teksten en animaties in beeld te brengen, zoals een animatie hoe de motorkap geopend en vastgezet moet worden.

Tijdens het project (afgerond in de eerste helft van 2011) zijn diverse onderhoudstaken gedefinieerd en voorzien van scenario's met daarin de uitgesproken handelingen (spraaksynthese), de te tonen animaties (augmented reality) en de grafische interface die in de bril te zien is. De aansturing van de instructies vindt plaats met een aantal makkelijk te onthouden spraakcommando's, zoals: 'begin', 'volgende', 'herhaal'. De spraakherkenning maakt het ook mogelijk om het systeem te personaliseren (bijvoorbeeld het volume regelen). Aangezien het systeem gebruikt gaat worden in omgevingen met veel lawaai (garages) is het robuust gemaakt tegen achtergrondgeluid.

Het systeem zal worden toegepast voor demonstratiedoeleinden en in monteursopleidingen. Monteurs kunnen met dit systeem zelf oefenen met het uitvoeren van onderhoud en waar nodig de hulp van het systeem inroepen om te bepalen of ze de juiste acties in de juiste volgorde uitvoeren. Uiteindelijk verkort dit de

opleidingstijd van monteurs en kan een monteur met minder specialistische kennis meerdere verschillende voertuigen onderhouden. Vele toepassingen mogelijk, ook in huis! Het is mogelijk dit systeem te gebruiken voor onderhoud en bediening van allerlei systemen. Een mooi toekomstscenario is te bedenken vanuit maatschappelijk relevante thema's als vergrijzing en versnelling van technologische innovaties: augmented reality in combinatie met spraaktechnologie maakt het leren bedienen van apparaten thuis (televisie, video recorder, intercom, etc.) een stuk eenvoudiger en plezieriger!

¹ De Defensie Innovatie Competitie 2009 is mogelijk gemaakt door NIDV, BOM en Defensie Research & Development.

² nl.wikipedia.org/wiki/Augmented_reality

Hans Jongebloed
Dutcheer



"BATAVO was conservatief en braaf"

Maar: doelstellingen behaald binnen STEVIN

Het onderzoeksprogramma STEVIN heeft opgeleverd wat ervan verwacht werd. Zo'n tien jaar na de presentatie van de Basis-Taal&spraakVOorziening (BATAVO), de prioriteitenlijst voor de taal- en spraaktechnologie die aan de basis stond van het grote Vlaams-Nederlandse onderzoeksprogramma STEVIN, constateren hoofdauteurs Helmer Strik en Walter Daelemans tevreden dat vrijwel alle voorzieningen op de wensenlijst inmiddels bestaan en beschikbaar zijn.

Prioriteiten

Multilinguale corpora en voorverwerking van tekst prijken volgens Daelemans, professor Computerlinguïstiek aan de Universiteit van Antwerpen, op de prioriteitenlijst van BATAVO. "En natuurlijk een groot corpus geschreven Nederlands". Verder stonden op de lijst nog morfologische, syntactische en aspecten van semantische analyse, een monolinguaal semantisch lexicon, en benchmarks voor evaluaties van die basisvoorzieningen.

Spraaktechnoloog Helmer Strik van de Radbouduniversiteit vult aan: "Voor de spraak- kant was dat een spraakherkenningstool en

gespecialiseerde corpora en benchmarks.

Braaf

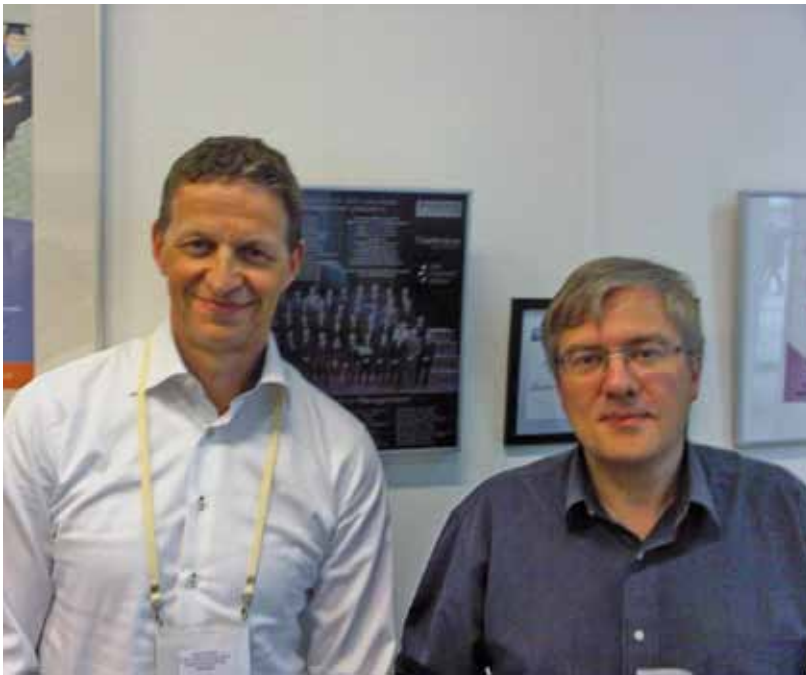
Daelemans: "Wat mij opviel, toen ik de lijst een paar dagen geleden nog eens bekeek, is dat het heel erg conservatief en braaf was. Het was duidelijk gericht op consensus. Zins- semantiek, bijvoorbeeld, zit er niet in. Dat werd gezien als te futuristisch." Strik: "Het idee was dat de projecten op de middellange termijn haalbaar moesten zijn. Dat ook het bedrijfsleven er op middellange termijn iets mee zou moeten kunnen doen." Het rapport stelt inderdaad expliciet dat voorkeur gegeven wordt aan data of modules 'waarvan de ontwikkeling haalbaar is met de huidige stand van zaken in wetenschap en technologie'.

Geld

Het rapport vormde de basis voor het onderzoeksprogramma STEVIN, dat door de Nederlandse en Vlaamse overheid gefinancierd werd. "Nadat het rapport gepubliceerd was hebben verschillende mensen voor deelstukjes een inschatting gemaakt voor de kosten. Het was belangrijk op aan dat wensenlijstje een kostenplaatje te hangen. Maar daarna hebben we in Nederland en Vlaanderen nog wel lang moeten lobbyen en het juiste moment moeten afwachten." Maar het volledige begrote bedrag, ruim 11 miljoen euro, wordt uiteindelijk ter beschikking gesteld. Daelemans vindt dat de lijst 'verbazend' goed is gerealiseerd. Daelemans: "de prijs-kwaliteitverhouding was erg goed." En Strik valt hem bij: "Er is heel veel uit het programma gekomen. Zelfs ook de infrastructuur, in de vorm van de TST-centrale. Daar wordt wel eens op gemopperd in verband met problemen rondom intellectueel eigendom, maar het is goed dat die er is."

Hoewel Strik niet in details wil treden, geeft hij aan dat het zeker niet vanzelfsprekend is dat wanneer het geld beschikbaar wordt gemaakt voor een project, de beloofde resultaten dan ook opgeleverd worden. "Dat is geen onwil", haast hij zich te zeggen, "maar bij grote projecten met veel partners is de samenwerking niet altijd even soepel. Bij STE-

Leonoor
van der Beek



Helmer Strik (links) en Walter Daelemans (rechts)

robuuste spraakherkenning. En synthese." Rond dat laatste onderwerp was nog wel wat discussie, omdat op dat gebied toch al best veel bereikt was. In een draft uit 2001 stond synthese nog niet op de lijst, maar de definitieve versie uit 2002 noemt spraaksynthese wel als prioriteit, naast tools voor betrouwbaarheidsmetingen, spreker- en taalidentificatie en annotatie van spraakcorpora, en

VIN was dat goed geregeld. Misschien waren er wat veel STEVIN-dagen, Taal in Bedrijven en workshops, maar dat was wel goed voor de coherentie”.

Opzet

Strik is enthousiast over de gekozen opzet van STEVIN: “ik geloof dat het een mooie mix geworden is van tenders, open calls en demo-projecten.” Daelemans vindt dat er ook wel wat te zeggen is voor het model dat gehanteerd werd voor het Corpus Gesproken Nederlands, waar van bovenaf bepaald werd wat er moest gebeuren, en er vervolgens onderzoeksgroepen gezocht werden om die doelstellingen te realiseren. “Ik denk dat dan de kans op duplicatie en lacunes kleiner is.” Maar ook met de gekozen opzet moeten de beide opstellers van het rapport lang nadenken over lacunes in het programma. Daelemans oppert uiteindelijk: “Annotatie van corpora met tekststructuur. Dat is er niet gekomen. Er waren wel goede voorstellen, maar er was gewoon te veel concurrentie. Bovendien was de tweede ronde applicatiegericht, waardoor er toen weinig ruimte was voor infrastructurele projecten.” Naslag van het rapport leert dat annotatie van tekststructuur ook niet expliciet als prioriteit benoemd was.

Google

Als er een vervolg op BATAVO zou komen, dan zou Daelemans zich op inferentie willen richtten. De tijd is er rijp voor, denkt hij. “Ik denk dat als je semantiek combineert met Kunstmatige Intelligentie, *inference* en *reasoning*, dat je daar misschien wel in 5 jaar tijd goede resultaten mee zou kunnen behalen”. Maar kan hij daar ook financiering voor vinden? “Nee”, denkt Daelemans. “Dan zullen ze zeggen dat ze dat bij IBM met Watson toch al lang kunnen.” Strik herkent dat argument. “Bij spraak hebben we dat ook. Google heeft enorm veel data. En als ze nieuwe data nodig hebben, dan zetten ze gewoon een gratis dienst op en halen zo de benodigde data binnen. Niemand kan bij die data, maar zij behalen er wel indrukwekkende resultaten mee. En dan krijg je die reactie: maar Google kan het toch al? Waar we vroeger op moesten boksen tegen Lernout en Hauspie, die zeiden dat ze alles al opgelost hadden, moeten we dat nu met Google doen.” Waarop Daelemans riposteert: “Met dat verschil dat Google het ook echt opgelost heeft.”

Toekomst

De kansen voor een tweede BATAVO of STEVIN lijken klein. “Er is geen geld meer voor fundamenteel TST-onderzoek”, denkt Strik.

“Helaas.” Hoewel Strik geniet van toepassingsgerichte onderzoeksprojecten, is hij ervan overtuigd dat fundamenteel onderzoek nodig blijft. “Met toepassingsgerichte projecten help je het veld niet ver voorruit. Bij spraakherkenning werken we al heel lang met dezelfde techniek. Die kruipt maar heel langzaam voorruit, en dan nog alleen door de beschikbaarheid van meer data.” Bovendien was STEVIN specifiek gericht op het Nederlands, en ook daarvoor zien de beide wetenschappers nu minder kansen. “Je kunt het argument dat het Nederlands niet ten onder mag gaan aan de technologische vooruitgang niet eindeloos blijven gebruiken”, aldus Daelemans. Er is ook een voordeel voor wetenschappers nu, ten opzichte van voor BATAVO: “Als je tien jaar geleden interessant onderzoek wilde doen, bijvoorbeeld op het gebied van semantiek, dan moest je beginnen met morfologie en parsing. Nu kun je gewoon meteen met semantiek beginnen. Daardoor is het veel gemakkelijker geworden om interessant onderzoek te doen.”

Taal in Bedrijf, een geslaagd evenement

Op 29 november vond in de Rotterdamse Doelen de vierde editie van Taal in Bedrijf plaats¹. De dag was een feestelijke afsluiting van het Vlaams-Nederlandse STEVIN-programma². Dat wil zeker niet zeggen dat er achterom werd gekeken. De dag was nadrukkelijk zo geprogrammeerd dat er vooral naar de toekomst werd gekeken. Wat heeft het STEVIN-programma aan mogelijkheden voor de toekomst geopend? En hoe zien bedrijven de toekomst die taal- en spraaktechnologie in hun toepassingen verwerken? Dat waren de scharnierpunten van de dag. En het zijn ook de elementen die u in deze DIXIT terugziet.

Henk van den Heuvel
CLST Radboud
Universiteit
Nijmegen

En het was een mooie dag. Ruim 200 aanmeldingen waren er, van een divers publiek: Kennisinstellingen, beleidsmakers, bedrijven en hun klanten. Goed overigens om te zien dat er ook zo'n gemêleerd publiek was, jonger en ouder, student en senior, man en vrouw.

Ondanks wat problemen met het al dan niet openbaar vervoer slaagden de meesten erin op tijd De Doelen te bereiken. Daar wachtte hun een mooi programma, dat werd geopend door Linde van den Bosch, voorzitter van het TST-bestuur van de Nederlandse Taalunie. Linde nam de ontwikkeling in de TST op een aansprekende wijze 'op de hak'.

Arjan van Hessen was een bedreven dagvoorzitter. Hij leidde de dag in met een plaatje dat liet zien hoe de technische mogelijkheden van TST en de verwachtingen in de loop der tijd aan het convergeren zijn. Op ontspannen wijze praatte hij de onderdelen aan elkaar en gaf verhelderende samenvattingen waar nodig.

De dag kende twee keynotes. Aan het einde van de dag sprak Yuri van Geest over singulariteit³, en 's ochtends vertelde Branimir Boguraev van IBM over Watson en de Jeopardy! uitdaging. Watson als open-domein vraag-antwoordsysteem werkt zonder onderliggende databases en ontologieën op basis van evidence retrieval en evidence scoring. Uit de inleidende opmerkingen werd wel duidelijk dat alles wat je er meer van zegt tot gerechtelijke vervolging kan leiden, dus laat uw correspondent het hier maar bij.

In de aansluitende paneldiscussie lieten Peter de Bie, Lou Boves, Kees Groeneveld en Paul Heisterkamp aan de hand van gerichte vragen hun gedachten gaan over Business Intelligence. Hoe kunnen bedrijven gebruik maken van wat in universiteiten wordt ontwikkeld, juist ook in het kader van programma's als STEVIN?

venmarkt die door NOTaS⁴ was georganiseerd en gesponsord. Hier kon men bij 13 bedrijven een blik werpen op de producten waar aan gewerkt wordt. Voor de aanwezige studenten was het een interessante mogelijkheid om in contact te komen met een toekomstige werkgever.

Na een voortreffelijke lunch begint het mid-dagprogramma dat werd gevormd door twee rondes van twee maal vier parallelsessies. De thema's van de sessies waren: Klantendiensten, Informatie-extractie (1 en 2), Communicatie en informatievoorziening, Overheid, Zorg, Onderwijs, Social Media. Deze thema's waren goed gekozen want men zag weinig volk bij de bedrijvenmarkt gedurende de sessies.

De dag werd afgesloten met een borrel. Deze werd aangeboden door NOTaS dat hiermee zijn 10-jarig bestaan op feestelijke wijze onder de aandacht bracht. Het was laat toen wij uit Rotterdam vertrokken.

¹⁾ <http://taalinbedrijf2011.org>

²⁾ <http://taalunieversum.org/taal/technologie/stevin>

³⁾ Zie het verslag van Leonoor van der Beek hierna

⁴⁾ www.notas.nl



Branimir Boguraev

Technologie voorbij de menselijke intelligentie

Een blik op de toekomst door Yuri van Geest

Stel je een wereld voor waarin je met je gedachten een apparaat kan bedienen. Of waarin je bij aankomst op een event je DNA-profiel laat scannen, waarna de rest van het programma, inclusief het eten, specifiek op jouw genen is aangepast. Deze toekomst, zo stelt Yuri van Geest, oprichter van Mobile Monday Amsterdam, Nederlands ambassadeur voor Singularity University en het jongste lid van minister Verhagens Topteam Creative Industry, die is er al. Ter afsluiting van het event Taal in Bedrijf schetst hij een spannend beeld van de technologische toekomst, en hoe taal- en spraaktechnologie daarin passen.

De boude voorspellingen van Van Geest zijn gebaseerd op het principe van singulariteit, dat stelt dat de groei van technologische capaciteit niet lineair is, maar exponentieel. Dat zorgt ervoor dat binnen afzienbare tijd computers slimmer worden dan de slimste mens op aarde. Eng? Wellicht, maar Van Geest legt de nadruk op de onbegrensde mogelijkheden van zulke technologie.

IBM's QA machine Watson en Apple's Virtueel Assistent SIRI zijn nog de meest bescheiden voorbeelden van de grote technologische vooruitgang. Yuri beschrijft toepassingen van DNA- en neuroprofielen voor recruitment, gezondheidszorg en de voedselindustrie. Mobiele DNA-tests voor 's werelds naaste ziektes zijn dankzij biotechnologie op komst, en dromen zullen we volgens Van Geest ooit direct vanuit ons brein kunnen downloaden.

Niet alleen zorgen technologische ontwikkelingen ervoor dat we deze informatie kunnen achterhalen en gebruiken, maatschappelijke en sociale verschuivingen zorgen er bovendien voor, zo stelt Van Geest, dat we meer en meer geneigd zullen zijn die informatie ook te delen via social media – wat weer tot nieuwe toepassingen zal leiden. Daarbij wordt de informatie die we achterhalen en die we delen steeds persoonlijker: van zorgvuldig geredigeerde blogs en foto's, via tweets en locaties naar persoonlijke medische informatie, DNA-profielen en herseninformatie.

TST, zo verwacht Van Geest, zal vooral een rol spelen in de interface tussen mens en machine. Daar zal het bovendien moeten concurreren en samenwerken met andere interfaces, zoals het aloude keyboard en touch screens en nieuwe interfaces zoals de al genoemde Brain Computer interface en augmented reality. Verschillende toepassingen vragen om verschillende interfaces. Opvallend genoeg gaat Van Geest niet in op de bijdrage die TST



Yuri met headset voor Brain Computer Interface.

kan leveren aan data mining, information retrieval en personalisatie.

De veranderingen gaan zó snel, stelt Van Geest, dat er een reëel gevaar is dat de oudere generatie niet mee kan komen. Hij ziet het dan ook als zijn taak om de wereld te informeren over technologie die er aankomt, en die er nu al is – ook al zijn die nieuwe ontwikkelingen nog weinig zichtbaar in het alledaagse leven. De jonge generatie moet de oudere gaan mentoren, vindt hij, zodat iedereen mee kan blijven doen. De TST-gemeenschap is gewaarschuwd.

Leonoor
van der Beek

Verslag Sociale Media Taal in Bedrijf 2011

De sociale media waren prominent aanwezig bij Taal in Bedrijf 2011 (TiB). Een van de thema-sessies was volledig aan de sociale media gewijd. Daarnaast was het mogelijk om via een groep op LinkedIn, het Twitter-account @taalinbedrijf en de hashtag #taalinbedrijf als, naar en over Taal in Bedrijf te communiceren.

Remco van Veenendaal

@rvanveenendaal
@taalinbedrijf

En het was een mooie dag. Ruim 200 aanmeldingen waren er, van een divers publiek: Kennisinstellingen, beleidsmakers, bedrijven en hun klanten. Goed overigens om te zien dat er ook zo'n gemêleerd publiek was, jonger en ouder, student en senior, man en vrouw.

Ondanks wat problemen met het al dan niet openbaar vervoer slaagden de meesten erin op tijd De Doelen te bereiken. Daar wachtte hun een mooi programma, dat werd geopend door Linde van den Bosch, voorzitter van het TST-bestuur van de Nederlandse Taalunie. Linde nam de ontwikkeling in de TST op een aansprekende wijze 'op de hak'.

Arjan van Hessen was een bedreven dagvoorzitter. Hij leidde de dag in met een plaatje dat liet zien hoe de technische mogelijkheden van TST en de verwachtingen in de loop der tijd aan het convergeren zijn. Op ontspannen wijze praatte hij de onderdelen aan elkaar en gaf verhelderende samenvattingen waar nodig.

De dag kende twee keynotes. Aan het einde van de dag sprak Yuri van Geest over singulariteit³, en 's ochtends vertelde Branimir Boguraev van IBM over Watson en de Jeopardy! uitdaging. Watson als open-domein vraag-antwoordsysteem werkt zonder onderliggende databases en ontologieën op basis van evidence retrieval en evidence scoring. Uit de inleidende opmerkingen werd wel duidelijk dat alles wat je er meer van zegt tot gerechtelijke vervolging kan leiden, dus laat uw correspondent het hier maar bij.

In de aansluitende paneldiscussie lieten Peter de Bie, Lou Boves, Kees Groeneveld en Paul Heisterkamp aan de hand van gerichte vragen hun gedachten gaan over Business Intelligence. Hoe kunnen bedrijven gebruik maken van wat in universiteiten wordt ontwikkeld, juist ook in het kader van programma's als STEVIN?

Tijdens de lunch kon men terecht op de bedrijvenmarkt die door NOTaS was georga-

niseerd en gesponsord. Hier kon men bij 13 bedrijven een blik werpen op de producten waar aan gewerkt wordt. Voor de aanwezige studenten was het een interessante mogelijkheid om in contact te komen met een toekomstige werkgever.

Na een voortreffelijke lunch begint het middagprogramma dat werd gevormd door twee rondes van twee maal vier parallelsessies. De thema's van de sessies waren: Klantendiensten, Informatie-extractie (1 en 2), Communicatie en informatievoorziening, Overheid, Zorg, Onderwijs, Social Media. Deze thema's waren goed gekozen want men zag weinig volk bij de bedrijvenmarkt gedurende de sessies.

De dag werd afgesloten met een borrel. Deze werd aangeboden door NOTaS dat hiermee zijn 10-jarig bestaan op feestelijke wijze onder de aandacht bracht. Het was laat toen wij uit Rotterdam vertrokken.

¹⁾ <http://taalinbedrijf2011.org>

²⁾ <http://taalunieversum.org/taal/technologie/stevin>

³⁾ Zie het verslag van Leonoor van der Beek hierna



Arjan van Hessen (links) en Branimir Boguraev (rechts).

Taal in Bedrijf

Het zal ergens in de zomer van 2011 zijn geweest dat Catia Cucchiarini mij belde met de mededeling dat 'ze' tijdens een TiB-vergadering bij de Taalunie (waar ik niet bij kon zijn) besloten hadden mij dagvoorzitter te maken. Eervol zo scheen het mij toe, maar uiteraard was vooral het budget hier leidend in! Naarmate eind-november naderde, werd ik toch wel zenuwachtig: natuurlijk was ik wel vaker dagvoorzitter geweest maar nog nooit voor ongeveer iedereen uit 'mijn wereld'. Als ik het nu niet goed doe, kan ik beter emigreren!

Het zal ergens in de zomer van 2011 zijn geweest dat Catia Cucchiarini mij belde met de mededeling dat 'ze' tijdens een TiB-vergadering bij de Taalunie (waar ik niet bij kon zijn) besloten hadden mij dagvoorzitter te maken. Eervol zo scheen het mij toe, maar uiteraard was vooral het budget hier leidend in! Naarmate eind-november naderde, werd ik toch wel zenuwachtig: natuurlijk was ik wel vaker dagvoorzitter geweest maar nog nooit voor ongeveer iedereen uit 'mijn wereld'. Als ik het nu niet goed doe, kan ik beter emigreren!

Programma Commissie

Het einde van het succesvolle Stevin-programma bestond uit twee onderdelen: een laatste bijeenkomst van de Programma Commissie met het International Advisory Panel: mensen uit het vak van buiten Nederland die de afgelopen zes jaar meer dan 70

voorstellen hebben gelezen en beoordeeld. Hun commentaar was dikwijls zeer vleidend geweest voor het Stevin-initiatief (goed, waarom doen wij dat 'thuis' niet?) en overtuigde de PC door te gaan met de ingeslagen weg.

De bijeenkomst vond plaats in de bijzonder prettige ambiance van Hotel New York in Rotterdam. Verschillende deelnemers (zowel aan de wetenschappelijke- als aan de demonstratieprojecten) kwamen kort vertellen over hun ervaring binnen het Stevin-programma: zowel de goede als de minder goede. Het was goed te zien dat veel bedrijven de moeite hadden genomen om iemand te sturen en nog beter om te horen dat de algemene perceptie van het project zeer positief was!

De avond werd afgesloten met een genoeglijk bordje pasta aan de overkant van het water. Tot mijn grote vreugde was daar ook de keynote spreker van de volgende dag aanwezig: Banimir Boguraev van IBM's TJ Watson Research Center, die een verhaal over WATSON zou houden. Eigenlijk was het de bedoeling die week nog een aantal interviews met allerlei bladen en zelfs de tv te doen, maar Banimir maakte mij (tot zijn eigen spijt) heel erg duidelijk dat dat echt niet kon. IBM heeft zoveel energie (lees: geld) in WATSON gestoken dat hij er alleen met allerlei bobo's erbij met de pers over mag praten. Jammer en een gemiste kans om in Nederland de zichtbaarheid van TST te verhogen, maar vanuit het gezichtspunt van IBM wel begrijpelijk. De rest van de avond besteed aan het nogmaals en nogmaals doornemen van de 'dag-erna'.

Taal in Bedrijf

De dag begon mooi met allerlei buitenlandse collega's aan het ontbijt die allemaal nog tijd genoeg hadden om eens rustig te ontbijten. Snel de spullen gepakt en per metro naar de Doelen: een lelijk doch functioneel gebouw

Arjan van Hessen
Telecats &
Universiteit
Twente



in de buurt van Rotterdam CS. Al bij de ingang veel bekenden gezien hetgeen mijn humeur goed deed maar de zenuwen iets minder.

Boven bij de grote zaal was het al een drukte van belang. Tijd voor gezellig bijbabbelen was er niet: ik werd direct door de dames van organisatiebureau Kuipers gekaapt om het dagprogramma door te nemen. Het bleek dat Kees Groeneweg, een van de leden in de paneldiscussie van die ochtend, ziek was geworden. Wat nu? Gelukkig had ik Leonoor van Beek die ochtend zien binnenkomen: een goede spreker en werkend op de grens van onderzoek en toepassingen. Ze bleek te overtuigen te zijn, zodat de paneldiscussie qua vorm gered leek te zijn.



De zenuwen gierden nu echt door mijn lijf, maar na de eerste woorden op het podium om de dag te openen en de eerste spreekster (Linde van den Bosch, directeur van de Nederlandse Taalunie) aan te kondigen, kwam de rust terug en begon ik er zelf ook plezier in te krijgen.

De toespraak van de zeer charmante Branimir was goed te volgen, maar moeilijk te begrijpen. Een team van 40 wetenschappers en ICT-deskundigen had jarenlang aan Watson gewerkt. Fascinerend om dit zo direct van een van de hoofdarchitecten te mogen horen.

Na de keynote kwam misschien wel het lastigste onderdeel van de ochtend: de paneldiscussie! Groot gevaar bij dit soort discussies is dat iedereen het met elkaar eens is, en de discussie zich beperkt tot het "I agree with the previous speakers". We hadden de vragen vooraf met de verschillende deelnemers doorgenomen en voor de presentatie steeds iemand uit de wetenschap en iemand uit het bedrijfsleven gekozen. Hierdoor werden de verschillende standpunten goed duidelijk en werd het zowaar een goede en inhoudelijk interessante discussie.

Bedrijvenmarkt

De ochtend sloot met de aankondiging van de bedrijvenmarkt: een verzameling stands van bedrijven en universiteiten die allemaal lieten zien waartoe ze in staat zijn op het gebied van toegepaste Taal- en Spraaktechnologie. Fascinerend om te zien hoeveel er in Nederland en Vlaanderen op dit gebied gebeurt en hoe wijdverspreid de toepas-

singen zijn. Natuurlijk niet alleen maar door Stevin, maar toch... zonder Stevin was veel nuttigs niet gebeurd.

Na de bedrijvenmarkt en lunch (geen hap kunnen eten: nog veel te gespannen) voorzitter van de themasessie: 'Intelligente ontsluiting / Informatie-extractie'. De zaal helemaal vol en goede verhalen. Helaas haalde ik twee zaken door elkaar: 15 minuten praatje en 5 minuten vragen (samen 20 minuten) werd door mij geïnterpreteerd als 20 minuten praatje EN 5 minuten vragen. Iedere presentatie liep dus 5 minuten uit zodat de laatste spreker slechts 5 minuten had: schandalig! Uiteindelijk hebben de mensen in de zaal hun koffiepauze opgeofferd, maar fraai was het niet.

Yuri van G.

De dag naderde z'n ontknoping: nog één onderdeel te gaan: Yuri van Geest met als titel van zijn verhaal: 'De impact van de combinatie van de singulariteit en TST op de creatieve sector'. Van Geest is lid van het 'Topteam Creatieve Industrie': het clubje mensen dat minister Verhagen bijstaat bij het bepalen waar de Nederlandse wetenschappelijke wereld zich de komende jaren mee moet bezig houden. Bij de introductie van Yuri vroeg ik een warm applaus voor Yuri van Gelder. Vond het niet erg grappig maar de zaal wel: pas bij terugkeer naar mijn stoel begreep ik deze flater. Later werd mij door een taalkundige uit Groningen duidelijk gemaakt dat de frequentie van het multi-word 'Yuri van Gelder' veel en veel hoger is dan die van 'Yuri van Geest', zodat deze blunder volstrekt volgens de regels was. Mooi, maar wel gênant.

Boeiend verhaal van Yuri kon niet verhullen dat zijn kennis van TST niet heel groot is, zodat er in de afsluitende discussie met de zaal zelfs sprake was van enige wederzijdse ergernissen. Begrijpelijk en het maakt direct duidelijk dat we nog een hoop aan de buitenwereld uit te leggen hebben. Zaken die voor 'ons' volstrekt duidelijk zijn, zijn dat zeker niet voor iedereen buiten onze TST-wereld.

Het zat erop! Ik werd heel hartelijk bedankt door Catia van de Nederlandse Taalunie en kreeg een mooi applaus. De complimenten bij de borrel maakte veel van de stress van de afgelopen dagen goed en moe, leeg maar erg voldaan werd de terugtocht naar Utrecht om 18:00 aangevangen.

Het was geweldig te hebben mogen doen!

Stichting NOTaS

Secretariaat:
Postbus 31070,
6503 CB Nijmegen
T: 024-352 88 88
W: www.notas.nl
E: info@notas.nl

NOTaS behartigt de belangen van bedrijven en kennisinstellingen die actief zijn op het terrein van Taal- en Spraaktechnologie (TST). Dit doet zij o.a. door middel van bijeenkomsten, lobbyactiviteiten en het tijdschrift DIXIT.

Deelnemers Stichting

CLST Radboud Universiteit Nijmegen

Erasmusplein 1, 6525 HT Nijmegen
Postadres:
Postbus 9103,
6500 HD Nijmegen
T: 024-361 16 86
W: www.ru.nl/clst
E: clst@let.ru.nl

Het Centre for Language and Speech Technology (CLST) onderzoekt en adviseert op het gebied van taal- en spraaktechnologie. Belangrijke speerpunten zijn robuuste ASR, audio- en tekstontsluiting, en de inzet van TST ten behoeve van het onderwijs en voor mensen met communicatieve beperkingen. Tevens valt onder CLST het Speech Processing Expertise Centre (SPEX), dat is gespecialiseerd in de productie en validatie van taal- en spraakdatabases.

Data Archiving and Networked Services (DANS)

Anna van Saksenlaan 10,
2593 HT Den Haag
Postadres:
Postbus 93067,
2509 AB Den Haag
T: 070-344 64 84
W: www.dans.knaw.nl
E: info@dans.knaw.nl

DANS bevordert duurzame toegang tot digitale onderzoeksgegevens. Hiertoe stimuleert DANS dat wetenschappelijke onderzoekers gegevens duurzaam archiveren en hergebruiken. Tevens biedt DANS met Narcis.nl toegang tot duizenden wetenschappelijke datasets, e-publicaties en andere onderzoeksinformatie in Nederland. DANS is een instituut van KNAW en NWO en deelnemer aan (inter)nationale projecten en netwerken.

Dedicon

Traverse 175,
5361 TD Grave
Postadres:
Postbus 24,
5360 AA Grave
T: 0486-486 486
W: www.dedicon.nl
E: info@dedicon.nl

Dedicon is in Nederland dé organisatie die lectruur en informatie toegankelijk maakt. Zo kan iedereen lezen wat hij wil in de vorm die hem past. Binnen de diensten van Dedicon speelt de moderne taal- en spraaktechnologie een belangrijke rol.

Dutchear

Olof Palmestraat 16-18,
2616 LR Delft
T: 015-219 11 11
W: www.dutchear.nl
E: info@dutchear.nl

Dutchear past spraaktechnologie toe voor het automatisch

doorzoeken van grote audio/video verzamelingen. Typische voorbeeldtoepassingen die we dagelijks doen zijn het zoeken in rechtzaken, vergaderingen en radio/televisieuitzendingen. Daarnaast levert en optimaliseert Dutchear spraaksynthese (TTS) en sprekerherkenning/-verificatie.

GridLine B.V.

Keizersgracht 520,
1017 EK Amsterdam
T: 020-616 20 50
W: www.gridline.nl
E: info@gridline.nl

GridLine is een succesvol IT-bedrijf dat gespecialiseerd is in taaltechnologie voor het Nederlands. Een aantal van onze producten en diensten: Klinkende Taal - plugin die ambtenaren helpt begrijpelijke brieven te schrijven; FAST; SharePoint; Lucene; NedLex - informatieplatform voor advocaten; de TaalServer - voor alle taalopgaven in uw bedrijf; opinion mining, information retrieval, kennismangement. GridLine, voor professionals die het Nederlands gebruiken.

Kecida, NFI

Laan van Ypenburg 6,
2497 GB Den Haag
T: 070-888 64 00
W: www.forensischinstituut.nl
E: m.israel@nfi.minvenj.nl

Als onderdeel van het Nederlands Forensisch Instituut (NFI) ondersteunt het Kennis- en expertisecentrum voor intelligente data-analyse (Kecida) overheidsinstellingen uit de openbare orde- en veiligheidssector (OOV) bij de informatiegestuurd uitvoering van opsporing-, handhaving en toezichthoudende taken. Diverse state-of-the-art technologieën voor datamining, tekstmining en sociale netwerk-analyse worden ingezet om nieuwe inzichten te geven.

Knowledge Concepts

De Handboog 9,
5283 WR Boxtel
T: 0411-610 802
W: www.knowledge-concepts.com
E: sales@knowledge-concepts.com

Gedreven door passie en ervaring in dit vakgebied ontwikkeld en realiseert Knowledge Concepts vooruitstrevende systemen ten behoeve van informatieverwerking. Onze oplossingen onderscheiden zich in aspecten als: verzamelen, samenvoegen, uitlezen, opslaan, ontsluiten en verspreiden van informatie en kennis. Dit realiseren wij met name door innovatief gebruik van taaltechnologie in combinatie met o.a. zoekmachines, document management systemen en portals.

Linguistic Systems BV

Bijleveldsingel 58,
6524 AE Nijmegen
T: 024 -322 63 02
W: www.euroglot.nl
E: info@euroglot.nl

Linguistic Systems BV, opgericht in 1985, is een van de eerste taaltechnologiebedrijven in Nederland. Het bedrijf ontwikkelde het product Euroglot, een meertalig vertaalmiddel, dat door grote multinationals in binnen- en buitenland dagelijks wordt gebruikt. Daarnaast ontwikkelde het bedrijf een parser-generator voor de constructie van morfologie- en grammatica-modules in zes Europese talen. Momenteel werkt het bedrijf aan de voltooiing van Euroglot Automatic, een systeem om hele teksten te vertalen.

OSTT/Sint Maartenskliniek

Hengstdal 3,
6574 NA Ubbergen
T: 024-365 97 18
W: www.ostt.eu/www.maartenskliniek.nl
E: ostt@maartenskliniek.nl

Het "Ontwikkelcentrum voor Spraak- en Taaltechnologie ten behoeve van spraak- en Taalpathologie en revalidatie algemeen" is erkend als expertisecentrum door het Ministerie van het VWS aan de Sint Maartenskliniek. In het kader van het OSTT werkt de Sint Maartenskliniek samen met de afdeling Taalwetenschap van de Radboud Universiteit en met het UMC St. Radboud in Nijmegen.

RightNow Technologies (Intent Guide)

'Diemercircle' Eekholt 40,
1112 XH Diemen
T: 020-531 38 00
W: www.rightnow.com
E: suzanne.dejong@rightnow.com

Understanding true customer intent through RightNow Intent Guide with Natural Language Technology will enable you to match customer goals with strategic corporate targets. You will make it easier for customers to engage during key moments of the customer journey in commercial and support-oriented transactions to create an optimal return per online customer interaction. By identifying and recognizing the intent of customer engagements these become more relevant, efficient and insightful, supporting your business goals in terms of transactions, cost-reduction and online engagement.

Telecats

Colosseum 42,
7521 PT Enschede
Postadres:
Postbus 92,
7500 AB Enschede
T: 053-488 99 00
W: www.telecats.nl
E: info@telecats.nl

Telecats ontwikkelt producten en diensten op basis van taal- en spraaktechnologie en VoIP. Met oplos-singen op basis van IVR en (open vraag) spraakherkenning worden bij telefoongesprekken bellers geïdentificeerd, gesprekken geïdentificeerd en (deels) geautomatiseerd. Met spraak-analyse kunnen inhoud en toon van live gesprekken en audioarchieven worden ontsloten. De taal- en spraaktechnologieën zijn als web-services beschikbaar voor derden die deze willen integreren in hun eigen propositie.

TST-Centrale, p/a Instituut voor Nederlandse Lexicologie

Matthias de Vrieshof 2-3,
2311 BZ Leiden
Postadres:
Postbus 9515,
2300 RA Leiden
T: 071-527 24 95
W: www.inl.nl/tst-centrale
E: servicedesk@inl.nl

Wanneer u op zoek bent naar (informatie over) digitale taalkundige bronnen dan bent u bij ons aan het juiste adres. Of u voor een kennisinstelling werkt of voor het bedrijfsleven, of u geïnteresseerd bent in taal of in spraak, of u een taalkundige bent of een spraaktechnologie, wij zijn u graag van dienst.

Universiteit van Tilburg - CIW/TICC

Warandelaan 2,
5037 AB Tilburg

Postadres:
Postbus 90153,
5000 LE Tilburg
T: 013-466 91 11
W: www.uvt.nl
E: uvt@uvt.nl

In onderwijs en onderzoek van het Departement Communicatie- en Informatiewetenschappen (CIW) ligt het accent op aspecten van menselijke communicatie, zowel talige als niet-talige, inclusief mens-machine interactie, kennisrepresentatie, text mining, machine learning, kunstmatige intelligentie, computationele semantiek, semantische annotatie, en dialoogmodellen en -systemen. Het departement omvat sinds September 2008 het Tilburg Center for Cognition and Communication (TICC). Het departement verzorgt onder andere de bloeiende bachelor- en masteropleidingen Bedrijfscommunicatie en Digitale Media, de masteropleiding Human Aspects of Information Technology, en samen met de Radboud Universiteit Nijmegen de research masteropleiding Language and Communication.

Universiteit Twente - HMI

Drienerloaan 5,
7522 NB
Postadres:
Gebouw Zilverling,
Postbus 217,
7500 AE Enschede
T: 053-489 37 40
W: hmi.ewi.utwente.nl
E: hmi_secr@ewi.utwente.nl

De HMI-groep van de Universiteit Twente (UT) is een onderzoeksgroep binnen de afdeling Informatica. HMI doet onderzoek op het gebied van mens-machine interactie, taal- en spraaktechnologie en zoektechnologie voor multimedia collecties (tekst, spraak, video). In de buurt van de campus van de Universiteit Twente bevinden zich enkele spin-off bedrijven die op hetzelfde terrein actief zijn.

Universiteit Utrecht - Utrecht Institute of Linguistics OTS

Trans 10,
3512 JK Utrecht
T: 030-253 60 06
E: uil-ots@uu.nl
W: www.hum.uu.nl/uilots

Het UIL OTS is een onderzoeksinstituut binnen de Faculteit Geesteswetenschappen van de Universiteit Utrecht en stelt zich ten doel onderzoek te doen naar menselijke taal. De volgende drie dimensies bepalen het onderzoek: (i) de architectuur van het menselijke taalsysteem, (ii) de cognitieve systemen die ten grondslag liggen aan de verwerving en de verwerking van taal, en (iii) de wijze waarop taal wordt gebruikt in communicatie tussen mensen en tussen mens en machine.

Ondersteuning

CumLingua Taal & Communicatie Full-service taalbureau

Bronkhorstweg 48,
5363 TZ Velp (N.B.)
T: 0486 - 471 554
W: www.cumlingua.com
E: info@cumlingua.com

- Copywriting, redactie- en vertaaldiensten
- Taalles op maat en in-company training
- Proefschrijftvertaling en correctie

Leonard Nijmegen bv

W: www.leonard.nl
E: info@leonard.nl

Maakt communicatie zichtbaar

Deykerhoff/Accountants

Rompertsebaan 70,
5231 GT, 's-Hertogenbosch
Postadres:
Postbus 2325,
5202 CH, 's-Hertogenbosch
T: 073-623 24 50
W: www.deykerhoff.nl
E: denbosch@deykerhoff.nl

Deykerhoff/Accountants is een accountantskantoor dat zich in de regio's 's-Hertogenbosch/Waalwijk en Nijmegen vooral richt op de dienstverlening aan ondernemingen en ondernemers in het MKB en het (medisch) vrije beroep. Onze proactieve instelling is gericht op een grote betrokkenheid en een op maat gesneden persoonlijk contact met de klanten.

Malta & de Keyzer

Toernooiveld 300,
6525 EC Nijmegen
T: 024-351 21 08
W: www.malta-online.nl
E: info@malta-online.nl

Malta & de Keyzer, specialisten in officemanagement en secretariaat.

Sponsoring/samenwerking

Nederlandse Taalunie

Lange Voorhout 19 2514 EB Den Haag
Postadres:
Postbus 10595,
2501 HN Den Haag
T: 070-346 95 48
W: www.taalunie.org
E: info@taalunie.org

De Nederlandse Taalunie is een beleidsorganisatie waarin Nederland, België en Suriname samenwerken op het gebied van de Nederlandse taal, onderwijs en letteren.

STEVIN

Lange Voorhout 19, 2514 EB Den Haag
Postadres:
Postbus 10595,
2501 HN Den Haag
(Ter attentie van Peter Spyns)
T: 070-346 95 48
W: www.stevin-tst.org
E: pspyns@taalunie.org

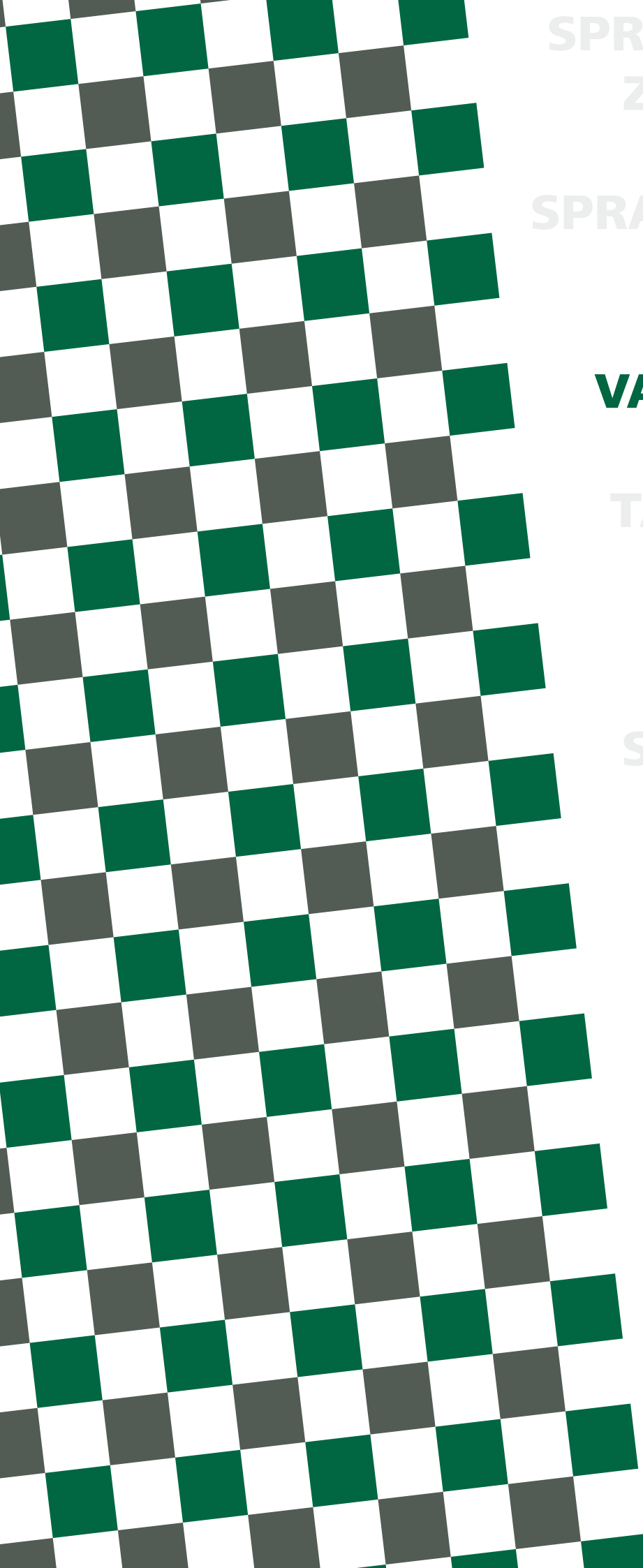
STEVIN is een meerjarig onderzoeks- en stimuleringsprogramma voor Nederlands-talige taal- en spraaktechnologie dat gezamenlijk door de Vlaamse en Nederlandse overheid wordt gefinancierd. STEVIN wordt gecoördineerd en financieel beheerd door de Nederlandse Taalunie.

NWO

Laan van Nieuw Oost-Indië 300,
2593 CE Den Haag
Postadres:
Postbus 93138,
2509 AC Den Haag
T: 070-344 06 40
W: www.nwo.nl
E: nwo@nwo.nl

Contactpersoon NOTaS:
a.dijkstra@nwo.nl

De Nederlandse Organisatie voor Wetenschappelijk Onderzoek heeft tot taak het bevorderen van de kwaliteit en vernieuwing van wetenschappelijk onderzoek, het initiëren en stimuleren van nieuwe ontwikkelingen in het wetenschappelijk onderzoek alsmede de bevordering van de overdracht van kennis van de resultaten van door haar geïnitieerd en gestimuleerd onderzoek ten behoeve van de maatschappij.



SPRAAKHERKENNING
ZEGT U HET MAAR
DIGITALISEREN
SPRAAKTECHNOLOGIE
DIALOOG
REPRESENTEER
VANZELFSPREKEND
VOICE RESPONSE
TAALTECHNOLOGIE
SPRAAKANALYSE
OPEN VRAAG
VERSTAAN
SPRAAKSYNTHESE
AUDIOMINING
BEGRIJPEN
EMOTIEDETECTIE
NASPREKEN
OPLIJNEN
LUISTEREN
HOREN
IVR
HERKEN
VOIP



TELECATS