

DIXIT

TIJDSCHRIFT OVER TOEGEPASTE TAAL- EN SPRAAKTECHNOLOGIE

Conferentiefetisjisme

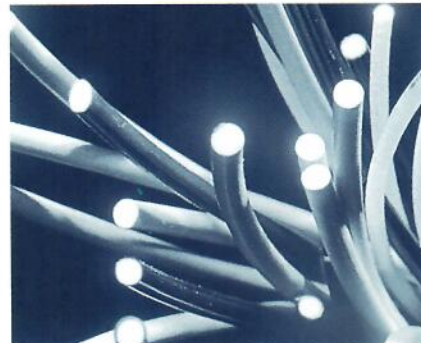
Wat betekent de vraag?
Natuurlijke taaltechnologie weet het!



Voice response:

*besparen door te
investeren*

DE VOICE RESPONSE SPECIALISTEN



Van eenvoudige automatische telefonistes tot de meest complete web-enabled interactieve voice response applicaties voor uiteenlopende doeleinden en diverse organisaties. Alle applicaties zijn voorzien van de modernste technieken zoals spraakherkenning en text-to-speech.

ORCAvoice
computer telephony

ORCAvoice levert stand-alone voice response systemen en interactieve voice response applicaties zoals:

- Automatische telefoniste
- Informatielijnen
- Bestelsystemen
- Reserveringssystemen
- Banksaldolijnen
- Betalingsverificatiesystemen
- Voice mail
- Fax-on-demand

Bel voor meer informatie naar
0412 - 695100 of surf naar
www.orcavoice.nl

ORCApharma
ivr solutions

ORCApharma levert voice response applicaties voor de farmaceutische industrie op basis van een interactief clinical trial systeem. Toepassingen zijn onder andere:

- Randomisatie
- Supply management
- Patient diaries
- Patient alerts
- Patient recruitment
- Emergency code break
- Maatwerkoplossingen

Bel voor meer informatie naar
0412 - 695195 of surf naar
www.orcapharma.com

AudioCom
0800/0900 diensten

AudioCom levert audiotext diensten op basis van 0800/0900 nummers die vanaf een centraal voice response platform worden geëxploiteerd. Deze diensten worden onder andere ingezet voor:

- Bel&Win acties
- SMS&Win acties
- Radio- en tv acties
- Sales promotion
- Productinformatie
- Wervingscampagnes
- Calamiteitenlijnen

Bel voor meer informatie naar
0412 - 695111 of surf naar
www.audiocom.biz

Van standaard oplossingen tot maatwerk applicaties.



INHOUD

Op basis van Natuurlijke Taal-technologie de exacte betekenis van een vraag achterhalen

Ook in de zorgverzekeringsbranche staat het thema 'klantcontacten' hoog op de agenda. De meeste zorgverzekeraars hebben te maken met toenemende (piek)druk in het aantal klantvragen, vooral per telefoon en per e-mail. Dit legt een druk op de service (contact center) organisatie en daarmee op het niveau van de dienstverlening aan de klant. In hoeverre kan nieuwe technologie uitkomst bieden?

4

Leeralgoritmen en Verse Vijgen aan de Adriatische Zee

De 'European Conference on Machine Learning (ECML)' vond plaats in september 2003 in Cavtat, Kroatië aan de Adriatische zee. De ECML is de Europese variant van de 'International Conference on Machine Learning (ICML)', dé conferentie op het gebied van automatische leeralgoritmen en de toepassingen ervan.

8

Babylonische spraakverwarring

Een voice response-systeem dat bellers in 17 talen tegelijkertijd te woord staat, kan al gauw leiden tot een Babylonische spraakverwarring. Toch wilde Organon een systeem dat inderdaad bellers uit 17 landen in hun eigen taal zou toespreken. Met een haast paradoxaal resultaat, want juist door het gebruik van die 17 talen is het systeem 100 procent betrouwbaar geworden in termen van gerealiseerde doelstelling. Kortom: iedereen kan in de eigen moedertaal communiceren met het voice response-systeem.

10

Rubrieken

Uitgelicht	6
Column	7
Wetenswaardig	13
Uitgelicht	15
Uit de boeken	16
TST Dictionary	17
NOTaS-nieuws	18

We drinken een glas...

Nederland wil een kenniseconomie worden. Sterker: moet een kenniseconomie worden. Kennispartij bij uitstek – D'66 – wierp zich bij de kabinetsformatie op als waar voorvechter van deze economische realiteit, resulterend in het Innovatieplatform en een werkgroep onder leiding van JP zelf. En toen? Toen helemaal niets. De digitale pendant van het Innovatieplatform lijkt aan te geven hoe urgent kennis en kennisontwikkeling voor de politiek is. *Hier komt de website van het Innovatieplatform* – luidt de welkomsttekst op de website. En JP is momenteel veel te druk met de waarden en normen, blijkbaar kan het kennishuis niet zonder een degelijk waarden-en-normen-fundament gebouwd worden.

Laten we voor de broodnodige context even terug in de tijd gaan. Vroeger, zeg rond de jaren '90, kochten bedrijven als Philips en PTT Telecom (een deel van) hun R&D bij universiteiten. Dat leverde kennis op met toegevoegde waarde. Maar binnen Philips hebben zich koerswijzigingen voorgedaan en hoe het met PTT Telecom (nu KPN) is verlopen, weten we allemaal. Bedrijven als Philips doen aan R&D met een *product-to-market* horizon van misschien wel twintig jaar. Kleine bedrijven, starters, en MKB – de kurk van de economie – hanteren een kortere horizon. En omdat universiteiten er zo hun eigen planning en speerpunten op na hielden, die volstrekt niet strookte met *wat* bedrijven *wanneer* wilden, bleek deze weg linksom al snel doodlopend. Maar gelukkig bewandelden de universiteiten tegelijkertijd ook het pad rechtsoom: men riep stichtingen in het leven om hun vindingen te vermarkten. Een pad dat om dezelfde reden doodlopend bleek. De discrepantie tussen wat bedrijven nodig hadden en wat de universiteiten boden, bleek wederom onoverbrugbaar.

Tijd dus voor een kentering. Door het TST-platform, als onderdeel van de Nederlandse Taalunie, is gewezen op de ernstige noodzaak om te komen tot Nederlandstalige TST-bouwstenen.

Door de overheid gestimuleerde ontwikkeling van deze basisvoorzieningen kan het Nederlandse TST-bedrijfsleven een internationale push geven.

De Stichting NOTaS stelt voor om een platform op te richten waar bedrijfsleven en universiteiten gaan samenwerken. Een platform dat reguleert en ervoor zorgt dat de geldstromen op de juiste plaats terecht komen. Bovendien zou er een soort makelaardij in het leven moeten worden geroepen die de producten die bedrijfsleven én universiteiten hebben ontwikkeld, licenseert en distribueert. Waarmee geld wordt gegenereerd dat nieuw onderzoek kan financieren. Een perpetuum mobile als het ware, aangejaagd door kennis!

Een goed plan vond iedereen, zo goed dat het inmiddels min of meer is afgeserveerd. Er is weliswaar een commissie opgericht die de samenwerking tussen universiteiten en bedrijven moet bevorderen, maar dan wel een commissie waarin de representant van het overgrote deel van het TST-bedrijfsleven en universiteiten schittert door afwezigheid. Dus de geldstromen lopen nog altijd vanuit EZ en OC&W via het NWO naar de universiteiten, met gerede kans dat ze als vanouds opnieuw hun doel volkomen missen.

'Het moet anders', zegt de politiek. 'Mee eens', zeggen ambtenaren, NWO's en universiteiten in koor. En ze hieven het glas, deden een plas, en alles bleef zoals het was.

Pierre Pieterse
Hoofdreducteur DIXIT

Self Service via website Zorg en Zekerheid loont

Op basis van Natuurlijke exacte betekenis van

Ook in de zorgverzekeringsbranche staat het thema 'klantcontacten' hoog op de agenda. De meeste zorgverzekeraars hebben te maken met toenemende (piek)druk in het aantal klantvragen, vooral per telefoon en per e-mail. Dit legt een druk op de service (contact center) organisatie en daarmee op het niveau van de dienstverlening aan de klant. In hoeverre kan nieuwe technologie uitkomst bieden zonder de kwaliteit van de service aan te tasten? Kunnen de klantvragen bijvoorbeeld ook (deels) worden beantwoord via de website?

| ARJO VAN OOSTEN &
ASTRID SMEDEMA |

Terwijl vrijwel alle organisaties in de verzekeringsbranche nu deze mogelijkheden analyseren, heeft Zorg en Zekerheid de knoop al doorgemaakt en de eerste resultaten behaald. Dankzij de online-marketing en self-service-applicaties, gebaseerd op natuurlijke taaltechnologie, beantwoordt het bedrijf via de website momenteel rond de 1.200 klantvragen per dag.

DE BESTE SERVICE, ALTIJD EN VOOR IEDEREEN?

"Het gaat om de afweging tussen het leveren van 'de beste service, altijd en voor iedereen' en de kosten die je hiervoor als bedrijf moet maken." Aan het woord is Arjo van Oosten. Van Oosten is directeur van Contact Experience en was als 'channel manager' tot het begin van dit jaar verantwoordelijk voor het *Front Office* van Zorg en Zekerheid.

In het derde kwartaal van 2003 stelde Zorg en Zekerheid vast dat het aantal telefoontjes van klanten naar het call center met 25 procent was gestegen. Het aantal e-mails was in dat jaar zelfs met 200 procent toegenomen.

Per 1 januari van 2004 zouden bovendien de wettelijke veranderingen in de dekking van de ziekenfondsen ingaan. Het gevolg: minder dekking en hogere premies. Als reactie op deze veranderingen introduceerde Zorg en Zekerheid op die datum diverse nieuwe modulaire verzekeringspakketten, met nieuwe namen en nieuwe voorwaarden.

Veel veranderingen dus, over onderwerpen die de klanten direct raken. En dan kun je voorspellen dat er nog veel meer vragen van klanten gaan komen... Het werd duidelijk dat de wachttijden voor klanten bij het call center te ver zouden oplopen als er geen actie werd ondernomen.

CROSS-SELLING KANSEN

Een te laag serviceniveau aan de klanten heeft negatieve gevolgen voor de verkoopkansen in de markt. Van Oosten: "De service aan onze klanten is altijd al één van de sterke punten geweest van Zorg en Zekerheid. Dat past helemaal bij de strategie van 'klantintimiteit'. Dat zij deze service op peil wilden houden was niet alleen van belang voor de loyaliteit van de klanten, maar ook voor de kansen in de markt. Probeer maar eens om een aanvullende verzekering aan een klant te verkopen, als hij net vijf minuten 'in de wacht' heeft gestaan om een antwoord te krijgen op zijn vraag."

MAATREGELEN

Op basis van deze omstandigheden stelde Zorg en Zekerheid een pakket aan maatregelen samen, met als voornaamste doelstellingen: de druk op het *front office* (het *contact center*) te verlichten en de dienstverlening aan de klant te waarborgen. In de ideale situatie zouden de contact center-medewerkers zich alleen nog bezighouden met de kansrijke contacten, waarbij de 'standaard' vragen op een andere manier



Taaltechnologie de een vraag achterhalen

worden afgehandeld.

Een belangrijk onderdeel van dit pakket aan maatregelen is de implementatie van een self-service-applicatie om de website van Zorg en Zekerheid gericht in te zetten om klantvragen te beantwoorden. Vanzelfsprekend werd hiervoor eerst de business case opgesteld.

BUSINESS CASE

Van Oosten: "De business case gaf een duidelijk beeld. Na een inhoudelijke analyse van de vragen die via het contact center binnenkomen, werd duidelijk dat 20 procent van de vragen per telefoon 'standaard'vragen zijn. Voor de e-mails gold dit zelfs voor 75 procent. Op basis van deze cijfers berekenden we dat de terugverdientijd voor de investering in de software zo'n drie maanden was. De leverancier van de software, Q-go, gaf ons hiervoor trouwens ook de financiële garantie dat de investering binnen zes maanden terugverdiend zou worden."

Voorwaarde voor Zorg en Zekerheid in dit rekenmodel was wel dat de klantvragen ook echt door het systeem kunnen worden afgehandeld en de klanten niet na een vruchteloos bezoek aan de website alsnog in het contact center terechtkomen. Van Oosten: "Het systeem moet 'begrijpen' wat de klant precies vraagt, juist daarom kozen we voor Q-go."

De oplossingen van Q-go werken op basis van *Natural Language Processing*, ofwel Natuurlijke Taaltechnologie. Het systeem analyseert de vraag op basis van grammaticale regels en achterhaalt zo de exacte betekenis van de vraag van de klant.

Dankzij deze nauwkeurige analysemethode, herleidt het systeem ook de 'vraag achter de vraag'. Zo benadert het systeem optimaal de toepassing van de menselijke taal. En kan een klant op een intuïtieve manier communiceren met het systeem: hij

kan zijn vraag stellen in eigen bewoordingen en wordt nog 'begrepen' ook!

INVOERING VAN HET SYSTEEM: CONTENT EN KANAALKEUZE

De eerste stap bij de invoering van een self-service-systeem is het bepalen van de 'content'; de informatie die via de self-service-software beschikbaar is. Van Oosten: "In dit proces werden we gedwongen om na te denken over de klantvragen die we via self service wilden beantwoorden. Een mooi voorbeeld hiervan is 'opzeggingen'. Wil een klant de verzekering bij Zorg en Zekerheid opzeggen, of juist bij zijn oude verzekeringsmaatschappij? Daarnaast is een opzegging ook een kans om de situatie om te buigen en een klant te behouden. En heb je daarvoor niet een 'live' contact met een *contact center* medewerker nodig? Kortom: je moet heel goed nadenken over welk type vragen je via welk medium wilt beantwoorden."

CONTENT EN VALIDERING

Van Oosten: "In eerste instantie heeft Zorg en Zekerheid er voor gekozen om de vragen en antwoorden in de self-service-oplossing te concentreren rondom de vier belangrijkste thema's: loongrensoverschrijding, dekking, premieverhogingen en de nieuwe producten van Zorg en Zekerheid. Overigens is het zo dat die vragen en antwoorden behoorlijk afwijken van de 'frequently asked questions' die ooit door de marketingafdeling waren opgesteld. Hieraan zie je al hoe belangrijk het is om in dit soort trajecten een hechte samenwerking tot stand te brengen tussen het *contact center* en de marketingafdeling."

Na het bepalen van de content volgde de validering en training. Van Oosten zegt: "Tijdens deze validering hebben we met z'n allen vragen 'afgevuurd' op het sys-





Astrid Smedema is Consultant bij Q-go. Arjo van van Oosten is directeur van Contact Experience en was als channel manager tot het begin van dit jaar verantwoordelijk voor het Front Office van Zorg en Zekerheid

teem. Hiermee hebben we het systeem als het ware 'getraind' in het snappen van Zorg en Zekerheid-taal. Hierbij moet je denken aan de namen van onze producten, maar ook, meer in het algemeen, aan de begrippen die in de verzekeringswereld worden gebruikt. Het proces is wel te vergelijken met het leerproces van een taal door een kind; nieuwe begrippen en verbanden worden 'toegevoegd' aan de woordenschat en de grammatica. Overigens ging dit proces erg snel door de goede samenwerking met de leverancier van het systeem."

DIRECT RESULTAAT

Zonder ruchtbaarheid te geven aan de nieuwe self-service-mogelijkheden, kreeg Zorg en Zekerheid vanaf dag één al circa 250 vragen binnen via de self-service-optie op de website. Van Oosten: "We moesten overigens de mensen wel stimuleren om hun vragen ook echt te stellen in hele zinnen en in eigen bewoordingen. Ze zijn niet gewend aan deze 'luxe', maar blijken snel

te leren. Opvallend was daarnaast, dat de vraagonderwerpen echt per week kunnen verschillen. Dankzij het self-service-systeem kunnen we de inhoud van onze website afstemmen op het onderwerp dat op dat moment 'hot' is. Hiermee ondersteunen we ook weer de marketing- en salesactiviteiten."

In januari 2004 werden er al rond de 1.000 tot 1.200 vragen per dag beantwoord via het self-service-systeem. Het aantal telefoontjes was inmiddels zo'n 20 procent gedaald ten opzichte van dezelfde periode in het jaar daarvoor.

Zorg en Zekerheid verwacht dat deze cijfers nog dit jaar substantieel verder zullen stijgen. Van Oosten: "Mijn conclusie is dat self service een volwassen distributiekanaal is, naast de andere kanalen. Als je het op de juiste manier inzet, is het een prima oplossing voor het vraagstuk van klantenservice versus kostenbeheersing. Cruciaal hierin is het begrijpen en goed beantwoorden van de klantvraag."

Uitgelicht!

Lezingenreeks over taal- en spraaktechnologie

Ook dit jaar organiseert de onderzoeksgroep Taal en Spraak van de Katholieke Universiteit Nijmegen weer een lezingenreeks. In deze lezingenreeks presenteren medewerkers van onderzoeksinstituten en bedrijven hun projecten, producten en ontwikkelingen met betrekking tot taal- en spraaktechnologie. De lezingen zijn niet alleen gericht op onderzoek, er is ook ruimte voor toegepaste taal- en spraaktechnologie. Verder ligt het accent zo veel mogelijk op presentaties over state-of-the-art-projecten, producten en ontwikkelingen.

De lezingen staan open voor iedereen met interesse in taal- en/of spraaktechnologie. Sprekers worden dan ook nadrukkelijk gevraagd rekening te houden met de brede achtergrond van het publiek en bijvoorbeeld voorzichtig om te gaan met erg technische zaken. Na elke lezing is er tijd voor het stellen van vragen. De thema's voor de komende drie bijeenkomsten zijn de volgende:

- 7 april: From technology to application. Sprekers: Christophe Couvreur (ScanSoft) en Berry Eggen (TUE).
- 12 mei: Multimodal information extraction. Sprekers: Maarten de Rijke (UvA) en Walter Daelemans (UvT).
- 9 juni: Relations to other disciplines. Sprekers: Louis Vuurpijl (NICI) en Ludo Verhoeven (KUN).

Voor meer informatie kunt u een email sturen aan Folkert de Vriend - f.devriend@let.kun.nl - of een kijkje nemen op <http://lands.let.kun.nl/>, onder de rubriek 'news'.

Q-go groeit

Q-go heeft over het boekjaar 2003 een omzetstijging van 75 procent gerealiseerd. Marcel Smit, General Manager van Q-go: "Deze groei is exact in lijn met onze ambities. Steeds meer organisaties kiezen voor de kwaliteit van onze oplossingen om on-line klantvragen te begrijpen en beantwoorden. In Nederland beantwoorden we voor onze klanten zo'n 450.000 vragen per maand. We zijn hard op weg om de de-facto standaard in self-service applicaties te worden voor de top van de financiële sector in Europa. In 2004 gaan we onze positie in de markt uitbouwen, zowel naar andere branches als naar nieuwe geografische markten."

Essent Storingslijn

TeleCats heeft in samenwerking met Essent Kabelcom een Interactieve Voice Response-systeem ontwikkeld voor de Storingslijn van Essent Netwerk Noord. Dit IVR systeem zorgt voor de geautomatiseerde afhandeling van telefoonverkeer van klanten van Essent die bellen om een storing door te geven.

Klanten die een storing willen melden, bellen het daarvoor bestemde storingsnummer en kunnen dan zelf met behulp van druktoetsen aangeven waarvoor ze bellen. Het IVR-systeem zal, wanneer er voor het geselecteerde product of dienst een grote storing is, de beller hiervan op de hoogte stellen. Mocht deze melding niet van toepassing zijn, dan wordt de beller direct doorgeschakeld naar een medewerker van de storingsdienst die verantwoordelijk is voor het geselecteerde product.

Conferentiefetisjisme

Wie wil, kan tegenwoordig gemakkelijk het hele jaar van huis zijn. Beroepshalve wel te verstaan! De onbedaarlijke wild-groei aan conferenties en workshops kan je gemakkelijk voltijds bezig houden. Vroeger had je in mijn vakgebied eens in de vier jaar een grote internationale conferentie: de 'International Conference of Phonetic Sciences', en eens in de twee jaar 'Eurospeech'. Die zijn er nog steeds. Maar als topje van een immense ijsberg. Inmiddels kan ik maart 2004 achtereenvolgens doorbrengen bij de 'Spring Meeting van de Acoustical Society of Japan' in Atsugi, de 'Joint Congress of the French and German Acoustical Societies' in Straatsburg, de 'ISCA Workshop' over *Speech Prosody* in Nara (Japan) en de workshop 'Biometrics on the Internet' in Vigo (Spanje). In april haast ik mij dan naar het 'International Congress on Acoustics' in Kyoto, de 'JEP-workshop' in Fez (Marokko), en de 'SigDial workshop on Discourse and Dialogue' in Boston, waarop in mei volgt... enzovoort. Waarbij aangetekend dat ik absoluut geen enkele aanspraak doe op volledigheid. Wat is er in zeg vijftien jaar gebeurd? "De mogelijkheden voor individuele onderzoekers tot onderlinge informatie-uitwisseling op een groot aantal terreinen van hun wetenschappelijke expertise zijn sterk verbreed" – zou een ronkende formulering zijn om het fenomeen onder woorden te brengen. Feit is echter dat diezelfde individuele onderzoeker door de bomen het bos niet meer ziet. Laten we eens wat praktische consequenties van het conferentiefetisjisme op een rijtje zetten. Om te beginnen kun je dus het hele jaar onderweg zijn. De paar dagen die je nog op je werkplek doorbrengt, kun je volledig besteden aan de verslaglegging van je vorige reis en de voorbereiding van de volgende. Nu bestaat er, althans aan mijn universiteit, een behoorlijke rem op extravagante reislust. Allereerst krijg je niet alle kosten vergoed en bovendien word je geacht op een conferentie ook wat te presenteren. Je moet dus van tevoren een artikel geschreven hebben dat door de organisatie van het spektakel geaccepteerd is. Maar ook dit heeft opmerkelijke gevolgen. Allereerst is het schrijven van conferentiebijdragen natuurlijk alleen maar een extra

aan reizen gerelateerde tijdsbesteding. Net zoals het aanvragen van beurzen en het bestellen van tickets dat zijn. Het schrijven van papers houd je wel thuis maar alleen met het oog op de volgende reis. Maar er is meer. Wat gebeurt er met die papers? Je zult de onderzoekers de kost geven die hetzelfde paper naar verschillende conferenties (tegelijk) opsturen in de hoop en redelijke verwachting dat het allicht ergens zal worden geaccepteerd. Duidelijk is, dat een dergelijke 'artikelensпам' de kwaliteit van het geleverde werk (laat ik mij vooral voorzichtig uitdrukken) niet automatisch doet stijgen. De overstelpende hoeveelheid artikelen die voor al die bijeenkomsten geproduceerd wordt, verlamt het eigenlijke werk van de onderzoeker nog op een andere wijze. Hij (v/m) wordt namelijk ook steeds vaker als *reviewer* opgetrommeld. Een beetje wetenschappelijk onderonsje wil er uiteraard wel mee pronken dat elk ingediend werk door twee, liefst drie, *reviewers* is nageplozen. Die moeten hier dus telkens tijd voor zien vrij te maken. Terwijl ze net zo druk waren met hun eigen papers! Wat doen we hieraan? Zelf heb ik besloten om elk jaar maar één conferentie en misschien nog één workshop te bezoeken. Een voorbeeld dat navolging verdient? Zolang het aantal bijeenkomsten vrolijk floreert, maakt mijn weliswaar milieuvriendelijke uitgangspunt de kans er niet groter op dat ik 'relevante' collega's tegenkom als die toevallig een andere gelegenheid hebben verkozen om hun laatste (?) ontdekkingen wereldkundig te maken. Overigens zou over niet al te lange tijd het aantal bijeenkomsten stevig kunnen slinken door toedoen van zieke geesten, zieke lichamen of zieke kippen. Het daagt in het oosten... Hoe dan ook, het moge duidelijk zijn dat het aantal bijeenkomsten inmiddels allicht eerder een belemmering dan een bevordering van het wetenschappelijk onderzoek tot gevolg heeft. Te hopen valt dat de onderzoeker zijn eigenlijke werk trouw blijft, namelijk grondige analyses maken van de gegevens waar hij mee bezig is, met alle eenzaamheid die daar soms bij komt kijken. Ook hier zijn de woorden van Pascal van toepassing: 'Wie niet alleen op zijn kamer kan zitten, moet zich ook maar niet onder de mensen begeven.'

'Wie niet alleen op zijn kamer kan zitten, moet zich ook maar niet onder de mensen begeven.'

Henk van den Heuvel is verbonden aan het CLST van de Katholieke Universiteit Nijmegen

De complexe waarheid achter machine learning

Toepassingen van leeralgoritmen

Tijdens ECML'03 lijkt de verscheidenheid aan soorten algoritmen die mij verraste tijdens mijn eerste deelname aan ECML'97, verder toegenomen. (Of ben ik mij in de loop der jaren bewuster geworden van wat er gaande is op zo'n conferentie?). Voor een buitenstaander lijkt machine learning soms, op het eerste gezicht, op een soort hocus-pocus-onderzoek. Data gaan erin en plots ontstaat er bijvoorbeeld een accurate *part-of-speech tagger*. Inderdaad, soms lijkt het hierop als je bepaalde, tamelijk oppervlakkige papers moet geloven. Echter, de waarheid achter machine learning blijkt veel complexer te zijn. De verscheidenheid aan algoritmen lijkt te benadrukken dat het niet om een vorm van tovenarij gaat: ieder algoritme vertoont bepaalde preferenties en past beter bij bepaalde soorten taken dan andere algoritmen en vice versa. Dat is tevens één van de hoofdconclusies van een paper dat over taalkundige taken gaat. De kunst is dus een algoritme te ontwikkelen dat bij de taak het beste past, wat overigens nog steeds veel kennis en inspanningen vereist.

| KHALIL SIMA'AN |

Leeralgoritmen en Verse Vijgen aan de Adriatische Zee

De 'European Conference on Machine Learning (ECML)' vond plaats in september 2003 in Cavtat, Kroatië aan de Adriatische zee. Cavtat (Tsavtat) is een klein, mooi dorpje in de buurt van de stad Dubrovnik. De ECML is de Europese variant van de 'International Conference on Machine Learning (ICML)', dé conferentie op het gebied van automatische leeralgoritmen en de toepassingen ervan. Khalil Sima'an doet verslag.

Dit is de tweede keer dat ik de ECML bezoek. De eerste keer was in 1997, in Praag. Tot mijn genoegen moest ik constateren dat er in 2003 meer aandacht is voor taalverwerking en aanverwante technologieën dan in 1997. Het lijkt of taalverwerking de laatste jaren de leeralgoritmen heeft ontdekt. Echter, de ontdekking lijkt echt wederzijds: taalkundige taken en toepassingen kenmerken zich door beschikbaarheid van grote hoeveelheden data (denk bijvoorbeeld aan de bestaande corpora, geannoteerd of niet, aan teksten op internet), en dat is een belangrijke voorwaarde voor het onderzoeken van leeralgoritmen.

Probabilistische modellen, support vector machines, en memory-based learning. De toename van het aantal publicaties over de toepassingen van leeralgoritmen op taalkundige taken mag eigenlijk geen verrassing meer zijn. De *proceedings* van conferenties als de 'Annual Meeting of the Association for Computational Linguistics' (ACL), 'International Conference on Computational Linguistics' (COLING) en de 'Computational Natural Language Learning' (CoNLL) zijn de laatste jaren doorspekt geraakt met papers over *probabilistische modellen*, *support vector machines*, *memory-based learning* en andere classificatie-algoritmen. Net als bij spraakherkenning, waar data-gebaseerde methoden de overhand hebben, lijkt de computerlinguïstiek de laatste decennia dit punt te hebben bereikt. Ik ervaar dit als een gigantische verrijking, omdat het onderzoek hiermee een stadium bereikt waarmee je in één klap zowel de wiskundige als theoretische kant kunt bestuderen als ook de toegepaste en empirische kant.

Natuurlijke taalverwerking

Sommige keynote speakers tijdens ECML'03 blijken erg boeiend te zijn. Leo Breiman (University of California at Berkely), bekend van *Classification and Regression Trees* (CART, ook bekend als *decision trees*), geeft een boeiend voordracht over de nieuwe *random forests*, een robuuste uitbreiding van CART. Zoals altijd vereisen zulke voordrachten bekendheid met het specifieke onderwerp. Andere keynote speakers geven meer algemene voordrachten die echter even inspirerend zijn. Wat echter opvalt aan de ECML'03, is dat natuurlijke taalverwerking nog steeds veel te bieden heeft voor *machine learning*-onderzoek. Ik vermoed dat de taalkundige taken zoals 'parseren' nog steeds te complex zijn voor de gangbare *machine learning*-algoritmen. Het is dus niet opmerkelijk dat op dit moment de enige succesvolle methodes voor 'parseren' gebaseerd zijn op formele

op taalkundig terrein

probabilistische grammatica. Op ECML'03 is er weinig aandacht voor dit soort complexere taken waarin het aantal 'classes' oneindig is. Het enige onderdeel waar enige aandacht wordt geschonken aan het leren van formele grammatica is de ICGI-workshop (ICGI staat voor: 'International Colloquium on Grammar Induction'). Maar deze workshop blijkt meer aandacht te hebben voor de abstracte zogenaamde *learning in the limit* dan voor echte leeralgoritmen op echte data. Desondanks blijft het onderwerp van ICGI erg interessant vanuit mathematisch oogpunt: inductie van grammatica uit platte sequenties van symbolen (woorden).

Zon, zee, en informele ideeënuitswisseling

Vermeldenswaard aan ECML'03 is de locatie: Cavtat/Dubrovnik. De conferentiezalen bevonden zich in een hotel boven op de top van een groen heuveltje boven Cavtat. Iedere ochtend, op weg naar de eerste voordrachten, ging ik even langs de markt om een handjevol verse vijgen te kopen. En 's middags, wanneer ik uit het conferentiecentrum naar het dorp afdaalde om te lunchen, ontdekte ik verscheidene collega's die een duik namen in het heldere water van Cavtat, of lagen te zonnebaden in de mooie septemberzon. Volop gelegenheid dus om collega's te ontmoeten in een ontspannen sfeer. En dat is natuurlijk ook een belangrijke reden om zulke conferenties bij te wonen.

(The 14th European Conference on Machine Learning
(Cavtat, September 2003 -
<http://www.cs.kuleuven.ac.be/conference/ecmlpkdd>)

Khalil Sima'an is universitair docent natuurlijke taalverwerking en statistische leeralgoritmen en is verbonden aan de Universiteit van Amsterdam

Is dit uw oplossing om het groeiende aantal binnenkomende klantcontacten het hoofd te bieden?

Niet nodig!

De nieuwe generatie intelligente* Web Selfservice op basis van ASP

Kostenbesparend, servicegericht en gemakkelijk



SOLUTIONS www.asknow.nl

AskNow! Solutions b.v. Saal van Zwanenbergweg 1 5026 RM Tilburg T • 013-4653430 F • 013-4653439 E • info@asknow.nl

* Gebaseerd op de laatste ontwikkelingen op het gebied van de Nederlandse taaltechnologie, kunstmatige intelligentie en document retrieval

Voice response: besparen door te investeren

Babylonische spraakverwarring

Een *voice response*-systeem dat bellers in 17 talen tegelijkertijd te woord staat, kan al gauw leiden tot een Babylonische spraakverwarring. Een Koreaan die in het Mandarijn te woord wordt gestaan, of een Fin die op basis van vragen in het Grieks de juiste toets maar moet weten te vinden. Toch wilde farmaceut Organon een systeem dat inderdaad bellers uit 17 landen in hun eigen taal zou toespreken. Met een haast paradoxaal resultaat, want juist door het gebruik van die 17 talen is het systeem 100 procent betrouwbaar geworden in termen van gerealiseerde doelstelling. Kortom: iedereen kan in de eigen moedertaal communiceren met het *voice response*-systeem.

| ERIC REUSER |



De ontwikkeling van een nieuw medicijn is voor een farmaceut een enorme investering. Een investering die pas kan worden terugverdiend als blijkt dat het nieuw ontwikkelde medicijn inderdaad werkt en geen bijwerkingen heeft. Het moge duidelijk zijn dat het ontwikkeltraject uiterst zorgvuldig doorlopen moet worden, omdat de gezondheidsrisico's tot het absolute minimum beperkt moeten worden. Het gemiddelde ontwikkel-goedkeuringstraject duurt 8 tot 12 jaar, waarin in de laatste fase het medicijn aan patiënten wordt gegeven om praktijkresultaten te krijgen.

In het meest vroege stadium van de ontwikkeling van het medicijn worden de benodig-

de patenten al vastgelegd, om te voorkomen dat andere farmaceuten hetzelfde medicijn op de markt brengen zodra blijkt dat de werkzame stoffen inderdaad effectief zijn. Een nieuw medicijn kan maximaal 15 jaar gepatenteerd worden. Dit betekent dus – gezien het feit dat de ontwikkeltijd 12 jaar kan bedragen – dat er maar heel weinig tijd is om de investering terug te verdienen. En dit betekent weer dat hoe korter de laatste fase duurt, hoe langer de terugverdientijd is. Door de laatste fase voor een belangrijk deel te automatiseren met een *voice response*-systeem wordt de doorlooptijd aanzienlijk verkort, waardoor het medicijn eerder op de markt gebracht kan worden.

CLINICAL TRIAL STUDIE

Een *clinical trial*-studie is dus de laatste fase van een lange-termijnonderzoek voordat medicijnen op de markt komen. Een aantal artsen vanuit de gehele wereld wordt bij dit testen betrokken. Deze artsen zijn in zoverre op de hoogte van de doelstellingen van de onderhavige studie dat zij geschikte patiënten om aan die studie mee te doen, kunnen selecteren. Niet alleen het aanmelden van de patiënten gaat via een *voice response*-systeem maar ook de medicijnverstrekking. De arts/onderzoeker krijgt dus van het *voice response*-systeem de opdracht medicijn A, B, of C te verstrekken, waarbij de arts/onderzoeker dus absoluut niet weet of hij een pla-

Voice response in 17 talen

Naast het opvangen van de arts/onderzoekers in de eigen taal, heeft het *voice response*-systeem nog veel meer functies in zich: het functioneert als een compleet managementsysteem! Met als voornaamste andere functies:

- Voorraadbeheer medicijnen
- Bewaken houdbaarheidsdata medicijnen
- Optimale medicijn productie
- Vastleggen van de door de FDA verlangde gegevens
- Registreren van patiëntgegevens
- Noodprocedures voor complicaties bij patiënten
- Registreren van randomisatie gegevens
- Eventueel doorverbinden naar de helpdesk

cebo of een medicijn met werkzame stof X of een medicijn met werkzame stof Y verstrekt.

ORCAvoice voert in opdracht van Organon het beheer over een wereldwijd opererend *voice response*-systeem waarop dit soort studies worden uitgevoerd.

TAALBARRIÈRES

Een van de problemen waar je dan tegenaan loopt, is uiteraard het taalprobleem.

Natuurlijk kun je gebruik maken van het Engels als standaardtaal, omdat alle artsen/onderzoekers deze taal machtig zijn. Toch is gekozen om vrijwel alle artsen/onderzoekers in hun eigen taal te woord te staan, omdat dit de kans op miscommunicatie minimaliseert.

Het *voice response*-systeem spreekt eerst in het Engels een welkomstekst uit en vraagt de persoonlijke identificatiecode in te toetsen. Aan de hand van de ingetoetste gegevens schakelt de computer over in de taal van die arts/onderzoeker en vanaf dat moment worden alle *voice files* in de moedertaal afgespeeld.

Er zijn op dit moment studies operationeel die behalve in de voor de hand liggende Europese talen bijvoorbeeld ook kunnen antwoorden in het Chinees, het Thais, het Koreaans en het Japans.

CHINEES KNIPPEN

Voor de programmering van het *voice response*-systeem zijn al die talen natuurlijk geen probleem, in de *flow* wordt vastgelegd uit welke directory de verschillende *voice files* moeten komen, waarbij iedere taal zijn eigen directory heeft. Voor de operators van de geluidsstudio daarentegen was het probleem veel lastiger.

Nadat de teksten zijn ingesproken, moet met behulp van een *edit*-programma iedere *voice file* in het juiste formaat geknipt worden. In een voor de operator verstaanbare taal kan, wanneer tenminste de tekst op papier aanwezig is, beoordeeld worden of de tekst juist is ingesproken en waar de tekst afgebroken moet worden. In het Chinees is dit een stuk problematischer: de klanken zijn dermate onbekend dat zelfs met de tekst op papier (overigens in Chinese tekens) het onduidelijk is waar de tekst ophoudt en waar de volgende tekst begint.

Dit probleem is opgelost door de insprekers iedere tekst te laten beginnen met het inspreken van de *voice file*-nummers in het

engels. De operator hoefde slechts dit nummer van de tekst af te knippen en had daarmee de juiste tekst geïdentificeerd.

Om het beeld compleet te hebben, de gebruikte talen in het draaiende *voice response*-systeem zijn: Engels, Frans, Duits, Spaans, Deens, Fins, Zweeds, Noors, Grieks, Italiaans, Hongaars, Pools, Russisch, Slowaaks, Turks, Chinees (Mandarijns), Thais, Koreaans en Japans.

NAUWKEURIGER EN TOCH BESPAREN

Door de introductie van *voice response* in dit soort medicijnstudies, wordt die studie veel betrouwbaarder. De arts/onderzoeker kan eigenlijk geen fouten maken, het *voice*-systeem bewaakt of de ingevoerde gegevens correct kunnen zijn en vraagt eventueel direct om een correctie of extra bevestiging. Wellicht nog belangrijker: in het verleden werden de patiëntgegevens gefaxt en vervol-

Artsen in hun eigen taal te woord staan minimaliseert de kans op miscommunicatie

gens op het hoofdkantoor ingevoerd, waar uiteraard tikfouten gemaakt konden worden. Door deze beide feiten is de patiëntuitval, dus die patiënten die niet meer voor een studie gebruikt konden worden omdat de data niet klopte, een stuk lager geworden. Omdat de medicijnen die verstrekt worden nog in het research-stadium verkeren, zijn ze extreem duur, dus iedere patiënt die uitvalt, kost handen vol met geld.

Een gemiddelde studie kan circa 2000 patiënten bevatten. De kosten per patiënt bedragen ongeveer € 5000,-. Bij een normale studie rekent men met een uitval van circa 10 procent dat betekent dus een kostenpost van € 1.000.000,-. Door het gebruik van *voice response* wordt de uitval gereduceerd tot 4 procent. Een besparing dus van € 600.000,-.

Eric Reuser is algemeen directeur van ORCAvoice

Grensoverschrijdend

De hier beschreven case is specifiek voor de farmaceutische industrie, maar het is duidelijk dat voor alle internationaal opererende bedrijven eenzelfde systematiek gevolgd kan worden. Den bijvoorbeeld aan een *order and try*-lijn voor een pan-Europees dealernetwerk of een wereldwijde consumentenlijn.

'Di•gi•taal'

**Taalsoftware op
elke werkplek**

**ook voor
intranet**



van Dale

IN BEDRIJF

www.vandale.nl/kantoren

Wat zingt de operazangeres?

Wie wel eens naar de opera gaat, zal zich bij het luisteren naar een door een sopraan gezongen aria dikwijls hebben afgevraagd: wat zingt ze nu? Natuurlijk helpt het als je Italiaans of Duits verstaat maar dan nog is het bijna onmogelijk om te verstaan wat er gezongen wordt. Hoe komt dat eigenlijk? Oftewel hoe op basis van taal- en spraaktechnologie onverklaarbare zaken opeens heel gewoon blijken.

Spraak bestaat uit klinkers en medeklinkers, maar een gezongen aria bestaat voornamelijk uit klinkers. Bij klinkers wordt de lucht door het strottenhoofd geperst waarbij de stembanden de lucht in trilling brengen. De trillende lucht verlaat de zanger via de geopende mond, en bereikt het oor van de luisteraar.

Herz-gehalte

Welke trillingen maken we in spreken en zingen? Net als bij een muziekinstrument gaat het in eerste instantie om de grondtoon: de laagste frequentie van een snaar of het aantal malen dat de stembanden per seconde open en dicht gaan. Behalve deze grondtoon worden ook tonen gemaakt met een veelvoud van de frequentie van de grondtoon. Een pianosnaar die 100 keer per seconde trilt produceert een toon van 100Hz maar ook tonen van 200Hz, 300Hz, 400Hz enz. Deze veelvouden van de grondtoon (boventonen) zijn meestal niet allemaal even sterk. Hun belang wordt bepaald door de resonantie-eigenschappen van de klankkast van het instrument: dit is een van de redenen waarom de ene vleugel anders klinkt dan de andere. Hetzelfde geldt bij de mens. Wij hebben als klankkast onze mond- en keelholte. Twee mannen die allebei een /a/ uitspreken klinken toch anders omdat ze een verschillende mond en keel hebben. Door de stand van de tong en de mond te veranderen, worden de onderlinge verhoudingen

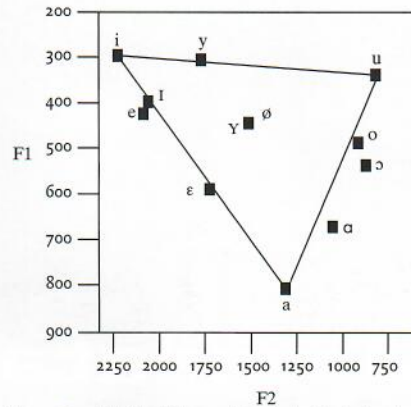


Fig. 1. Gemiddelde klinkerposities in het Nederlands. De illustratie is gebaseerd op de gemiddelde formantfrequenties van 50 Nederlandse mannen uit de studie van Pols, Tromp & Plomp (1973)

van de verschillende boventonen gewijzigd en klinkt er een andere klinker. Door articulatie maken mensen de verschillende klinkers en medeklinkers en kunnen ze praten. De ene klinker verschilt van de andere door de verschillende onderlinge verhoudingen van hun boventonen, die hun oorzaak vinden in verschillen in resonanties.

Resonerende klinkers

Spraak en zang worden gekenmerkt door vijf van deze resonanties, formanten genoemd, waarvan vooral de eerste twee de klinker bepalen. In figuur 1 zien we de frequenties van de eerste twee formanten voor de klinkers in het Noord-Nederlands. De eerste formant onderscheidt onder andere de klinkers /a/ (als in laat) van /i/ (lied) en /u/ (hoed), de tweede formant is belangrijk voor het verschil tussen /i/ en /u/.

Hier komen we op een belangrijk punt voor het waarnemen van spraak en zang: de perceptie van de luisteraar wordt bepaald door de mate waarin het oor (en de hersenen) in staat zijn de positie van deze formanten te achterhalen. Hoe beter dit gaat, hoe beter de luisteraar de klinker 'hoort'.

Dit proces lijkt een beetje op het uitzetten van een weg: hoe meer punten van het tracé vastliggen, hoe duidelijker de weg bepaald is.

De Klinker van Maria Callas

Hetzelfde gebeurt bij het luisteren. Hoe lager de grondtoon, hoe meer boventonen in het gebied dat we voor spraak gebruiken ($\approx 60\text{Hz} \leftrightarrow 4000\text{Hz}$). Immers een grondtoon van 100Hz (spreekstem van de man) heeft bij elk veelvoud van 100Hz een boventoon en dat zijn er in dit geval dus 39. Een vrouw (grondtoon van de spreekstem 200Hz) heeft er tot 4000 Hz 'slechts' 19, maar dat is toch voldoende voor de luisteraar om met succes de formanten te kunnen schatten. Een operazangeres die zingt op een grondtoon van 800 Hz, heeft daarvoor maar 4 boventonen (1600, 2400, 3200 en 4000 Hz). Om als luisteraar dan de positie van de formanten (die aangeven welke klinker gezongen is) te bepalen, is schier onmogelijk. Dat is de reden dat een operazangeres eigenlijk niet te verstaan is wanneer ze erg hoog zingt.

We zien dit duidelijk in figuur 2. Het linker plaatje toont de boventonen (het spectrum) van een door een man gesproken /a/ (grondtoon ≈ 120 Hz), het rechter plaatje is van een klinker van de zangeres Maria Callas die op een grondtoon van ≈ 1000 Hz zingt. We zien de grondtoon (linker piek) en de boventonen (de overige drie duidelijke pieken) maar kunnen daaruit geen schatting maken van de positie van de formanten. We verstaan haar dus niet.

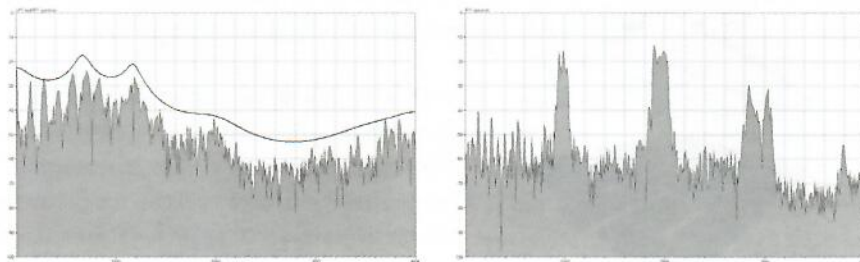


Fig. 2. Het spectrum van een /a/ (als in laat) van een mannelijke spreker met een grondtoon van 120 Hz (links), en een klinker door een zangeres gezongen met een grondtoon van 1000Hz. Goed te zien is dat links een goede schatting van de pieken (formanten) in de omhullende (de dikke lijn) te maken valt, terwijl dit rechts onmogelijk is. De pieken rechts zijn de grondtoon (Linker piek op 1000Hz) en twee boventonen (2000Hz en 3000Hz). Hier een omhullende doorheentrekken is niet goed mogelijk.

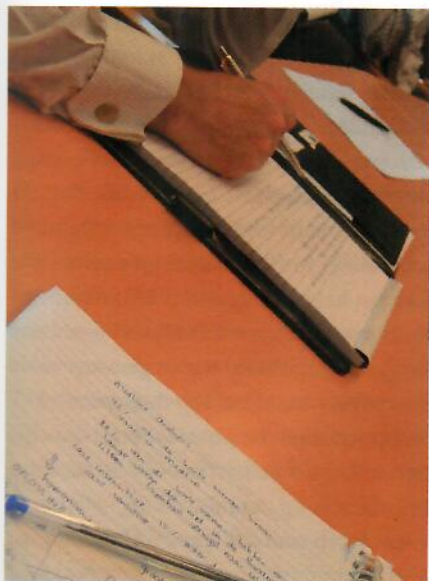
Arjan van Hessen is verbonden aan de Universiteit van Twente. Gerrit Bloothoof is aan de Universiteit van Utrecht, en is gepromoveerd op het onderwerp zangers en formanten en geldt als internationaal expert op dit gebied

Wij werken aan de derde generatie Intranet



Willem van Beuningenlaan 3, 5261 NG Vught.
Tel: +31(0)73 6840199 | Fax: +31(0)73 6840328
E-mail: info@knowledge-concepts.com | Web: www.knowledge-concepts.com

Uitgelicht!



Medisch taalgebruik niet eenduidig

Op 4 februari heeft in Geldermalsen het symposium Medische Taaltechnologie plaatsgevonden. Bijdragen werden geleverd door Collexis (www.collexis.com), de Biosemantics groep van het departement Medische Informatica van de Erasmus Universiteit (MIEUR, www.biosemantics.org) en Polderland Language and Speech Technology (www.polderland.nl). De centrale vraag tijdens het symposium was: 'Op welke manier kan taaltechnologie een rol spelen in het structureren en analyseren van medische documenten?'

Onderzoek van MIEUR ontzenuwt de mythe dat medisch taalgebruik eenduidig zou zijn. Zelfs bij gennamen komt ambiguïteit en homonymie (een woord heeft twee betekenissen) veelvuldig voor. Daarnaast levert inconsistente spelling veel ruis op, waardoor spellingnormalisatie essentieel is. Tevens zal bij het definiëren van de relatie tussen twee termen in de toekomst Natural Language Processing (NLP) een rol kunnen gaan spelen. Om dit goed te kunnen doen, zijn taaltechnologische tools met een hoge accuratesse nodig. Taaltechnologie is dus onontbeerlijk om een voldoende niveau van analyse te bereiken voor het classificeren en indexeren van documenten. Omgekeerd bevat medisch taalgebruik voor de taaltechnologie interessante uitdagingen.

Human Inference in andere handen

Onder de kop 'Prime Technology Ventures investeert in Europees leiderschap Human Inference', wordt wereldkundig gemaakt dat de lang verwachte management buy-out zijn beslag heeft gekregen. Desgevraagd laat directeur Hugo Verwijs weten blij te zijn met deze constructie, omdat met name door de financiële injectie van 3,5 miljoen euro van Prime versneld invulling gegeven kan worden aan de strategie om internationaal door te groeien.

Zie je wel, hij begrijpt me niet

Hoewel IVR en spraaktechnologie in opmars zijn, blijkt uit onderzoek van bijvoorbeeld Forrester dat de consument kritisch, zo niet sceptisch tegenover 'machinepraat' staat. Maar waarom eigenlijk? Het werkt (in veel gevallen) en heeft onmiskenbare voordelen, in termen van snelheid of bereikbaarheid. Reden voor Telecats en de Vereniging Contactcenters Nederland om onderzoek te doen naar de beleving van consumenten ten aanzien van spraakherkenning. Daaruit blijkt dat men eigenlijk vooral wat onwennig staat tegenover spraakherkenning. Men is nog onbekend met het fenomeen, en onbekend maakt onbemind. Op basis van het onderzoek komen auteurs Arjan van Hessen en Ernst Kruize tot de volgende dwingende aanbevelingen: maak het systeem niet te goed, geef de gebruiker de kans eraan te wennen; reduceer de complexiteit, zodat men niet verdwaalt; kondig altijd aan dat bellers een computer aan de lijn krijgen; meld herkende spraak terug, geef direct bevestiging; en, last but not least, bied altijd de gelegenheid om direct te switchen van machine naar mens! (bron: Telecommerce)

TST-dag 2004

Op vrijdag 2 april 2004 vindt in Rotterdam een Vlaams-Nederlandse TST-dag plaats, georganiseerd door de Nederlandse Taalunie in samenwerking met INL, NWO, AWI en IWT.

Begonnen wordt met de presentatie van de Centrale voor Taal en Spraaktechnologie. Deze TST-centrale is verantwoordelijk voor het verwerven, beheren, onderhouden en verspreiden van diverse spraak- en taalmaterialen die voor het Nederlandstalig taalgebied met overheidsmiddelen zijn ontwikkeld. Vervolgens komen de resultaten aan bod van de projecten van het INL, de Commissie Lexicografische Vertaalsoorzieningen en het Corpus Gesproken Nederlands. Deze materialen worden vanaf dit jaar bij de TST-centrale opgenomen en aan de verschillende gebruikersgroepen ter beschikking gesteld. (Aanmelden via http://taalunieversum.org/taal/technologie/tst-dag2004_inschrijfformulier.doc)

Bedrijven bezorgd over beveiliging spraak- en datanetwerken

Uit onderzoek (in opdracht van Avaya) blijkt dat zowel Nederlandse als andere Europese bedrijven bezorgd zijn over de beveiliging van spraak- en datanetwerken. Het onderzoek van onderzoeksbureau GTMI toont aan dat meer dan de helft van de Europese bedrijven niet beschikt over het gewenste beveiligingsniveau.

Hoewel het belang van beveiliging in de directiekamers is doorgedrongen, is de veelzijdigheid van de risico's nog niet bij alle bedrijven bekend. Meer dan de helft van de respondenten blijkt zich niet bewust van de beveiligingsrisico's. De zorgen kindjes: hackers (64 procent) en de idee dat het spraaknetwerk even goed beveiligd is als het datanetwerk (59 procent). Het gebrek aan standaarden en een duidelijk stappenplan voor bedrijven onderstreept het belang van een duidelijk beveiligingsbeleid en de behoefte aan een marktorgaan dat de standaardisatie van beveiliging aanstuurt.

Uit de boeken

L E E S W I J Z E R

De praktijk is altijd een goede leermeester, maar zonder gedegen achtergrondkennis. Antal van den Bosch – verbonden aan de Universiteit van Tilburg, faculteit ILK/Computational Linguistics – bespreekt voor DIXIT elke keer een aantal relevante boeken of artikelen. Ditmaal *Machine Learning* van Tom Mitchell, en *Foundations of statistical natural language processing* van Christopher Manning en Hinrich Schuetze.

Romans worden doorgaans van kافت tot kافت gelezen. Hoe anders ver gaat het veel academische boeken. Het stereotiepe beeld is dat van een verstrooide professor die een boek uit de kast plukt en er vervolgens wild in gaat bladeren op zoek naar die ene formule. Een uitzondering zijn tekstboeken die bedoeld zijn als achtergrondmateriaal voor academische training, die daarom meestal een duidelijke opbouw hebben en die soms lezen als een trein – of, inderdaad, als een roman. In deze aflevering van DIXIT nemen we maar eens de trein, dus twee tekstboeken die allebei achtergrondtechnieken introduceren die de laatste tien jaar in groeiende mate zijn geïntegreerd in de taaltechnologie: *machine learning* en statistische modellen.

Machine Learning in sneltreinvaart

Machine Learning van Tom Mitchell (McGraw-Hill, 1997) is een lezenswaardig tekstboek dat inderdaad leest als de spreekwoordelijke trein. Het onderzoeksgebied binnen de kunstmatige intelligentie dat zich 'machine learning' noemt, schrijft Mitchell, 'is concerned with the question of how to construct computer programs that automatically learn with experience'. In een volkomen heldere stijl vertrekt het boek om hoofdstuk voor hoofdstuk de belangrijkste families van *machine learning*-algoritmen te behandelen. Via een algemeen begin (hoe kan een lerend systeem een concept leren, hoe kan het generaliseren uit voorbeelden) voert de trektocht langs beslissingsbomen, kunstmatige neurale netwerken, Bayesiaans leren, voorbeeld-gebaseerd leren, genetische algoritmen, regelinductie, analytisch leren, en *reinforcement learning*. In tussentijdse hoofdstukken wordt stilgestaan bij de praktische kant van evaluatieve methoden, en de theoretische achtergrond van *machine learning*, de computationele leertheorie. Mitchell doet in het hele boek zijn stinkende best om volstrekt begrijpelijk te zijn. Op sommige punten helpt

een wiskundige achtergrond, maar veel verder dan middelbare-schoolwiskunde en eenvoudige kansberekening gaat het niet.

De kracht van het boek is dat ieder hoofdstuk gezien kan worden als een vriendelijke maar grondige introductie in het onderwerp van dat hoofdstuk. Door gebruik te maken van voorbeelden worden abstracte concepten concreet gemaakt. Daarnaast wordt ieder hoofdstuk afgesloten met een overzicht van literatuur dat als tweede stap gelezen zou moeten worden mocht de lezer op het onderwerp verder willen. Mitchell is een erudiet wetenschapper met het engelengeuld om alles uit te leggen, en dat met een vriendelijkheid die veel andere erudiete figuren niet opbrengen. Om niet in gejubel te eindigen moet ik wel vaststellen dat het boek wat *state of the art* mist; het is ruim zes jaar oud en vermeldt weinig tot niets over, bijvoorbeeld, kernelmethode (support vector machines), boosting of ensemble-methoden. Maar gelukkig duurt het tegenwoordig hoogstens enige minuten websurfen voordat je daar de eerste tutorials en overzichtsartikelen over hebt gevonden.

Taalfundamenten

Foundations of statistical natural language processing van Christopher Manning en Hinrich Schuetze (MIT Press, 1999) is een dikke pil van 679 pagina's die weliswaar bedoeld is als tekstboek (inclusief oefeningen bij ieder hoofdstuk), maar dat mede door zijn compleetheid en daarmee gecorreleerde dikte meedingt naar de titel van heus naslagwerk. Het exemplaar dat ik hier op mijn bureau heb liggen, net terug van uitgeleend te zijn geweest aan een collega, staat bol van de papiertjes die op strategische posities voor bladwijzer spelen. Uit de posities van de papiertjes maak ik op dat mijn collega alles wilde weten over het Viterbi-algoritme. Wil je weten wat Good-Turing smoothing ook al weer is, en waar het goed voor is? Wat doet de backing-off estimator van Katz ook al weer? Het antwoord is terug te vinden in het eerste deel van het boek, dat waarschijnlijkheidsleer introduceert en methoden voor het schatten van kansen van voorkomens op grond van metingen.

Het tweede deel neemt taal bij de horens en begint bij het begin: woorden. Welke woorden komen statistisch vaak voor met welke andere woorden, of hoe kun je statistisch de verschillende betekenissen van een woord bepalen. Na woorden komen uiteraard zinnen. Het derde deel van het boek beschrijft de theorie van Markovmodellen en de 'fruitvliegjes'-toepassing die statistische modellen populair maakten in de taaltechnologie, part-of-speech tagging (het toekennen van syntactische woordsoorten aan woorden). Vervolgens worden probabilistische grammatica's en parsers besproken. In deel vier ten slotte worden de vleugels uitgespreid naar bekende toepassingen van statistische modellen in automatisch vertalen, informatie retrieval, en tekstclassificatie. Het boek is op een aantal plekken 'foundational' en gaat daarmee voor de geïnteresseerde leek soms wat te diep, maar op andere plekken beschrijft het methoden die enorm veel gebruikt worden (Markovmodellen bijvoorbeeld) op een heldere, tutorial-achtige manier die je op niet veel andere plekken zult aantreffen. Lezen dus

Heeft u suggesties voor een boekbespreking? Mail dat aan de redactie: redactie@diaalooiguitgevers.nl

TST Dictionary

Dictionary: zie **Spraaksynthese** en **Spraakherkenning**; zie ook **Fonetische Transcriptie**, zie ook...

Zoals elk vakgebied heeft ook TST (taal- en spraaktechnologie) een compleet eigen vocabulaire, inclusief volstrekt duistere afkortingen. Handig natuurlijk om een zeker aureool rond het vakgebied te garanderen en in stand te houden, maar bepaald niet bevorderlijk voor de onderlinge communicatie tussen TST-expert en leek (lees: klant). 'Multimodaal'? Hij bedoelt toch zeker 'multimediaal'! En IVR kunnen we inmiddels wel duiden, en ook TTS zal een belletje doen rinkelen, maar ASR, VXML, NLDS of DTMF?

Arjan van Hessen (Universiteit Twente) zal de mythische woorden- en begrippen-sluier, en daarmee het TST-vakgebied, voor u oplichten, resulterend in een heuse TST-dictionary (= de basis van iedere spraaksynthese en spraakherkenningapplicatie is het dictionary: een lijst met correct geschreven woorden en hun fonetische transcriptie(s). Alleen woorden die in het dictionary staan, kunnen herkend worden).

amplitude De amplitude is de grootte van de uitwijking van een geluidsgolf. Veel mensen denken dat een spraakherkenner je beter verstaat wanneer je harder en langzamer gaat praten. Oversturing van het signaal leidt dan tot nog vreemdere herkenningen, waardoor veel mensen nog harder gaan praten. Een gebruikersvriendelijke herkenner zou in dat geval eigenlijk moeten antwoorden: ik ben niet doof!

evaluatie Het meest onderschatte onderdeel van veel dialoog en *information retrieval*-systemen. Het dure en dikwijls saaie bepalen van de werking van een systeem wordt dikwijls vermeden met de uitroep: hij doet het best goed of 80 procent van de uitkomst is juist. Goede evaluatie is noodzakelijk om iets zinnigs over dialoogsystemen te zeggen: wat gaat er precies goed en vooral: wat gaat er fout!

information retrieval Het ophalen van de gewenste informatie uit een verzameling informatie. Een voorbeeld van IE is het ophalen van N krantenartikelen (uit een enorme verzameling kranten) die het meest overeenkomen met een ingevoerd stuk tekst. Begrippen die hierbij horen zijn *precision* en *recall*. *Precision* is de mate waarin de artikelen daadwerkelijk de gevraagde

informatie bevatten, *recall* de mate waarin alle relevante artikelen gevonden worden. Wanneer alle artikelen teruggegeven worden, is de *recall* hoog (namelijk 100 procent), maar de *precision* laag. Wanneer een zeer relevant artikel wordt teruggegeven, is de *precision* hoog maar de *recall* laag. Beoordeling van IR-systemen vereist een uiterst kostbare evaluatie: experts moeten ieder artikel op z'n relevantie beoordelen gegeven een gestelde vraag. Het is helaas de enige manier om de *precision* en *recall* van IR-systemen te beoordelen.

normalisatie Woorden die eigenlijk het zelfde zijn, maar slechts van elkaar verschillen in schrijfwijze (pannekoek en pannenkoek) moeten eerst herschreven (genormaliseerd) worden tot één en hetzelfde woord voordat ze in een statistisch taalmodel gebruikt kunnen worden. De talloze verschillende manieren waarop in het Nederlands het woord café geschreven wordt (cafe, caffee, kafe, kaffé enzovoort) kan tot gevolg hebben dat geen enkele versie van het woord bij de 10.000 meest gebruikte woorden hoort. Hierdoor zou het gesproken woord 'café' niet herkend kunnen worden.

spraaksynthese De generatie van spraak met behulp van in de computer opgeslagen regels en/of 'deeltjes

dictatorship [dik'tetəʃɪp] 0.1 **dictation** [ˈdikʃən] 0.1 **dictie** ⇒ *voordwoordkeus, dictie*.
dictionary [ˈdɪkʃənri] (mv.: -ies) 0.1 **con**.
dictum [ˈdɪktəm] (mv.: ook dicta) 0.1

spraak'. Deze deeltjes kunnen de elementaire klanken van een taal zijn (fonemen) combinaties van fonemen (difonen), lettergrepen of zelfs gehele woorden. De modernste spraaksynthese gebruikt een combinatie van dit alles. Het is handig om bijvoorbeeld de 10.000 meest gebruikte woorden van een taal in z'n geheel op te slaan. Vervolgens worden ook de 10.000 meest gebruikte woorddelen opgeslagen (bijvoorbeeld: ge-, be-, voor-, anti-, -ing, -heid). Woorden die niet met een combinatie van deze 20K items gemaakt kunnen worden, worden vervolgens gegenereerd met behulp van difonen. Deze manier van spraaksynthese vereist veel geheugen en snelle zoekalgoritmes maar levert in de regel goede resultaten.

taalmodel De statistische waarschijnlijkheid van woorden en woordcombinaties gegeven een bepaalde context. Zonder taalmodel is de kans op ieder woord even groot, zodat de woorden slechts bepaald worden door de opeenvolging van hun klanken. Een spraakherkenner kan dan geen onderscheid maken tussen de woorden 'word', 'wordt' en 'wort'. Een voorbeeld van spraakherkenning zonder taalmodel vinden we bij het noemen van je naam aan het loket. De kans dat je Hessen, Hesse, Esse of Essen heet, is wat de loketbeambte betreft even groot. Zeer duidelijke uitspraak of spelling moet in dit geval uitkomst bieden. Is de naam Hesse zeer veel voorkomend, dan is het echter waarschijnlijk dat de beambte toch Hesse hoort, hoe duidelijk je ook Hessen zegt (zeker wanneer de conversatie over de telefoon verloopt; zijn taalmodel (gebaseerd op het aantal keren dat hij de naam Hesse al gehoord heeft) *overruled* als het ware de naam Hessen.

Wijzigingen bestuur Stichting NOTaS

De afgelopen periode hebben er een paar wijzigingen plaatsgevonden in het bestuur van de stichting NOTaS. Na een aantal jaren energie, enthousiasme en doorzettingsvermogen te hebben gestoken in het opstarten van de stichting zijn Theo van de Heuvel, Geert Kobus en Norbert Mergen ermee gestopt. Er zijn nieuwe bestuursleden geïnstalleerd, te weten Arjan van Hessen (Universiteit van Twente & Telecats), Rob Boeyink (AskNow Solutions) en Debbie Jackson (Polderland).

Arjan en Rob gaan zich bezig houden met de communicatie en Debbie is voorzitter geworden. De andere bestuursleden zijn: Harry Bunt (Universiteit van Tilburg; secretaris) en Daniel Nachtegaal (Vocaliber; penningmeester). Theo, Geert en Norbert bedankt al jullie inspanningen van de afgelopen jaren!

Ook het secretariaat is inmiddels in andere handen overgegaan. Na een aantal jaren enthousiast het secretariaat van de NOTaS te hebben gerund heeft María Camarasa besloten te stoppen met haar werkzaamheden. Het secretariaat is nu in handen van Claudia Linsen. Wij danken María voor haar inzet en enthousiasme en wensen haar veel succes in de toekomst.

Team voor Taal

Woensdag 21 januari vond de eerste 'Team voor Taal' plaats, een NOTaS-bijeenkomst waar diverse deelnemers elkaar op de hoogte brachten van hun activiteiten. Een uitgelezen evenement voor inspiratie, netwerken en inhoudelijke verdieping.

Wat meteen duidelijk werd, is dat NOTaS niet alleen een prima vehikel is om de belangen van de aangesloten deelnemers 'breed' te behartigen maar dat ook tussen de verschillende deelnemers de nodige synergie valt te kapitaliseren. Met name waar het gaat om – al dan niet toegepast – wetenschappelijk onderzoek dat direct zijn weg kan vinden in commerciële toepassingen. Waarbij inmiddels 'proven' toepassingen weer een inspiratiebron kunnen vormen voor meer toegespitst onderzoek.

Door de kort maar krachtige presentaties van al aanwezige NOTaS-leden werden bedrijfsprofielen, oplossingen en academische doelstellingen uiteengezet waarbij ook de toegepaste technologieën werden besproken; intelligente woordenboeken, spraak-retrieval, self service, search engines, en text to speech-oplossingen.

Het bestuur van NOTaS greep deze gelegenheid te baat om enige wisselingen in de bestuurswacht mee te delen: Geert Kobus (Knowledge Concepts) heeft zijn voorzittershamer doorgegeven aan Debbie Kenyon-Jackson (Polderland), en Theo van de Heuvel (Polderland) zijn 'rekenmachine' aan de nieuwe penningmeester Daniel Nachtegaal (Vocalibur).

N.V. BOM wenst Stichting NOTaS en DIXIT veel succes!

Nu N.V. BOM om afscheid heeft genomen van Stichting NOTaS, wil zij een laatste woord richten aan de ondernemers van Stichting NOTaS, zowel de bedrijven als ondernemende universiteiten.

Na een aanloopperiode van drie jaar is Stichting NOTaS een volledig zelfstandig orgaan dat verschillende activiteiten ontplooit. Stichting NOTaS groeit en heeft een positie verworven in de Nederlandse taal- en spraaktechnologiesector. Dit is vooral te danken aan het enthousiasme van een vijftal ondernemers (GIOS Voice Professionals, Knowledge Concepts, ORCAvoice, Polderland Language and Speech Technology, Van Dale Data) en een tweetal universiteiten (Technische Universiteit Eindhoven en Universiteit van Tilburg) die aan de basis hebben gestaan van de stichting. Deze bedrijven en universiteiten hebben samen met de BOM het risico genomen om te investeren in versterking van de sector.

N.V. BOM houdt zich bezig met het initiëren en aanjagen van samenwerkingsprojecten in de regio Brabant. Het vergt vaak trek- en duwkracht om middellange en lange termijn projecten met bedrijven van de grond te krijgen. De ondernemers in NOTaS hebben bewezen dat ook MKB in staat is om risico's te nemen in activiteiten die niet direct resultaat en winst opleveren maar op langere termijn resulteren in nieuwe bedrijvigheid. Voorbeelden hiervan zijn de verschillende bilaterale en multilaterale samenwerkingsverbanden die zijn voortgekomen uit NOTaS en direct business hebben opgeleverd voor betrokken partijen. De infrastructuur van NOTaS heeft er nu tevens toe geleid dat netwerk en structuur zichtbaar worden door een tijdschrift: DIXIT.

N.V. BOM wenst Stichting NOTaS en Dixit een vruchtbare toekomst toe en bedankt betrokken partijen voor de samenwerking.

NOTaS

Deelnemers van Stichting NOTaS zijn: AskNow Solutions
Comsys - EliTech - GIOS Voice Professionals - Human Inference -
Institute for Logic Language and Computation (ILCC) - Knowledge
Concepts - KUB Faculteit Letteren - Het Nijmeegs Instituut voor
Spraak- en Taaltechnologie (NISTT) - ORCAvoice - PAT Document
Systemen - Polderland Language and Speech Technology - Q-go -
Technische Universiteit Eindhoven - TeleCats - TNO TPD -
Universiteit Twente Faculteit Informatica - Van Dale Lexicografie -
Vocalibur Language & Speechtechnology
Meer informatie: www.stichtingnotas.nl



Het is een gemiste kans en u weet het: klanten gebruiken uw website om producten en diensten af te nemen. Maar hoe vaak staan uw klanten toch nog in de telefonische wachtrij, terwijl een antwoord op hun vragen online beschikbaar is?

De Online Marketing & Self Service software van Q-go is gebaseerd op een unieke vorm van Natural Language Processing. Hierdoor kan uw klant intuïtief in eigen woorden zijn vragen stellen op de website. Het resultaat? Uw klant wordt begrepen en direct online geholpen.

Het voordeel voor u? Betere service, tevreden

klanten, inzicht in hun vraaggedrag en bovendien aanzienlijke kostenbesparingen. Q-go garandeert dat u uw investering in de software binnen 6 maanden zult terugverdienen. Klanten als Postbank, Rabobank, TPG, La Caixa, Deutsche Telekom en Telefónica maken succesvol gebruik van Q-go.

Wilt u meer weten en een ROI berekening laten maken?

Surf naar www.q-go.com of stuur een email naar astrid.smedema@q-go.com. Of bel naar 020 531 3800.

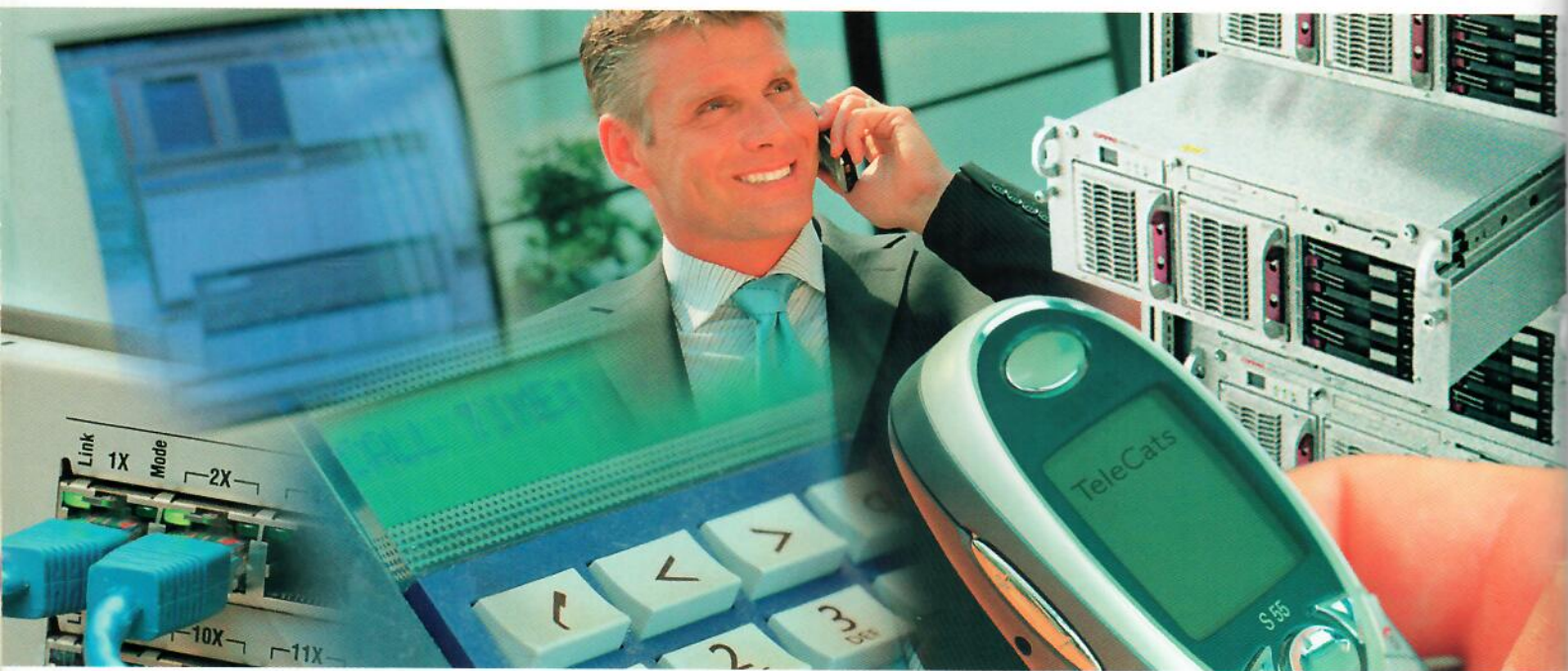


Online Marketing & Self Service

Spraakherkenning

Selfservice Telefonie

Interactive voice response



Veel organisaties besparen enorme kosten door het inzetten van selfservice telefoniesystemen. Selfservice telefoniesystemen handelen telefoongesprekken geheel of gedeeltelijk af m.b.v. IVR en spraakherkenning.

TeleCats is de partij om zaken mee te doen als we praten over selfservice met behulp van IVR en spraakherkenning. Organisaties die dat dagelijks ondervinden zijn o.a.: De Telegraaf, Eneco Energie, OHRA, SNT, KWF, het Ministerie van VROM, De Belastingdienst en NS reizigers.

Wilt u weten hoeveel kosten u kunt besparen door de selfservice telefoniesystemen van TeleCats in te zetten? Maak dan een afspraak met ons.



Internet: www.TeleCats.nl info@TeleCats.nl Telefoon: 053 4889900