

Proiect de cercetare-dezvoltare în Tehnologia Vorbirii

- îndrumar de proiect -

Ş.l. Dr. Ing. Horia Cucu

CUPRINS

1. ARHITECTURA GENERALĂ A UNUI SISTEM DE RECUNOAȘTERE AUTOMATĂ A VORBIRII .	5
1.1 Introducere în recunoașterea automată a vorbirii	5
1.2 Formalismul recunoașterii automate a vorbirii	8
1.3 Arhitectura unui sistem de recunoaștere automată a vorbirii	9
2. RESURSE FONETICE, ACUSTICE ȘI LINGVISTICE NECESARE ÎN CONSTRUCȚIA UNUI SISTEM DE RECUNOAȘTERE A VORBIRII CONTINUE	13
2.1 Generalități	13
2.2 Elemente de fonetică a limbii române	14
2.3 Dicționarul fonetic	14
2.4 Baza de date de vorbire	17
2.5 Corpusul de text	20
2.6 Concluzii	21
3. PRINCIPIILE DE BAZĂ ALE MODELĂRII ACUSTICE	23
3.1 Parametri acustici	23
3.2 Platforma HMM/GMM	24
3.3 Selecția, dependența de context și clusterizarea unităților de vorbire	26
4. PRINCIPIILE DE BAZĂ ALE MODELĂRII LIMBAJULUI NATURAL	29
4.1 Modele de limbă statistice	30
4.2 Gramatici cu stări finite (FSG – finite state grammar)	32

4.3	JavaSpeech Grammar Format (JSGF)	33
5.	EVALUAREA SISTEMELOR DE RECUNOAȘTERE A VORBIRII CONTINUE	45
5.1	Introducere	45
5.2	Evaluarea sistemului de recunoaștere din punctul de vedere al domeniului vorbirii	45
5.3	Evaluarea completă a sistemului de recunoaștere a vorbirii	48
5.4	Concluzii	57
6.	CONSTRUCȚIA UNUI SISTEM SIMPLU DE RECUNOAȘTERE A VORBIRII CONTINUE	59
6.1	Introducere	59
6.2	Înregistrarea unei baze de date de clipuri audio	61
6.3	Crearea dicționarului fonetic și a listei de foneme	62
6.4	Antrenarea modelului acustic	64
6.5	Crearea modelului de limbă (gramaticii)	69
6.6	Evaluarea sistemului de recunoaștere și interpretarea rezultatelor	70
6.7	Optimizarea modelului acustic	74
6.8	Construcția unui sistem de recunoaștere independent de vorbitor	83
6.9	Redactarea raportului final al proiectului de cercetare-dezvoltare	83
6.10	Desfășurarea activității pe server-ul de dezvoltare	85
	BIBLIOGRAFIE	87
	ANEXA 1. DIAPOZITIVE DE PREZENTARE	89

1. ARHITECTURA GENERALĂ A UNUI SISTEM DE RECUNOAȘTERE AUTOMATĂ A VORBIRII

1.1 Introducere în recunoașterea automată a vorbirii

Procesul recunoașterii automate a vorbirii (RAV) vizează transformarea unui semnal audio ce conține vorbire într-o succesiune de cuvinte. Înțelegerea automată a vorbirii extinde scopul RAV (generarea unei secvențe de cuvinte) încercând să producă informații de natură semantică privind înțelesul propoziției generate de sistemul de RAV. În cazul în care semnalul acustic de intrare conține vorbire provenind de la mai mulți vorbitori, procesul de RAV poate fi privit ca având două etape: diarizare (etapă ce încearcă să răspundă la întrebarea: “cine și când a vorbit?”) și transcriere (etapă ce încearcă să răspundă la întrebarea: “ce a spus?”).

Recunoașterea automată a vorbirii are o gamă largă de aplicabilitate. Domeniul cel mai important pare a fi acela al interfețelor hands-free și eyes-free. Există multe aplicații în care utilizatorii au nevoie să-și folosească mâinile și ochii pentru altceva, iar vorbirea rămâne singura lor alternativă de a fi eficient în dialogul cu calculatorul. Mai mult decât atât, vorbirea este cel mai natural mijloc de comunicare pentru ființele umane. Alte domenii importante în care RAV se aplică cu succes sunt sistemele de dialog pentru call-centere și sistemele de traducere speech-to-speech. Traducerea speech-to-speech este în acest moment un subiect foarte „la modă” în mai multe centre de cercetare academice și industriale. Nu în ultimul rând, RAV se poate utiliza și pentru dictare: transcrierea unui monolog al unui vorbitor anume. Transcrierea unui astfel de monolog este utilă în diverse domenii, de la cea mai simplă utilizare pentru editarea documentelor de către jurnaliști, scriitori, etc., până la utilizarea pentru transcrierea discuțiilor din timpul proceselor de judecată. Fiecare dintre aceste aplicații este de obicei mult mai restrictivă decât problema generală care impune transcrierea automată a vorbirii naturale, continue, provenind de la un vorbitor necunoscut, în orice mediu (zgomotos sau nu). Diferitele surse de variabilitate în vorbire, care vor fi discutate în continuare, fac sarcina generală de RAV una foarte dificilă. Cu toate acestea, în multe situații practice, variabilitatea este restricționată. De exemplu, vorbitorul ar putea fi cunoscut (în loc un număr oarecare de vorbitori necunoscuți sistemului) sau vorbirea ar putea fi dictată (în loc de vorbire spontană) sau mediul de înregistrare ar putea ne-zgomotos și ne-reverberant. În recunoașterea automată a vorbirii se face o distincție între modulul care

abordează variabilitatea acustică (modelul acustică) și modulul care tratează incertitudinea lingvistică (modelul de limbă).

Unul dintre cei mai importanți factori de care depinde dificultatea procesului de transcriere este sarcina de recunoaștere. Aceasta include specificitatea limbii, numărul de cuvinte ce pot fi pronunțate și incertitudinea lingvistică a sarcinii de recunoaștere. Diversele limbi vorbite prezintă provocări diferite pentru un sistem de recunoaștere a vorbirii. Pentru un număr covârșitor de mare de limbi nu există suficiente resurse acustice (baze de date de vorbire) și lingvistice (corpusuri de text). Aceste așa-numite limbi slab dotate sunt vorbite de un număr mare de oameni, însă până acum s-au achiziționat prea puține resurse acustice și lingvistice pentru a putea dezvolta un sistem de RAV. În consecință, proiectarea și crearea unui sistem de RAV trebuie să prevadă și colectarea acestor resurse.

Alte limbi “suferă” de o morfologie foarte complexă. De exemplu limbile cu morfologie complexă, cum ar fi Franceza sau Româna, au vocabulare de dimensiuni mai mari decât limbile cu morfologie simplă, cum ar fi Engleza. Timpul prezent al verbului *a merge* are cinci forme diferite în limba Română: *merg*, *mergi*, *merge*, *mergem* și *mergeți*, și șase forme diferite în Franceză: *vais*, *vas*, *va*, *allons*, *allez*, *vont*. În Engleză, același verb, la același timp, are doar două forme morfologice diferite: *go* și *goes*. Limbile Germană și Turcă sunt două dintre așa-numitele limbi aglutinative. În aceste limbi se pot forma cuvinte noi prin concatenarea altor cuvinte sau morfeme. Astfel dimensiunea vocabularului de cuvinte pentru o limbă aglutinativă este mult mai mare, acest lucru complicând sarcina de recunoaștere automată a vorbirii.

Dimensiunea vocabularului este de asemenea un factor important care influențează dificultatea sarcinii de RAV. Este evident faptul că o sarcină de recunoaștere de cifre (cu un vocabular de 10 cuvinte) este mult mai simplă decât, de exemplu, recunoașterea vorbirii spontane (cu un vocabular de 64.000 de cuvinte). Cu toate acestea, dimensiunea mare a vocabularului nu înseamnă neapărat că sarcina de recunoaștere va fi mai dificilă. Incertitudinea lingvistică a potențialelor discursuri ce trebuie recunoscute constituie de asemenea un factor important. De exemplu, o sarcină de RAV specifice turismului (cu un vocabular de 64 de mii de cuvinte conținând în general nume de locuri, hoteluri, restaurante, etc.) nu este atât de dificilă ca o sarcină de recunoaștere a vorbirii spontane (cu un vocabular de aceeași dimensiune). Incertitudinea lingvistică (perplexitatea) scăzută a primeia o face mult mai simplă.

Procentul de cuvinte incorecte obținut pentru diverse sarcini de recunoaștere este prezentat în Tabelul 1.1. Datele se referă la sisteme state-of-the-art de RAV pentru limba Engleză [Jurafsky, 2009]. Rata de cuvinte eronate (word error rate - WER) este criteriul de performanță standard utilizat în evaluarea sistemelor de RAV. Pentru mai multe informații despre acest criteriu de performanță consultați platforma „Evaluarea unui sistem de recunoaștere a vorbirii continue”.

Tabelul 1.1. Rezultate WER raportate în anul 2005 pentru diverse sarcini de RAV [Jurafsky, 2009]

Sarcina de RAV	Dimensiunea vocabularului	WER [%]
TI Digits	11 cuvinte	0.55
Wall Street Journal read speech	5000 cuvinte	3.0
Wall Street Journal read speech	20000 cuvinte	<6.6
Broadcast News	64000+ cuvinte	9.9
Conversational Telephone Speech	64000+ cuvinte	20.7

Un alt factor important care influențează dificultatea procesului de RAV este stilul vorbirii. Stilul vorbirii se referă la cât de fluentă, naturală sau conversațională este vorbirea ce trebuie recunoscută. În mod evident, recunoașterea de cuvinte izolate, unde cuvintele sunt separate de pauze de liniște, este mult mai simplă decât recunoașterea vorbirii continue, unde cuvintele sunt pronunțate legat și vorbirea trebuie segmentată înainte de recunoaștere. De fapt, primele sisteme de recunoaștere automată a vorbirii puteau rezolva problema localizării granițelor între cuvinte numai dacă vorbitorul făcea pauze generoase în pronunțarea lor: Dragon Dictation [Baker, 1989] este un bun exemplu de un sistem de recunoaștere de cuvinte izolate cu vocabular extins.

Sarcinile de recunoaștere a vorbirii continue pot fi mai departe clasificate în funcție de dificultatea lor. De exemplu, sarcina de recunoaștere a vorbirii citite este mult mai simplă decât sarcina de recunoaștere a vorbirii conversaționale sau spontane. Variabilitatea acustică mai mare face ca sarcina urmă să fie mai dificilă. Această diferență în dificultate între sarcinile de vorbire continue este reflectată în ratele crescute de eroare pentru recunoașterea vorbirii spontane, comparativ cu recunoașterea vorbirii citite (a se vedea Tabelul 1.1).

Dificultatea și, în consecință, acuratețea procesului de recunoaștere a vorbirii este, de asemenea, influențată de mediul acustic în care este înregistrată vorbirea și de canalul de transmisie. În afara studiourilor de înregistrare sau laboratoarelor de cercetare, există, de obicei, mai multe surse acustice, inclusiv alți vorbitori, zgomot de fundal, etc.. În cele mai multe cazuri separarea semnalele acustice diferite este o problemă foarte dificilă. Microfonul folosit pentru înregistrare are, de asemenea, un impact semnificativ asupra acurateții recunoașterii vorbirii.

Sistemele comerciale de dictare și majoritatea experimentelor de cercetare în domeniul recunoașterii vorbirii sunt realizate cu microfoane de înaltă calitate. Alte tipuri de microfoane vin cu dezavantaje diferite, care contribuie la calitatea sistemului de RAV. Variațiile în canalul de transmisie apar din cauza mișcărilor capului vorbitorului relativ la microfon și din cauza transmisiei printr-o rețea de telefonie sau internet. Cea mai mare diferență dintre acuratețea sistemelor de RAV comparativ cu acuratețea cu care omul recunoaștere vorbirea apare în situația în care semnalul acustic este poluat cu zgomot aditiv

mare, sau provine de la mai multe surse acustice, sau este înregistrat într-un mediu reverberant. Recunoașterea semnalului de vorbire zgomotos, cu un raport semnal-zgomot scăzut, poate avea rezultate de 2 până la 4 ori mai slabe, comparativ cu recunoașterea semnalului de vorbire nezmotos.

În cele din urmă, caracteristicile vorbitorilor au, de asemenea, un impact semnificativ asupra preciziei unui sistem de recunoaștere a vorbirii. Variabilitatea acestor caracteristici include accentul vorbitorului, limba sau dialectul folosit, faptul că vorbitorul este nativ sau nu (dacă limba utilizată limba lui mamă), rapiditatea pronunției, vârsta vorbitorului, și, desigur, diferențele anatomice și fiziologice care influențează producerea vorbirii. Mai mult decât atât, vorbitori diferiți prezintă grade diferite de variabilitate intrinsecă, bazate pe starea emoțională, probleme temporare de sanatate, etc. Variabilitatea inter-vorbitori poate fi soluționată prin proiectarea de sisteme de RAV dependente de vorbitor (specifice fiecărui vorbitor). Dezavantajul aici este faptul că pentru fiecare vorbitor nou trebuie antrenat un nou model acustic. Acest lucru conduce la un sistem mai complex și ridică mai multe probleme de antrenabilitate (date insuficiente de antrenare pentru fiecare vorbitor și altele). Pe de altă parte, sistemele de RAV independente de vorbitor sunt mai simple și mai flexibile (acestea pot fi folosite pentru a recunoaște vorbirea oricărui vorbitor). Cu toate acestea, un sistem independent de vorbitor este mai puțin performant pentru un anumit vorbitor în comparație cu un sistem dependent de vorbitor creat special pentru respectivul vorbitor (în cazul în care există suficiente date de antrenare pentru respectivul vorbitor). Cu toate că algoritmi de adaptare la vorbitor au progresat mult în ultimii 15 de ani, adaptabilitatea și robustețea sistemelor de RAV pentru recunoașterea diverselor voci este încă foarte limitată în comparație cu performanța umană.

Paradigma state-of-the-art în domeniul recunoașterii vorbirii continue cu vocabular extins este modelul Markov ascuns (Hidden Markov Model - HMM). Infrastructura HMM a fost introdusă ca un candidat viabil pentru partea de modelare acustică a recunoașterii vorbirii în 1975 [Baker, 1975]. În particular pentru sistemele de RAV cu vocabular extins, modelul acustic bazat pe HMM este utilizat împreună cu un model de tip n-gram, care este responsabil pentru partea de modelare limbă. Modelele lingvistice statistice (n-gram) au devenit soluția stat-of-the-art pentru modelarea limbajului în primul rând datorită extinderii extraordinare a Internetului, care a furnizat date suficiente pentru antrenarea corespunzătoare a acestor sisteme.

1.2 Formalismul recunoașterii automate a vorbirii

Procesul recunoașterii automate a vorbirii (RAV) vizează transformarea unui semnal audio ce conține vorbire într-o succesiune de cuvinte. Recunoașterea automată a vorbirii este unul dintre primele domenii în care tehnicile de modelare statistică (în care modelele se creează pe baza unor catințări mari de date) s-au impus ca și standard. Platforma statistică

pentru recunoașterea automată a vorbirii a fost creată și dezvoltată de-a lungul a două decenii de către Baker [Baker, 1975], o echipă de cercetători de la compania IBM [Jelinek, 1976; Bahl, 1983] și o echipă de cercetare-dezvoltare de la compania AT&T [Levinson, 1983; Rabiner, 1989].

Procesul de transformare a vorbirii în text poate fi formulat într-o manieră statistică astfel: *Care este cea mai probabilă secvență de cuvinte W^* în limba L , dat fiind mesajul vorbit X ?*

Reprezentarea formală utilizează funcția argmax, care selectează argumentul ce maximizează probabilitatea secvenței de cuvinte:

$$W^* = \arg \max_W p(W | X) = \arg \max_W \frac{p(X | W)p(W)}{p(X)} = \arg \max_W p(X | W)p(W) \quad (1.1)$$

Dezvoltarea din Ecuația 1.1 are la bază regula Bayes și s-a făcut ținând cont de faptul că probabilitatea mesajului vorbit $p(X)$ este independentă de secvența de cuvinte W . Ultimul rezultat evidențiază doi factori care pot fi estimați direct. Problema inițială (găsirea secvenței de cuvinte pe baza mesajului vorbit) a fost împărțită în două sub-probleme mai simple: a) estimarea probabilității apriori a secvenței de cuvinte $p(W)$ și b) estimarea probabilității mesajului vorbit dată fiind secvența de cuvinte pronunțată $p(X|W)$. Primul factor poate fi estimat utilizând exclusiv un *model de limbă*, iar cel de-al doilea poate fi estimat cu ajutorul unui *model acustic*. Cele două modele pot fi construite independent așa cum se va vedea în secțiunea următoare, dar vor fi folosite împreună pentru a decoda un mesaj vorbit, așa cum arată Ecuația 1.1.

1.3 Arhitectura unui sistem de recunoaștere automată a vorbirii

Arhitectura generală a unui sistem de recunoaștere automată a vorbirii (RAV) este prezentată în Figura 1.1. Din figură reies două lucruri esențiale în ceea ce privește procesul de recunoaștere a vorbirii: a) recunoașterea se face utilizând o serie de parametri vocali extrași din semnalul vocal (nu direct, utilizând mesajul vorbit) și b) recunoașterea se face pe baza unor modele (acustic, fonetic și lingvistic) dezvoltate în prealabil.

Modelul acustic are rolul de a estima probabilitatea unui mesaj vorbit, dată fiind o succesiune de cuvinte. În sistemele de RVC actuale modelul acustic nu folosește cuvinte ca unități acustice de bază pentru că: a) fiecare sarcină de RAV are un vocabular de cuvinte diferit pentru care nu există modele deja antrenate și nici date de antrenare disponibile și b) numărul de cuvinte diferite dintr-o limbă este prea mare. În locul cuvintelor, se utilizează unități acustice de bază sub-lexicale, de exemplu foneme, sau, mai recent, chiar unități sub-fonemice numite senone. Prin urmare, modelul acustic este format dintr-un set de modele

pentru foneme (sau senone) care se conectează în timpul procesului de decodare pentru a forma modele pentru cuvinte și apoi modele pentru succesiuni de cuvinte. Acestea sunt utilizate în final pentru a estima probabilitatea ca mesajul rostit de vorbitor să fie format dintr-o succesiune sau alta de cuvinte. Experiența a arătat că această abordare generativă pentru modelele acustice poate fi foarte bine implementată utilizând modele Markov ascunse (hidden Markov models - HMM). Pentru mai multe informații despre modul cum funcționează modelele acustice consultați platforma „Principiile de bază ale modelării acustice”.

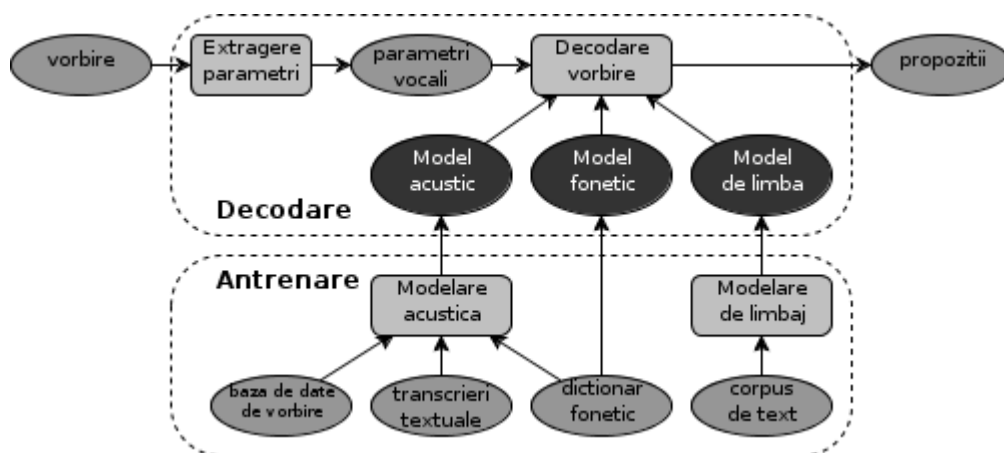


Figura 1.1. Arhitectura generală a unui sistem de recunoaștere a vorbirii

Modelul de limbă este utilizat în timpul decodării pentru a estima probabilitățile tuturor secvențelor de cuvinte din spațiul de căutare. În general, rolul unui model de limbă este de a estima probabilitatea ca o secvență de cuvinte $W = w_1, w_2, \dots, w_n$, să fie o propoziție validă a limbii. Probabilitatea acestor secvențe de cuvinte ajută foarte mult modelul acustic în procesul de decizie. De exemplu, aceste două fraze: *ceapa roșie este sănătoasă* și *ce apar oști ie este sănătoasă* sunt similare din punct de vedere acustic, însă cea de-a doua nu are niciun sens. Rolul modelului de limbă este acela de a asigura o probabilitate mult mai mare primei secvențe de cuvinte și, prin urmare, de a ajuta sistemul de RAV să decidă în favoarea acesteia.

Modelul fonetic are rolul de a conecta modelul acustic (care estimează probabilitățile acustice ale *fonemelor*) cu modelul de limbă (care estimează probabilitățile secvențelor de *cuvinte*). De cele mai multe ori modelul fonetic este un dicționar de pronunție care asociază fiecărui cuvânt din vocabular una sau mai multe secvențe de foneme adecvate, reprezentând modul în care se poate pronunța respectivul cuvânt.

Figura 1.1 ilustrează, de asemenea, procesele implicate în dezvoltarea unui sistem de RAV, dar și resursele necesare creării modelelor acustice, lingvistice și fonetice. Modelul acustic se construiește strict pe baza unui set de clipuri audio înregistrate asociate cu transcrierea textuală a mesajelor vorbite și a unui dicționar fonetic ce cuprinde toate cuvintele din respectiva transcriere textuală. În cazul sistemelor de RVC cu vocabular extins se

folosesc modele de limbă statistice, care se construiesc utilizând corpusuri de text de dimensiuni cât mai mari și cât mai adaptate domeniului din care fac parte mesajele vorbite ce trebuie decodate. În sistemele de RVC cu vocabular redus se folosesc preponderent modele de limbă de tip gramatică cu reguli, iar pentru construcția acestora nu este nevoie de resurse textuale.

2. RESURSE FONETICE, ACUSTICE ȘI LINGVISTICE NECESARE ÎN CONSTRUCȚIA UNUI SISTEM DE RECUNOAȘTERE A VORBIRII CONTINUE

2.1 Generalități

Sistemele actuale de recunoaștere a vorbirii continue (RVC) transformă semnalul vocal în text pe baza unor modele (acustic, fonetic și lingvistic) dezvoltate în prealabil. Figura 2.1 ilustrează arhitectura generală a sistemelor de recunoaștere a vorbirii continue punând accent pe resursele necesare antrenării celor trei modele.

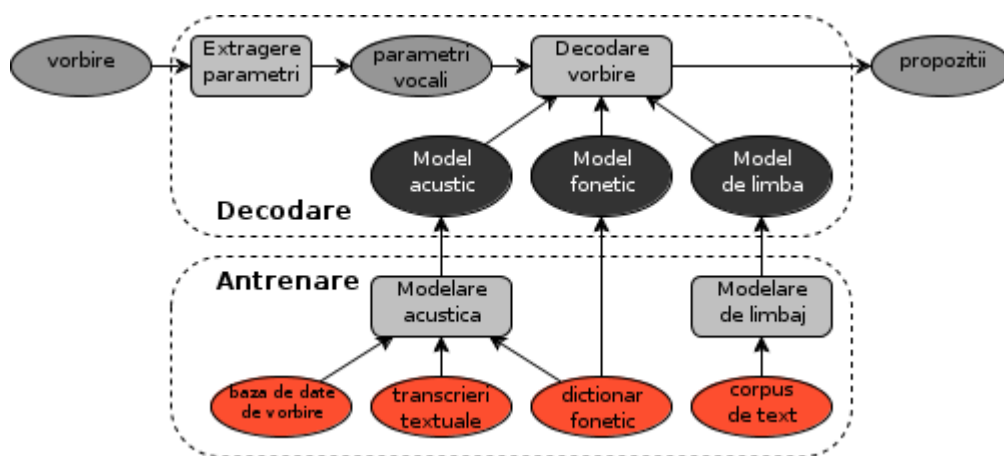


Figura 2.1. Resursele necesare construcției unui sistem de RVC

Așa cum arată Figura 2.1, modelul acustic, care modelează de obicei fonemele limbii, se contruiește pe baza unui set de clipuri audio înregistrate și a transcrierilor textuale aferente. De asemenea, pentru antrenarea acestui model este necesară și existența unui dicționar fonetic care să specifice modul de pronunție al cuvintelor din transcrierile textuale.

Dicționarul fonetic stă și la baza dezvoltării modelului fonetic necesar în procesul de decodare. Dacă dicționarul fonetic cuprinde mai multe pronunții pentru același cuvânt, în cadrul modelului fonetic aceste pronunții sunt comasate formând un model de pronunție pentru respectivul cuvânt.

În cele din urmă, modelul de limbă necesar unui sistem de recunoaștere a vorbirii continue se construiește utilizând un corpus de text de mari dimensiuni. Din această resursă se extrag statisticile caracteristice limbii, iar acestea se vor utiliza în procesul de decodare pentru a atribui probabilități diverselor cuvinte și secvențe de cuvinte propuse de modelul acustic.

2.2 Elemente de fonetică a limbii române

În sistemele actuale de recunoaștere a vorbirii continue modelul acustic nu modelează acustica cuvintelor limbii în mod direct, ci indirect, folosind unități acustice sub-lexicale, numite foneme sau chiar unități acustice sub-fonemice.

Fonemul este unitatea de sunet fundamentală din limbile vorbite care ajută la diferențierea cuvintelor și morfemelor. Prin modificarea unui fonem al unui cuvânt, se generează fie un cuvânt inexistent, dar perceput ca diferit de către vorbitorii limbii, fie un cuvânt cu alt sens.

În limba română se folosesc 7 vocale de bază și două vocale împrumutate, 4 semi-vocale și 22 de consoane, conform reprezentării din Tabelul 2.1.

2.3 Dicționarul fonetic

Importanța dicționarului fonetic într-un sistem de recunoaștere automată a vorbirii este evidențiată în Figura 2.1. Se poate observa că dicționarul fonetic este utilizat în partea de dezvoltare a unui sistem de recunoaștere, mai precis la antrenarea modelului acustic și modelului fonetic.

Un dicționar fonetic este un instrument lingvistic care specifică modul în care se pronunță cuvintele unei limbi. Altfel spus, dicționarul fonetic face corespondența între forma scrisă și forma fonetică a cuvintelor unei limbi. Forma fonetică a unui cuvânt este o succesiune de foneme. Într-un sistem de recunoaștere a vorbirii continue, dicționarul fonetic are rolul de a face legătura între modelul acustic (care modelează modul de producere a *sunetelor specifice limbii*) și modelul de limbă (care modelează modul în care se succed *cuvintele limbii*). În consecință, dicționarul fonetic trebuie să conțină toate cuvintele posibile pentru o anumită sarcină de recunoaștere a vorbirii alături, bineînțeles, de transcrierile fonetice ale acestor cuvinte.

Tabelul 2.1. Fonemele limbii române

Fonemul			Exemple de cuvinte	
Tip	Simbol IPA	Simbol intern	Forma scrisă	Forma fonetică
vocale	a	a	sat	s a t
	e	e	mare	m a r e
	i	i	lift	l i f t
	o	o	loc	l o c
	u	u	șut	s 1 u t
	ə	a1	gură	g u r a 1
	ɨ	i2	între	i 2 n t r e
vocale împrumutate	ɣ	y	ecru	e c r y
	ø	o2	bleu	b l o 2
semivocale	ɛ̃	e1	deal	d e 1 a l
	j	i3	fiară	f i 3 a r a 1
	ɔ̃	o1	oase	o 1 a s e
	w	w	sau	s a w
consoane	c	k2	chem	k 2 e m
	b	b	bar	b a r
	p	p	par	p a r
	k	k	acum	a k u m
	tʃ	k1	cenușă	k 1 e n u s 1 a 1
	g	g	galben	g a l b e n
	ʈʂ	g1	girafă	g 1 i r a f a 1
	ʃ	g2	unghi	u n g 2
	d	d	dar	d a r
	t	t	tot	t o t
	f	f	fața	f a t 1 a
	v	v	vapor	v a p o r
	h	h	harta	h a r t a
	ʒ	j	ajutor	a j u t o r
	ʃ	s1	coș	k o s 1
	l	l	lac	l a c
	m	m	măr	m a 1 r
	n	n	nas	n a s
	s	s	sare	s a r e
	z	z	zar	z a r
	r	r	risc	r i s k
	ʦ	t1	țaran	t 1 a 1 r a n
consoană palatalizată	j	i1	tari	t a r i 1

Procesul de construcție al unui dicționar fonetic presupune fonetizarea (crearea formei fonetice) cuvintelor vocabularului sarcinii de recunoaștere a vorbirii. Regulile de fonetizare ale cuvintelor limbii române au fost studiate și explicate destul de puțin în literatura de specialitate. De aceea, orice proces de construcție al unui dicționar fonetic trebuie să înceapă cu stabilirea unor reguli de fonetizare și excepții de la aceste reguli.

Construcția dicționarului fonetic se poate realiza manual, semi-automat sau chiar automat. În cazul în care vocabularul sarcinii de RVC este redus (câteva zeci-sute de cuvinte), crearea dicționarului fonetic se poate realiza manual într-un timp suficient de scurt. Dacă însă sarcina de RVC necesită utilizarea unui vocabular extins (mii sau chiar zeci de mii de cuvinte), atunci fonetizarea manuală a acestui vocabular nu mai este o opțiune.

Din fericire pentru limba română corespondența între forma scrisă și pronunția cuvintelor este de cele mai multe ori bijectivă. În general fiecărei litere îi corespunde întotdeauna același fonem. Există doar câteva clase de ambiguitate (vezi Tabelul 2.2) în care aceeași literă se pronunță diferit în funcție de contextul în care apare.

Tabelul 2.2. Clasele de ambiguitate în fonetizarea cuvintelor limbii române

Litera	Fonemul	Exemple
e	e	cer, rest, rece
	e1	deal, te-am, ceainic, gheorghe
	i3 e	eu, el, eram, erai
i	i	fir, citit, prins
	i1	tari, rupi, deci
	i3	iarbă, doi, fier
o	o	gol, potrivit
	o1	voal, soare
u	u	sud, furtun
	w	băcăuan, nouă
c	k	casă, sac
	k1	cer, circ
	k2	chiar, ochelari
g	g	gară, olog
	g1	girafă, george
	g2	unghi, ghem
h	h	hartă, rahat
	mut	ochi, unghi
x	k s	pix, fix
	g z	exemplu, examen

Odată ce setul de reguli a fost stabilit, fonetizarea se poate realiza și automat, utilizând aplicații software care implementează setul de reguli de fonetizare stabilit a priori. Chiar dacă la prima vedere ar părea simplă, implementarea acestor reguli nu este deloc trivială. Pentru rezolvarea celor mai des întâlnite clase de ambiguități (de exemplu ambiguitatea vocală/semivocală) este nevoie de o informație suplimentară: modul de despărțire al

cuvintelor în silabe. Fonetizarea manuală se bazează pe faptul că expertul uman cunoaște aceste reguli de despărțire în silabe, însă, în cazul în care se dorește automatizarea acestui proces, regulile de despărțire trebuie și ele formalizate.

O alternativă la aplicațiile de fonetizare automată pe bază de reguli sunt aplicațiile de învățare automată statistică. Utilizând un dicționar fonetic creat manual se poate antrena un model statistic de fonetizare, care poate fi utilizat în continuare pentru fonetizarea unor vocabulare mult mai extinse.

2.4 Baza de date de vorbire

Importanța bazei de date de vorbire într-un sistem de recunoaștere automată a vorbirii (RAV) este evidențiată în Figura 2.1. Se poate observa că această bază de date este utilizată în partea de dezvoltare a unui sistem de recunoaștere, mai precis la antrenarea modelului acustic. Baza de date de vorbire este utilizată și în procesul de evaluare a unui sistem de recunoaștere de vorbire, proces ce nu este ilustrat în Figura 2.1.

O bază de date de vorbire completă cuprinde trei componente:

- un set de clipuri audio ce conțin vorbire;
- un set de fișiere text ce conțin textul ce a fost pronunțat în clipurile audio și eventual marcaje temporale pentru fiecare cuvânt/fonem;
- informații suplimentare privind stilul și domeniul vorbirii, identitatea vorbitorului, etc.

În majoritatea cazurilor clipurile audio ce conțin vorbirea sunt eșantionate cu frecvența de 16kHz, dimensiunea eșantioanelor fiind de 16 biți, iar formatul fișierelor ms wav.

Baza de date de vorbire utilizată în antrenarea unui sistem de RAV este un element cheie care afectează în mod direct performanța sistemului: rata de eroare și viteza de recunoaștere. Cele mai importante caracteristici ale bazei de date de vorbire (cele care determină în mod direct calitatea acesteia) sunt:

- dimensiunea bazei de date (numărul de ore de vorbire, numărul de vorbitori),
- stilul vorbirii (cuvinte izolate, vorbire continuă citită, vorbire conversațională),
- variabilitatea (calitatea înregistrărilor, zgomotul de fundal, variabilitatea vorbitorilor)

Datorită faptului că modelarea acustică este abordată statistic în sistemele RAV moderne, dimensiunea bazei de date de antrenare este foarte importantă. Tutorialul de antrenare acustică CMU Sphinx sugerează niște valori numerice concrete în ceea ce privește dimensiunea unei bazei de date de vorbire (Tabelul 2.3). Sarcina de comandă și control este

o sarcină tipică de recunoaștere a vorbirii pseudo-continuă cu vocabular redus (RVC-VR), în timp ce sarcina de dictare este una tipică de recunoaștere a vorbirii continue cu vocabular extins (RVC-VE). De reținut este faptul că dezideratul independenței de vorbitor pare foarte dificil de atins, deoarece pentru aceasta este nevoie de material vorbit de la aproximativ 200 de vorbitori. Acest număr ar putea părea exagerat, dar variabilitatea vorbitorului este într-adevăr un factor important și poate fi depășită doar prin modelarea diferitelor pronunții posibile ale fonemelor. Acest lucru poate să fie realizat prin înregistrarea mai multor vorbitori.

Tabelul 2.3. Dimensiuni ale bazelor de date de vorbire sugerate de CMU Sphinx

Sarcina RAV	Sistem dependent de vorbitor	Sistem independent de vorbitor
comandă și control (RVC-VR)	1 oră de înregistrări, 1 vorbitor	5 ore de înregistrări, 200 de vorbitori
dictare (RVC-VE)	10 ore de înregistrări, 1 vorbitor	50 ore de înregistrări, 200 de vorbitori

Achiziționarea unei baze de date de vorbire se poate face în două moduri: a) înregistrarea unor texte predefinite sau b) etichetarea unor materiale audio ce contin vorbire. În funcție de aplicația ce urmează a fi dezvoltată, vor prima avantajele sau dezavantajele uneia dintre cele doua metode de achiziție.

2.4.1 Înregistrarea unor texte predefinite

Înregistrarea unor texte predefinite implică alegerea locului de înregistrare (studio, laborator, etc.), alegerea microfonului, alegerea stației de înregistrare (placa audio folosită pentru achiziție, amplitudinea semnalului, etc.) și, în sfârșit alegerea textelor ce urmează a fi înregistrate. Alegerea textelor este un proces foarte important de care depinde timpul și costul achiziției bazei de date, dar și corectitudinea ei (corespondența între texte și materialul înregistrat trebuie să fie perfectă). De aceea, în alegerea textelor trebuie să se țină cont de următoarele aspecte:

- frazele ar trebui să fie corecte din punct de vedere lexical și gramatical. În afara greșelilor gramaticale trebuie evitate și schimbările de topică a propoziției, schimbări ce pot pune vorbitorul în dificultate;
- frazele ar trebui să fie scurte (maxim 15-20 de cuvinte, în funcție de lungimea cuvintelor) pentru a putea fi pronunțate fără ezitări sau respirații;
- frazele ar trebui să conțină numai cuvinte uzuale în limbajul de zi cu zi, pentru a putea fi înțelese și pronunțate ușor și fără ezitări;
- frazele ar trebui să conectate semantic pentru ca vorbitorul să înțeleagă cât mai repede sensul frazei și astfel să poată obține o fluentă în pronunție (trebuie evitate schimbările de subiect);

- frazele nu ar trebui să conțină cuvinte împrumutate din limbi străine (de obicei ambigue din punct de vedere al pronunției) sau interjecții (de obicei pronunțate foarte diferit de fiecare vorbitor în parte);
- frazele ar trebui să conțină toate semnele de punctuație (pe care vorbitorul le va utiliza pentru a citi cu intonație);
- frazele nu ar trebui să conțină numere scrise cu cifre (care pot fi ambigue din punct de vedere al pronunției).

În cazul acestei prime metode de achiziție cele mai importante avantaje sunt următoarele:

- proiectanții dețin controlul asupra materialelor înregistrate, asupra calității înregistrărilor și asupra variabilității vorbitorilor;
- achiziția se poate face într-un timp relativ scurt, fără ca cei implicați în achiziție să fie antrenați special în acest sens;
- erorile de pronunție sunt de regulă rare, astfel că baza de date achiziționată în acest mod poate fi presupusă ca fiind corectă.

Totuși, un dezavantaj important este faptul că, utilizând această metodă de achiziție, nu se pot obține baze de date de vorbire spontană, conversațională. Vorbitorul știe a priori textul pe care îl va pronunța și este mult mai concentrat să îl pronunțe clar și corect, fără ezitățile, bâlbele sau lipsa de fluiditate caracteristice vorbirii conversaționale.

2.4.2 Etichetarea unor materiale audio ce conțin vorbire

Etichetarea unor materiale audio ce conțin vorbire implică obținerea materialelor audio, aducerea lor la un format comun (aceeași frecvență de eșantionare, aceeași dimensiune a eșantioanelor, etc.), selectarea părților ce conțin vorbire, fragmentarea materialelor audio și, în final, etichetarea lor.

Avantajul cel mai important al acestei metode de achiziție este posibilitatea obținerii unei baze de date de vorbire spontană, în care vorbesc diverși vorbitori, în diverse contexte, având diverse stări de spirit, emoții, etc. Prețul plătit pentru a obține o astfel de bază de date de vorbire este unul destul de mare, el fiind direct influențat de dificultatea procesului de achiziție, procesare și etichetare pentru aceste clipuri audio. Timpul necesar etichetării unei ore de vorbire spontană este de cel puțin 3-4 ori mai mare decât timpul necesar înregistrării unei ore de vorbire. Mai mult, etichetarea nu poate fi făcută fără un antrenament prealabil, iar erorile de etichetare sunt mult mai frecvente decât erorile de pronunție.

În concluzie, în cazul în care se dorește recunoașterea vorbirii continue citite cea mai potrivită metodă de achiziție este înregistrarea textelor predefinite. În cazul în care se dorește recunoașterea vorbirii continue spontane cel mai potrivit ar fi un proces de antrenare

ierarhizat în care să se utilizeze inițial o bază de date achiziționată folosind prima metodă, iar ulterior o bază de date de vorbire spontană.

2.5 Corpusul de text

Rolul și modul de utilizare al unui corpus de text într-un sistem de recunoaștere automată a vorbirii este evidențiat în Figura 2.1. Se poate observa că această resursă lingvistică este necesară în partea de dezvoltare a unui sistem de recunoaștere, mai precis la antrenarea modelului lingvistic.

Sistemele de RAV de comandă și control nu necesită un corpus de text pentru antrenarea modelului lingvistic pentru că în aceste sisteme se folosesc de obicei gramatici cu reguli. În cazul sistemelor de RVC cu vocabular extins se folosesc modele de limbă statistice, care se construiesc utilizând corpusuri de text de dimensiuni cât mai mari și cât mai adaptate domeniului din care fac parte mesajele vorbite ce trebuie decodate. Capacitatea de predicție a modelului lingvistic depinde de dimensiunea corpusului de text utilizat la antrenare, pentru că cu cât numărul cuvintelor și secvențelor de cuvinte dintr-o limbă este mai mare, cu atât probabilitățile de apariție ale acestora sunt mai dificil de estimat cu acuratețe. De aceea, pentru un sistem de RAV de bună calitate este absolut necesară achiziția unui corpus de text de dimensiuni cât mai mari (milioane-miliarde de cuvinte).

Achiziția unui corpus de text de asemenea dimensiuni nu poate fi făcută decât în mod automat, iar o cantitate atât de mare de text nu poate fi găsită decât pe Internet. În consecință, cele mai populare metode de dezvoltare a corpusurilor de text folosesc principiul Web-as-Corpus și implică identificarea, descărcarea și procesarea materialelor text de pe Internet.

Identificarea surselor de text disponibile pe Internet este relativ simplă (site-uri de știri online, blog-uri, subtitrări, etc.), însă dificultatea proceselor de descărcare, normalizare și restaurare a textelor depinde mult de sursa de text. De exemplu, descărcarea unui număr atât de mare de articole de știri trebuie făcută automat, însă nu toate ziarele online pun la dispoziție link-uri ușor de procesat în mod automat pentru articolele lor.

2.5.1 Normalizarea și restaurarea corpusului de text

Normalizarea și restaurarea corpusului de text presupune procesarea acestuia în vederea obținerii unui corpus de text util în antrenarea modelului lingvistic al unui sistem de RVC. Un sistem de RVC are rolul de a transcrie vorbirea, deci va scoate la ieșire text simplu ce conține numai cuvinte și simboluri care au un corespondent acustic (fără litere mari, fără simboluri speciale, fără cifre, fără semne de punctuație, etc.). De aceea modelul de limbă trebuie să modeleze numai probabilitățile de apariție ale *cuvintelor și succesiunilor de cuvinte pronunțabile* (nu ale altor caractere sau simboluri).

În funcție de sursa de achiziție a textului, normalizarea și restaurarea lor poate presupune mai multe sau mai puține operații, însă, de regulă, următoarele sunt indispensabile:

- eliminarea tuturor etichetelor HTML (în cazul în care textele au fost achiziționate de pe Internet);
- normalizarea caracterelor diacritice cu care este scris textul (înlocuind toate versiunile de diacritice cu un singur simbol pentru fiecare caracter diacritic în parte). Această operație se referă la textele scrise cu ș cu sedilă, respectiv ț cu sedilă;
- expandarea abrevierilor (cuvintele abreviate nu se pronunță așa cum se scriu, ci în forma lor neabreviată);
- transcrierea numerelor scrise cu cifre în numere scrise cu litere (în cazul numărului 1984, semnalul acustic nu are nimic de-a face cu cifrele propriu zise 1, 9, 8, 4, ci este pur și simplu pronunția secvenței de cuvinte o mie nouă sute opt zeci și patru);
- eliminarea sau înlocuirea semnelor de punctuație și a caracterelor speciale (de exemplu a) virgulele trebuie șterse, b) punctele, semnele de întrebare și semnele de exclamare trebuie înlocuite cu caracterul de linie nouă, c) parantezele trebuie șterse, iar textul din interiorul lor poate fi plasat pe o linie separată, d) liniuțele, cu excepția celor din interiorul cuvintelor trebuie șterse, e) caracterele procent pot fi înlocuite cu expresia la sută, etc.);
- transformarea tuturor literelor mari în litere mici.

O operație ceva mai dificilă, dar care trebuie abordată pentru multe dintre textele achiziționate de pe Internet este restaurarea diacriticelor. Cuvintele formate din aceleași litere, dar care au semne diacritice diferite (deci cuvintele ambigue din punctul de vedere al diacriticelor) au de cele mai multe ori sensuri diferite, se pronunță diferit și apar în contexte diferite. De aceea, ele nu pot fi modelate, nici din punct de vedere lingvistic, nici din punct de vedere acustic, neținând cont de diacritice. În concluzie, dacă textul achiziționat și normalizat nu conține diacritice ele trebuie neapărat restaurate pentru ca textul să devină util în antrenarea sistemului de RVC.

2.6 Concluzii

Sistemele moderne de recunoaștere a vorbirii continue utilizează modele statistice pentru a modela probabilitățile de apariție ale cuvintelor și succesiunilor de cuvinte dintr-o limbă și pronunția fonemelor. Pentru antrenarea modelelor statistice este nevoie de cantități mari de date reprezentative pentru fenomenul ce trebuie modelat. În cazul sistemelor de RVC este necesare baze de date de vorbire etichetată pentru modelarea pronunției fonemelor și corpusuri de text pentru modelarea statisticii apariției cuvintelor. Pe lângă aceste două resurse este nevoie și de un dicționar fonetic care să specifice modul de fonetizare al cuvintelor limbii.

3. PRINCIPIILE DE BAZĂ ALE MODELĂRII ACUSTICE

În platforma intitulată “Arhitectura generală a unui sistem de recunoaștere a vorbirii” am menționat faptul că sistemele de recunoaștere a vorbirii continue (RVC) state-of-the-art nu estimează direct probabilitatea mesajului vorbit, dată fiind o succesiune de cuvinte $p(X|W)$, ci estimează probabilitatea unor unități de vorbire mai mici, de obicei foneme. Prin urmare, modelul acustic (MA) este format dintr-un set de modele pentru foneme care sunt conectate, în timpul procesului de decodare, pentru a forma modele pentru cuvinte și apoi modele pentru succesiuni de cuvinte. Acestea sunt utilizate în final pentru a estima $p(X|W)$. S-a dovedit că această abordare generativă poate fi foarte bine implementată utilizând modele Markov ascunse (hidden Markov models - HMM) [Baker, 1975; Poritz, 1988; Rabiner, 1989; Jelinek, 1998].

HMM-urile sunt automate probabiliste cu număr finit de stări care pot fi combinate în mod ierarhic pentru a construi modele pentru secvențe de cuvinte din modele pentru unități de vorbire mai scurte. În sistemele de RVC cu vocabular extins modelele pentru secvențe de cuvinte sunt construite din modele pentru cuvinte, care, la rândul lor, sunt construite din modele pentru unități sub-lexicale (de obicei foneme dependente de context) utilizând un dicționar fonetic.

3.1 Parametri acustici

HMM-urile nu modelează semnalul acustic pe baza formei de undă a acestuia. Așa cum ilustrează Figura 1.1, un bloc de extragere de parametri este folosit pentru a calcula niște vectori de coeficienți care sunt apoi modelați de către modelul acustic.

Semnalul de vorbire este un semnal nestăționar. În consecință analiza spectrală nu poate fi făcută pe întreg semnalul, ci numai pe cadre scurte (20ms – 30ms), cvasi-stăționare. Semnalul inițial în domeniul timp este transformat prin ferestruire într-o secvență de ferestre în domeniul timp. Din fiecare dintre aceste ferestre sunt extrași în continuare parametri de tip spectral sau cesptral. Mai multe tipuri de parametri, ce pot fi utilizați pentru modelarea vorbirii, au fost propuși încă din anii '80. Cei mai folosiți parametri sunt cei de tip cepstral-perceptual, dintre care amintim *parametri Mel-cepstrali* (*Mel-Frequency Cepstrum*

Coefficients - MFCCs) și *parametri percepțuali obținuți prin predicție liniară (Perceptual Linear Prediction - PLP)*. Un avantaj particular al reprezentării cepstrale comparativ cu reprezentarea spectrală este decorelația coeficienților, comparativ cu gradul mare de corelație observat între coeficienții spectrali vecini. Extragerea parametrilor MFCC, respectiv PLP este detaliată în Figura 3.1 și Figura 3.2.

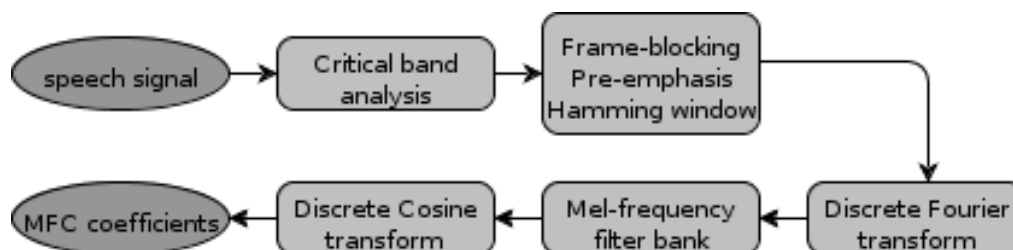


Figura 3.1. Extragerea parametrilor MFCC

Cu toate că fiecare set de coeficienți, fie ei MFCC sau PLP, este calculat pe o fereastră scurtă din semnalul de vorbire, informația conținută de dinamica temporală a acestor parametri este și ea utilă în RAV. Mai mult, energia totală a parametrilor și derivatele ei temporale par să fie utile în RAV. De aceea, cel mai utilizat set de parametri în RAV este format din 39 de coeficienți: 12 MFCC + energie, împreună cu derivatele temporale de ordin unu și doi.

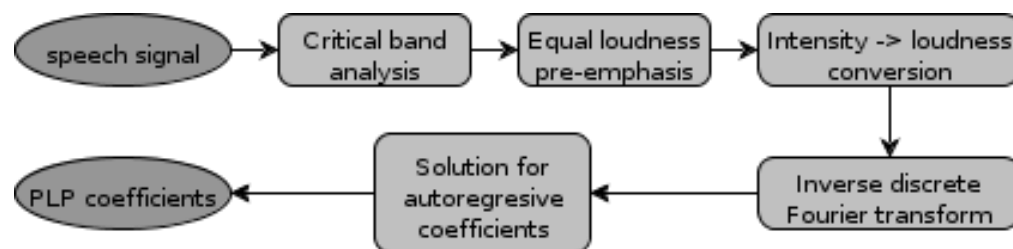


Figura 3.2. Extragerea parametrilor PLP

3.2 Platforma HMM/GMM

Un HMM este un automat cu stări finite format dintr-un set de stări conectate cu arce reprezentând tranziții. Într-un HMM secvența de stări este ascunsă, de aceea întregul proces Markov este considerat ascuns: secvența de stări nu este direct disponibilă observatorului. Observatorul are acces numai la secvența de vectori de parametri acustici generată de o funcție de densitate de probabilitate atașată fiecărei stări. Acest tip de proces Markov s-a dovedit a fi un foarte bun model al producerii vorbirii. O reprezentare detaliată a unui HMM este ilustrată în Figura 3.3. Așa cum se vede în figură, un HMM este caracterizat de următorii parametri:

- un set de stări $Q=q_1q_2...q_N$;

- un set de probabilități $A = a_{11}a_{12}...a_{NN}$ reprezentând probabilitățile de tranziție între stări. Fiecare $a_{ij} = p(q_j | q_i)$ reprezintă probabilitatea de tranziție din starea i în starea j ;
- un set de observații probabilistice $B = b_i(x_t) = p(x_t | q_i)$, fiecare exprimând probabilitatea ca observația x_t să fi fost generată de starea i .

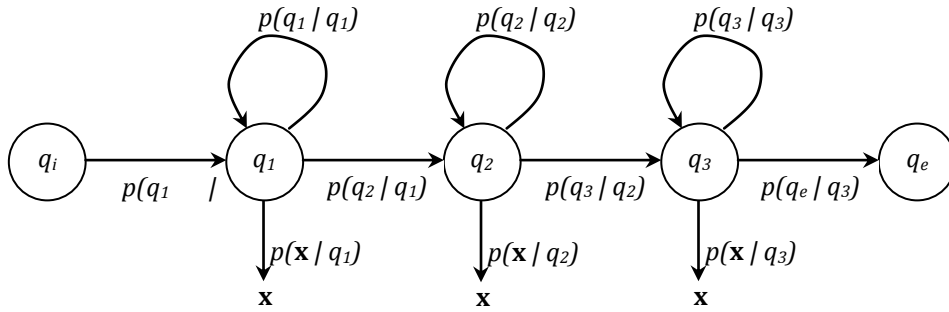


Figura 3.3. Reprezentarea unui HMM ca un automat probabilist cu număr finit de stări

Cu toate că definiția unui HMM permite tranziția din orice stare în orice altă stare, în RAV modelele sunt create astfel încât să interzică tranzițiile arbitrare. Asta pentru că este foarte important și util ca modelul să țină cont de caracterul secvențial al producerii vorbirii. De aceea modelul impune constrângeri suplimentare privind tranzițiile înapoi sau salturile peste stări.

Probabilitățile parametrilor generați de o stare q_i sunt de obicei modelate de o funcție densitate de probabilitate. De cele mai multe ori aceasta ia forma unei sume de K funcții Gaussiene (Gaussian mixture model - GMM) în d dimensiuni:

$$b_i(x) = p(x | q_i) = \sum_{k=1}^K c_{ik} N(x; \mu_{ik}, \Sigma_{ik}) \quad (3.1)$$

Modelarea vorbirii cu ajutorul HMM-urilor face două presupuneri importante [Renalds, 2010]:

- proces Markov: se presupune că tranziția între stările unui HMM este un proces Markov de ordin întâi în care probabilitatea tranziției următoare depinde numai de starea curentă.
- independența observațiilor: probabilitatea ca o stare să genereze un anumit vector de parametri este independentă de vectorii de parametri generați în prealabil.

Aceste două presupuneri pot conduce la un model al producerii vorbirii nerealist, însă ele sunt necesare pentru simplificările matematice și computaționale pe care le aduc. Problema estimării și decodării unui HMM nu pot fi abordate sau pot fi abordate într-un mod foarte complicat dacă se renunță la aceste două condiții simplificatoare.

Pentru rezolvarea celor trei probleme ale platformei HMM/GMM (evaluarea, decodarea și estimarea) există algoritmi standard ce utilizează cele două simplificări

menționate mai sus. Evaluarea (estimarea probabilității ca o secvență de observații să fi fost generată de un HMM) este rezolvată recursiv de *algoritmul Forward*. Problema decodării HMM-urilor, adică găsirea celei mai probabile secvențe stări care să fi generat o secvență de observații, este rezolvată de o variantă a *algoritmului Viterbi*. Diverșii parametri ai unui sistem HMM/GMM sunt estimați folosind *algoritmul Forward-Backward* sau *Baum-Welch* [Baum, 1972], o particularizare a algoritmului *expectation-maximization* (EM) [Dempster, 1977].

3.3 Selecția, dependența de context și clusterizarea unităților de vorbire

Am arătat deja că sistemele de RVC modelează unități de vorbire sub-lexicale folosind HMM-uri. Totuși, mai sunt niște aspecte ce privesc alegerea unităților de vorbire și care trebuie discutate. Autorii [Huang, 2001] văd trei caracteristici pe care o unitate de vorbire adecvată trebuie să le aibă:

- *precizie*; unitatea trebuie să reprezinte realizarea acustică care apare în diferite contexte.
- *antrenabilitate*; trebuie să existe destule date pentru a estima parametri respectivei unități.
- *generalizabilitate*; orice nou cuvânt trebuie să poată fi construit având la dispoziție un set de unități de bază (aceasta ar asigura independența de sarcină a sistemului de RAV).

Date fiind acestea, este acum foarte evident de ce *cuvintele* nu pot fi folosite ca unități de vorbire în sistemele de RVC: ele nu sunt nici antrenabile, nici generalizabile.

Ca o alternativă viabilă s-ar putea utiliza *fonemele* ca unități de vorbire de bază. Spre diferență de modelele pentru cuvinte, modelele pentru foneme nu prezintă probleme de antrenare pentru că există suficiente ocurențe pentru fiecare fonem chiar și într-o bază de date de câteva mii de fraze. Mai mult, fonemele sunt, prin natura lor, independente de vocabularul sarcinii de recunoaștere. În consecință fonemele sunt preferate în detrimentul cuvintelor pentru că ele sunt mai antrenabile și generalizabile. Totuși, modelul fonetic nu este tocmai adecvat pentru că el presupune că un fonem este pronunțat identic în orice context. Chiar dacă am încerca să pronunțăm fiecare cuvânt ca o succesiune de foneme independente, aceste foneme nu sunt produse independent pentru că sistemul articulator nu se poate mișca instantaneu dintr-o poziție în alta. Deci producerea unui fonem este puternic influențată de producerea fonemelor din imediata sa vecinătate. Dacă modelele pentru cuvinte nu sunt caracterizate de generalizabilitate, modelele pentru foneme generalizează prea mult și asta conduce la o scădere semnificativă de precizie.

Utilizarea modelelor fonetice dependente de context ar crește foarte mult precizia recunoașterii, în condițiile în care ar exista suficiente date de antrenare pentru a estima corect parametri acestor modele pentru unitățile de vorbire de bază (dependente de context). Fonele dependente de context au fost utilizate pe scară largă în sistemele de RAV cu vocabular extins în primul rând datorită preciziei și antrenabilității crescute. Contextul se referă de obicei la fonelele din imediata vecinătate (stânga și/sau dreapta) a fonemului curent.

Un model de tip *trifonem* este un model fonetic care ia în considerare ambele fonele vecine (stânga/dreapta). Dacă două fonele identice se regăsesc în mesajul vorbit în contexte diferite, ele sunt considerate trifoneme diferite. Diferitele realizări particulare ale unui fonem sunt denumite alofone. Trifonemele sunt exemple de alofone.

Modelele pentru trifoneme sunt mai puternice pentru că ele capturează cele mai importante efecte de co-articulare. Cu toate acestea, din cauza numărului ridicat de modele diferite, ce conduce la un număr mare de parametri ce trebuie estimați, antrenabilitatea lor devine o problemă delicată. Mai multe tehnici de partajare a parametrilor (legarea stărilor, clusterizarea modelelor, etc.) au fost utilizate pentru a echilibra balanța între precizia și antrenabilitatea modelelor. Hwang și Huang au introdus clusterizarea la nivel de stare [Hwang, 1991]. Fiecare cluster reprezintă un set de stări Markov similare denumit *senone* [Hwang, 1993]. În momentul de față senonele sunt cele mai utilizate unități de vorbire în sistemele de RAV.

4. PRINCIPIILE DE BAZĂ ALE MODELĂRII LIMBAJULUI NATURAL

În Capitolul 0 am menționat faptul că sistemele de recunoaștere a vorbirii continue cu vocabular extins (RVC-VE) utilizează în timpul decodării un model de limbă (ML) pentru a estima probabilitățile tuturor secvențelor de cuvinte din spațiul de căutare. În general, rolul unui model de limbă este de a estima probabilitatea ca o secvență de cuvinte $W = w_1, w_2, \dots, w_n$, să fie o secvență validă pentru sarcina de recunoaștere respectivă. Probabilitățile acestor secvențe de cuvinte ajută foarte mult modelul acustic în procesul de decizie. De exemplu, în cazul unei sarcini de recunoaștere de vorbire continuă (fără restricții) aceste două fraze: *casa ta e mare?* și *casat a Ema re?* sunt similare din punct de vedere acustic, însă cea de-a doua nu are niciun sens. Rolul modelului de limbă este acela de a asigura o probabilitate mult mai mare primei secvențe de cuvinte și, prin urmare, de a ajuta sistemul de RAV să decidă în favoarea acesteia.

Sistemele de recunoaștere de vorbire continuă cu vocabular mare de ultimă generație utilizează modele de limbă statistice de tip n-gram. Aceste modele de limbă se construiesc pe baza unor corpusuri mari de text specifice sarcinii de recunoaștere, estimând probabilitățile de apariție ale cuvintelor și secvențelor de cuvinte specifice respectivei sarcini. Modelele de limbă de tip n-gram se utilizează apoi în procesul de decodare (recunoaștere a vorbirii) pentru a calcula probabilitățile secvențelor de cuvinte propuse de modelul acustic. Modelele de limbă statistice vor fi detaliate în următoarea secțiune.

Pentru sarcinile de recunoaștere de tip comandă și control modelele de limbă statistice nu sunt foarte adecvate. În cazul în care vocabularul sarcinii de recunoaștere conține doar câteva zeci sau sute de cuvinte, iar modul în care acestea se pot succede este bine precizat, modelele de limbă de tip gramatică cu stări finite (FSG – finite state grammar) sunt mult mai potrivite. O gramatică cu stări finite este un model de tip graf în care nodurile reprezintă cuvinte ale vocabularului sarcinii de recunoaștere, iar tranzițiile între cuvinte sunt reprezentate de arcele grafului. Un astfel de model de limbă specifică în mod explicit toate secvențele de cuvinte permise de gramatica sarcinii de recunoaștere. Mai mult, fiecărui arc în parte îi poate fi asignat un cost, specificând astfel probabilitatea ca un cuvânt să fie precedat de un altul (sau altfel spus probabilitatea respectivei secvențe de două cuvinte).

4.1 Modele de limbă statistice

Așa cum spuneam, rolul unui model de limbă este de a estima probabilitatea ca o secvență de cuvinte $W = w_1, w_2, \dots, w_n$, să fie o secvență validă pentru sarcina de recunoaștere respectivă. Probabilitatea secvenței de cuvinte $W = w_1, w_2, \dots, w_n$ poate fi descompusă astfel:

$$p(W) = p(w_1, w_2, \dots, w_n) = p(w_1)p(w_2 | w_1) \dots p(w_n | w_1, w_2, \dots, w_{n-1}) \quad (4.1)$$

Aceasta înseamnă că problema estimării probabilității secvenței de cuvinte W este împărțită în mai multe probleme de estimare a probabilității unui singur cuvânt, dată fiind secvența de cuvinte anterioare lui (*istoria cuvântului*). Din motive de eficiență computațională, istoriile de cuvinte precedente nu pot include un număr indefinit de cuvinte, ci sunt limitate la ultimile m cuvinte. Altfel spus, se face presupunerea că doar un număr limitat de cuvinte afectează probabilitatea următorului cuvânt. Aceasta conduce la modelul de limbă convențional de tip *n-gram*, model ce este state-of-the-art în RVC-VE de peste 20 de ani [Renalds, 2010]. În mod tipic m este ales în funcție de cantitatea de text de antrenare disponibil (este nevoie de mai mult text pentru a putea estima cu o precizie bună probabilitățile pentru istorii mai lungi). De obicei se folosesc modele de limbă de tip trigram. Acestea iau în considerare numai ultimile două cuvinte pentru a-l prezice pe al treilea. Pentru aceasta este nevoie de colectarea unor statistici de apariție pentru secvențe de trei cuvinte, așa numitele 3-gramme (trigrame). Se pot construi modele de limbă și pentru secvențe de două cuvinte (bigrame), sau pentru secvențe formate din mai multe cuvinte (de ordin mai mare).

4.1.1 Construcția modelelor de tip n-gram și alte probleme specifice

Un model de limbă de tip *n-gram* se construiește estimând probabilitățile discutate mai sus pe baza unor corpusuri de text de mari dimensiuni. De exemplu, în cazul unui ML bigram trebuie estimate probabilitățile $p(w_j | w_i)$ pentru fiecare pereche de cuvinte (w_i, w_j) . Pentru a calcula aceste probabilități se utilizează principiul *maximum likelihood*. Se numără de câte ori cuvântul w_i este urmat de cuvântul w_j , comparativ cu alte cuvinte:

$$p(w_j | w_i) = \frac{\text{count}(w_i, w_j)}{\sum_w \text{count}(w_i, w)} \quad (4.2)$$

Pentru a estima corect aceste probabilități este nevoie de o cantitate mare de date de antrenare (tipic de ordinul a câteva zeci de milioane de cuvinte). Pentru construcția modelelor *n-gram* de ordin mai mare este de nevoie de și mai multe date.

O problemă cheie în construcția ML de tip *n-gram*, chiar și atunci când se dispune de corpusuri mari, este *data sparseness*. Indiferent cât de mare este corpusul de antrenare, vor fi *n-gramme* care nu vor apărea în acest text, însă care ar putea să apară în textul de evaluare. Conform Ecuației 4.2, în acest caz limită, probabilitatea asignată *n-gramelor* necunoscute este

0. În afară de acest caz există alte n -grame care apar de foarte puține ori (mai puțin de zece ori) în corpusul de antrenare. Această problemă devine mai importantă în cazul ML de tip n -gram de ordin mai mare. În toate aceste cazuri, probabilitățile care au fost estimate pe baza numărului de apariții ale n -gramelor în corpusul de antrenare, trebuie ajustate.

Metodele care se ocupă de acest proces de ajustare se numesc *metode de netezire*. Ele extrag o parte din probabilitatea alocată pentru n -gramele întâlnite la antrenare și o redistribuie n -gramelor necunoscute. Există mai multe metode de netezire care particularizează modul de redistribuție a probabilității. Dintre acestea, cea mai eficientă metodă este *Good-Turing*, numită și *netezirea Katz*.

Problema data sparseness este abordată într-o manieră oarecum diferită de către așa numitele *metode de back-off*. Aceste metode utilizează mai multe modele de limbă, care au avantaje diferite, pentru a crea un model de limbă interpolat, care să poată beneficia de toate părțile constituente. De exemplu, modelele de tip n -gram de ordin mai mare oferă un context de predictibilitate mai mare, însă modelele de ordin mai mic sunt mai robuste. Mai multe metode de back-off, ce încearcă să estimeze probabilitățile n -gramelor necunoscute pe baza probabilităților asiguate acestor n -grame de către modele de ordin mai mic, au fost propuse în ultimile două decenii. Dintre acestea, metoda *Kneser-Ney modificată* [Chen, 1998] pare să fie cea mai eficientă la ora actuală.

4.1.2 Criterii de performanță utilizate în evaluarea modelelor de limbă

În cazul în care un sistem de RAV este disponibil, cel mai utilizat criteriu de performanță pentru ML este *rata de eroare la nivel de cuvânt* (word error rate – WER). Alternativ, fără a implica un sistem de RAV, se poate evalua puterea de predicție a unui ML măsurând probabilitatea pe care el o asignează unor secvențe de cuvinte de evaluare. Un model de limbă bun ar trebui să atribuie o probabilitate mare unui text corect și o probabilitate mică unui text prost. În acest caz, cel mai comun criteriu de performanță este *perplexitatea*. O perplexitate mai mare pe o secvență de cuvinte particulară denotă o capacitate de predicție a secvenței mai mică pentru respectivul ML. În general, în contextul RAV, perplexitatea se corelează foarte bine cu WER-ul [Huang, 2001].

Este posibil ca unele dintre cuvintele din textul de evaluare să nu fi fost întâlnite deloc în corpusul de antrenare. Aceste cuvinte sunt numite *cuvinte în afara vocabularului* (out of vocabulary - OOV) și nu pot fi prezise de modelul de limbă. Cuvintele OOV fac evaluarea unui ML mai dificilă: din cauză că perplexitatea lor este infinită, aceasta nu poate fi adunată la perplexitățile celorlalte n -grame pentru a obține perplexitatea întregii secvențe de cuvinte. Astfel, pe lângă perplexitate, procentul de cuvinte OOV trebuie să fie și el specificat pentru ca evaluarea ML să fie completă.

4.2 Gramatici cu stări finite (FSG – finite state grammar)

Recunoașterea vorbirii continue este una dintre cele mai dificile sarcini de recunoaștere pentru că vorbitorul poate pronunța orice succesiune de cuvinte care are un înțeles în limba respectivă. De aceea, vocabularul acestei sarcini de recunoaștere ar trebui să conțină practic toate cuvintele limbii, iar modelul de limbă ar trebui să fie capabil să estimeze corect probabilitățile de apariție pentru fiecare succesiune de cuvinte posibilă. În aceste condiții modelele de limbă statistice de tip n -gram sunt foarte potrivite. În condițiile în care sarcina de recunoaștere prezintă restricții cu privire la vocabularul de cuvinte și la setul de succesiuni valide de cuvinte, atunci modelul n -gram nu mai este atât de potrivit, întrucât efortul de antrenare nu se justifică. Succesiunile valide de cuvinte și probabilitățile lor de apariție pot fi specificate direct folosind un model de limbă de tip gramatică cu stări finite (FSG). O gramatică cu stări finite este un model de tip graf în care nodurile reprezintă cuvinte ale vocabularului sarcinii de recunoaștere, iar tranzițiile între cuvinte sunt reprezentate de arcele grafului. Un astfel de model de limbă specifică în mod explicit toate secvențele de cuvinte permise de gramatica sarcinii de recunoaștere. Mai mult, fiecărui arc în parte îi poate fi asignat un cost, specificând astfel probabilitatea ca un cuvânt să fie precedat de un altul (sau altfel spus probabilitatea respectivei secvențe de două cuvinte).

4.2.1 Proiectarea unei gramatici cu stări finite

Pentru a ilustra mai bine conceptele descrise mai sus vom porni de la un exemplu simplu de construcție a unei gramatici cu stări finite, specifică unei anumite sarcini de recunoaștere: transcrierea secvențelor audio ce conțin cifre. Această sarcină de recunoaștere presupune că vorbitorul poate rosti numai cuvintele *zero*, *unu*, *doi*, *trei*, *patru*, *cinci*, *șase*, *șapte*, *opt* și *nouă*, în orice ordine și de oricâte ori. Definirea grafului ce va reprezenta gramatica cu stări finite începe cu definirea nodurilor de intrare și ieșire din gramatică. Aceste noduri vor fi notate cu N (de la *null*) pentru că trecerea prin aceste noduri nu corespunde afișării vreunui cuvânt. Aceste noduri fac parte din infrastructura gramaticii cu stări finite. Urmează crearea nodurilor specifice fiecărui cuvânt în parte și legarea lor de nodurile de intrare, respectiv ieșire utilizând arce direcționate (dinspre nodul de intrare către fiecare cuvânt și dinspre fiecare cuvânt către nodul de ieșire). În acest moment am definit o gramatică a cărei parcurgere poate fi făcută dacă se pronunță o singură cifră. Pentru sarcina de recunoaștere propusă (recunoașterea unei secvențe de cifre) mai avem de făcut un singur lucru: adăugarea unei tranziții înapoi (dinspre nodul de ieșire spre nodul de intrare). În final trebuie să mai adăugăm încă două noduri *null* pentru că cele precedente nu mai sunt noduri de intrare/ieșire propriu-zise (au și arce care intră în ele și arce care ies din ele).

În Figura 4.1 este reprezentată grafic gramatica cu stări finite a sarcinii noastre de recunoaștere. Se poate observa că modelul este format din 14 noduri, dintre care numai zece reprezintă cuvinte ale limbii (cifrele de la zero la nouă), iar celelalte patru sunt noduri de infrastructură, utilizate pentru a realiza „intrarea”, respectiv „ieșirea” din graf și pentru

tranziția înapoi. Tranzițiile arată modul în care se poate parcurge această gramatică și implicit secvențele de cuvinte permise: în fiecare clip audio se pot pronunța una sau mai multe cifre.

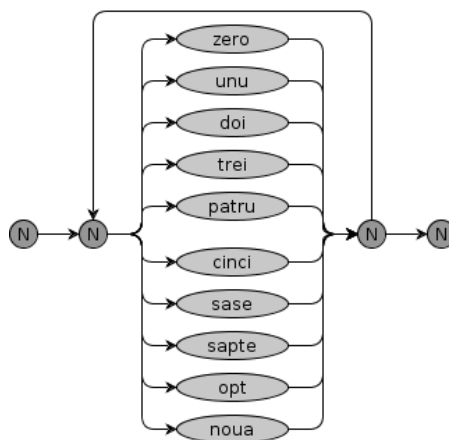


Figura 4.1. Gramatica cu stări finite a sarcinii de recunoaștere de cifre

Implementarea unei astfel de gramatici FSG poate fi făcută foarte simplu utilizând formatul Java Speech Grammar (JSGF) așa cum vom vedea în continuare.

4.3 JavaSpeech Grammar Format (JSGF)

Formatul de gramatică JavaSpeech definește o modalitate de descriere a gramaticilor de reguli (gramatici de comandă și control) independentă de platformă și de furnizor. Acest format folosește o reprezentare textuală a regulilor ce poate fi ușor citită și modificată atât de manual, cât și programatic.

O gramatică cu reguli specifică tipurile de fraze vorbite pe care un vorbitor le-ar putea rosti (o frază vorbită, în această accepțiune, este similară cu o frază scrisă). De exemplu, o gramatică simplă de control al unei ferestre dintr-un program ar putea să “asculte” comenzi precum: *deschide un fișier*, *închide fereastra*, etc.

Tipologia frazelor pe care utilizatorul le-ar putea rosti este foarte puternic dependentă de context: utilizatorul folosește o aplicație de email, formează un număr de telefon sau selectează un font? Aplicațiile cunosc contextul și, în consecință, pot furniza gramatica cea mai potrivită sistemului de recunoaștere a vorbirii.

În continuare vor fi prezentate câteva convenții utilizate pentru denumiri și apoi componentele de bază ale gramaticii (antetul și corpul gramaticii). În antetul gramaticii se specifică numele gramaticii și se listează alte reguli și gramatici importate. În corpul gramaticii se declară regulile ce formează gramatica. Aceste reguli sunt o combinație de text ce poate fi pronunțat și referințe către alte reguli. În final vom prezenta câteva exemple de definiții de gramatici.

4.3.1 Definiții

Fiecare gramatică definită prin JSGF are un nume unic declarat în antetul gramaticii. Numele gramaticii poate fi un nume complet (ce conține și denumirea pachetului din care face parte gramatica) sau un nume simplu. Denumirea pachetelor de gramatici și a gramaticilor se face similar cu denumirea pachetelor și claselor din limbajul de programare Java.

Orice gramatică este compusă dintr-un set de reguli care definesc complet tipologia frazelor ce pot fi pronunțate. Regulile sunt combinații de text (ce poate fi pronunțat) și referințe către alte reguli. Fiecare regulă are o denumire unică. O referință către o altă regulă se realizează folosind denumirea regulii respective, precedată de caracterul < și urmată de caracterul >. Numele regulilor sunt similare cu numele identificatorilor din limbajul de programare Java, respectând aceleași restricții și convenții.

Simbolurile sau terminalii definesc cuvinte sau secvențe de cuvinte care pot fi pronunțate de către vorbitor. De cele mai multe ori un terminal este asociat unui singur cuvânt, însă aceasta nu reprezintă o restricție (un terminal poate fi asociat și unui grup de cuvinte). Un terminal poate apărea izolat, ca în următoarele exemple:

salutare
bucurești

sau într-o secvență de terminali ca în exemplele de mai jos:

bună ziua
douăzeci și patru

Un terminal este de fapt o secvență de caractere delimitată de spații, ghilimele sau alte caractere care au o semnificație specială în JSGF: ; = | * + < > () [] {} /* */ //. Un terminal este în același timp o referință către un cuvânt din dicționarul fonetic al sistemului de recunoaștere automată a vorbirii. Utilizând pronunția terminalului din dicționarul fonetic, sistemul de recunoaștere automată a vorbirii poate transcrie vorbirea în text.

Un terminal nu trebuie neapărat să corespundă unui singur cuvânt din vocabular, ci poate fi o secvență de cuvinte sau un simbol. Pentru a defini un terminal multi-cuvânt sau un terminal de tip simbol trebuie folosite ghilimelele ca în exemplele de mai jos:

primăria municipiului "Baia de Aramă"
doi "+" trei

Un terminal multi-cuvânt este util când pronunția cuvintelor componente diferă în funcție de context. În acest fel se pot defini intrări separate în vocabularul sistemului de RAV, fiecare cu pronunția ei. Terminalii multi-cuvânt pot fi folosiți și pentru a simplifica procesarea rezultatelor, de exemplu obținerea unui singur terminal pentru expresii precum "București", "Baia Mare", "Cluj-Napoca".

Terminalii multi-cuvânt pot fi incluși în dicționarul fonetic al sistemului de RAV, exact ca orice alt terminal. În cazul în care terminalul multi-cuvânt este găsit în dicționarul fonetic, atunci se va folosi pentru recunoaștere pronunția respectivă. În caz contrar, se vor folosi pronunțiile pentru toți sub-terminalii acestuia (toate simbolurile separate de spații din terminalul multi-cuvânt).

Majoritatea sistemelor de recunoaștere de vorbire au capacitatea de a folosi și semne de punctuație sau diferite simboluri. De exemplu, sistemele de recunoaștere pentru limba română trebuie să poată gestiona cratima (într-o, dă-mi). Totuși, există multe situații care nu pot fi tratate corect și neambiguu de un sistem de RAV. În aceste cazuri, persoana care dezvoltă gramatica ar trebui să utilizeze terminali cât mai apropiați de modul în care se pronunță cuvintele respective și care apar în dicționarul fonetic. Următoarele exemple prezintă câteva cazuri uzuale:

- Numerele: terminalii scriși cu cifre “0 1 2” trebuie scriși cu litere “zero unu doi”. În mod similar o comandă de forma “formează 0740 213 449” ar trebui scrisă “formează zero șapte patru zero doi unu trei patru patru nouă”,
- Datele ce conțin cifre trebuie expandate. De exemplu data “25 aprilie 1984” trebuie scrisă sub forma “douăzeci și cinci aprilie o mie nouă sute opt zeci și patru”,
- Abrevierile și acronimele trebuie de asemenea expandate. De exemplu abrevierile de genul “d-l” sau “dl.” trebuie scrise “domnul”, iar acronimele de forma “SUA” trebuie transformate în “Statele Unite ale Americii”.

4.3.2 Antetul gramaticii

O gramatică de tip JSGF este definită într-un singur fișier. Sau, altfel spus, un fișier poate conține definiția unei singure gramatici. Definiția gramaticii conține două părți: antetul și corpul gramaticii. Antetul începe cu un identificator, definește numele gramaticii și conține eventuale importuri de reguli din alte gramatici. Corpul gramaticii definește regulile componente ale gramaticii, dintre care unele pot fi declarate ca fiind publice.

Orice fișier care definește o gramatică JSGF trebuie să înceapă cu un identificator care indică faptul că fișierul este o definiție JSGF și specifică versiunea formatului folosit. Opțional, identificatorul poate specifica în continuare setul de caractere utilizat în fișier, limba și zona geografică (locale-ul) aferente gramaticii. Identificatorul trebuie să se încheie cu caracterul punct și virgulă urmat de caracterul linie nouă. Formatul identificatorului este următorul:

```
#JSGF version char-encoding locale;
```

Iată mai jos câteva exemple de posibili identificatori:

```
#JSGF V1.0;  
#JSGF V1.0 ISO8859-5;  
#JSGF V1.0 JIS ja;
```

În primul exemplu nu se specifică setul de caractere sau limba. În consecință se vor folosi valorile implicite pentru acești parametrii (setul de caractere și limba implicită din sistemul de operare). În al doilea exemplu se specifică setul de caractere ISO8859-5 (caractere chirilice), iar limba este cea implicită. În ultimul exemplu se specifică atât setul de caractere (JIS – un set de caractere pentru limba japoneză), cât și limba (ja – limba japoneză).

Orice fișier care definește o gramatică JSGF trebuie să înceapă cu caracterul diez (#), iar toate caracterele din acest identifiator trebuie să fie codate ASCII.

După acest identificator despre care am discutat mai sus este obligatoriu să urmeze numele gramaticii. Numele gramaticii poate fi format din numele pachetului din care face parte gramatica urmat de numele gramaticii sau poate fi numai numele gramaticii (în cazul în care pachetul nu este specificat). Astfel, formatul liniei care definește numele gramaticii poate fi:

```
grammar numePachetGramatică.numeGramatică;
```

sau

```
grammar numeGramatică;
```

Iată mai jos câteva exemple de nume de gramaticii:

```
grammar org.etti.speech.apps.names;  
grammar ro.pub.speed.dates;  
grammar numbers;
```

Opțional, antetul gramaticii poate conține și declarații de import de alte reguli sau gramatici. Declarațiile de import trebuie să urmeze după declarația numelui gramaticii, dar înainte de corpul gramaticii (înainte de definiția regulilor). O declarație de import permite utilizarea în gramatica curentă a uneia sau tuturor regulilor dintr-o altă gramatică. Formatul unei astfel de declarații de import poate fi:

```
import <numeRegulă>;
```

sau

```
import <numeGramatică.*>;
```

Iată mai jos câteva exemple de declarații de import:

```
import <org.etti.speech.apps.names.feminine>;  
import <ro.pub.speed.dates.*>;
```

Primul exemplu arată modul în care se poate importa o singură regulă dintr-o gramatică: regula `<feminine>` din gramatica `org.etti.speech.apps.names`. Este obligatoriu ca regula importată să fie declarată publică în gramatica din care face parte.

Utilizarea simbolului asterisc în cel de-al doilea exemplu specifică faptul că toate regulile publice ale gramaticii `ro.pub.speed.dates` vor fi importate în gramatica curentă. Dacă, de exemplu, gramatica `ro.pub.speed.dates` definește trei reguli publice: `<shortDateNames>`, `<dayNames>` și `<fullDateNames>`, atunci toate aceste trei reguli vor putea fi referite local în gramatica curentă.

Observați că datorită faptului că atât numele gramaticii, cât și numele regulii (sau asterisc-ul) sunt necesare, o declarație de import nu va avea niciodată forma:

```
import <numeRegulă>; // declarație ilegală
```

O regulă importată cu o declarație de import poate fi utilizată în trei moduri: specificând numai numele regulii (de exemplu, `<shortDateNames>`), specificând numele gramaticii și numele regulii (de exemplu, `<dates.shortDateNames>`) sau specificând numele pachetului gramaticii, numele gramaticii și numele regulii (de exemplu, `<ro.pub.speed.dates.shortDateNames>`).

În cazul în care o regulă este utilizată întotdeauna specificând numele pachetului gramaticii, numele gramaticii și numele regulii, atunci declarația de import nu mai este necesară:

```
// importarea ro.pub.speed.dates este opțională  
<regulă> = <ro.pub.speed.dates.shortDateNames>;
```

4.3.3 Corpul gramaticii

Definiții de reguli

Corpul unei gramatici JSGF definește reguli. O regulă este definită odată ce definiția ei apare în gramatică, iar ordinea în care apar definițiile regulilor nu este importantă.

O definiție de regulă are formatul:

```
<numeRegulă> = expansiuneDeRegulă ;
```

sau

```
public <numeRegulă> = expansiuneDeRegulă ;
```

Astfel, cele cinci componente ale unei definiții de regulă sunt: 1) cuvântul cheie `public` (opțional), 2) numele regulii ce urmează a fi definită, 3) semnul egal, 4) forma expandată a regulii și 5) caracterul punct și virgulă. Spațiile goale (caracterele space și tab) au relevanță numai în interiorul formei expandate a regulii și sunt ignorate în rest.

Forma expandată a unei reguli specifică modul în care o regulă poate fi pronunțată. Ea este o combinație logică de terminali (text ce poate fi pronunțat) și referințe către alte reguli.

Orice regulă din gramatică poate fi definită ca fiind publică utilizând cuvântul cheie **public**. O regulă publică poate fi utilizată în trei scopuri:

- poate fi referită într-o altă gramatică,
- poate fi folosită ca regulă activă în procesul recunoașterii vorbirii. Cu alte cuvinte o regulă publică poate fi utilizată de un sistem de recunoaștere a vorbirii pentru a transcrie vorbirea în text.
- poate fi referită local de către orice altă regulă publică sau non-publică definită în gramatica curentă.

Dacă cuvântul cheie **public** nu este specificat atunci regula este implicit privată și poate fi referită numai local, de către o altă regulă a gramaticii curente.

Expansiuni de reguli

Cele mai simple expansiuni de reguli sunt referințe către un terminal sau o altă regulă. De exemplu:

```
<a> = elefant;  
<b> = <x>;  
<c> = <com.acme.grammar.y>;
```

Regula **<a>** se expandează într-un singur terminal (**elefant**). Astfel pentru a spune regula **<a>**, vorbitorul trebuie să pronunțe “elefant”.

Regula **** se expandează folosind regula **<x>**. Astfel, pentru a spune regula ****, vorbitorul trebuie să pronunțe ceva conform regulii **<x>**. În mod similar, pentru a spune regula **<c>**, vorbitorul trebuie să pronunțe ceva conform regulii **<com.acme.grammar.y>**.

Cu alte cuvinte, expansiunile legale de reguli fac parte din următoarele categorii:

- orice terminal
- o referință către orice regulă publică sau non-publică definită în gramatica curentă
- o referință către o regulă publică dintr-o altă gramatică (regulă ce a fost importată în prealabil în gramatica curentă)

Nu sunt permise definițiile goale, precum cea de mai jos:

```
<d> = ; // definiție ilegală
```

În continuare vom explica modurile în care se pot defini reguli mai complexe prin combinații logice de expansiuni de reguli utilizând:

- secvențe de expansiuni de reguli și seturi de expansiuni alternative,
- gruparea expansiunilor de reguli folosind paranteze rotunde și pătrate,
- operatori unari (pentru repetarea unei expansiuni de regulă),
- atașarea de etichete specifice aplicațiilor.

Compunerea regulilor

O regulă poate fi definită ca o secvență de expansiuni de reguli. O secvență de expansiuni legale separate prin spațiu este ea însăși o expansiune legală. De exemplu, datorită faptului că atât terminalii, cât și referințele la alte reguli sunt expansiuni legale, următoarele sunt de asemenea expansiuni legale:

<unde> = Eu locuiesc în România;
<declarație> = acest <obiect> este <OK>;

Pentru a pronunța complet o secvență, item-urile din secvență trebuie pronunțate (toate) în ordinea definită de regulă. În primul exemplu, pentru a spune regula <unde>, vorbitorul trebuie să pronunțe cuvintele “Eu locuiesc în România” chiar în această ordine. Al doilea exemplu combină terminali cu referințe către alte reguli (<obiect> și <OK>). Pentru a spune regula <declarație>, vorbitorul trebuie să spună cuvântul *acest*, urmat de ceva care să fie se potrivească cu regula <obiect>, apoi cuvântul *este* și în final ceva care să fie se potrivească cu regula <OK>.

Într-o secvență putem avea orice expansiune legală, inclusiv structuri mai complexe similare cu cele ce vor fi descrise în continuare pentru alternative, grupuri și așa mai departe.

O regulă poate fi definită ca un set de expansiuni alternative separate prin caracterul bară verticală (|) și spații, ca în exemplele de mai jos:

<nume> = Maria | Ion | Daniel | Valentina | <alteNume>;

Pentru a spune regula <nume> vorbitorul trebuie să pronunțe unul și numai unul dintre expansiunile specificate în setul de alternative. De exemplu, vorbitorul ar putea spune cuvântul *Maria*, sau cuvântul *Ion*, sau cuvântul *Daniel*, sau cuvântul *Valentina* sau orice altceva care să se potrivească cu regula <alteNume>. Vorbitorul nu poate spune însă *Ion Daniel*.

Secvențele au precedență în fața alternativelor. De exemplu,

<oraș>= Cluj Napoca | Baia Mare | Satu Mare;

este un set de trei alternative, fiecare dintre ele reprezentând numele unui oraș.

Nu este permisă utilizarea unei alternative vide ca mai jos:

```
<nume> = lon | Maria; // ilegal
<nume> = lon | Maria | ; // ilegal
```

De obicei nu toate variantele dintr-un set de alternative sunt egal probabile. Fiecărei variante dintr-o alternativă îi poate fi asociată o pondere care să indice probabilitatea ca respectiva alternativă să fie pronunțată. O pondere este un număr aflat între două caractere slash (de exemplu /3.14/). Cu cât numărul este mai mare, cu atât probabilitatea ca acea alternativă să fie pronunțată este mai mare. Ponderea se plasează în fața fiecărei variante din setul de alternative, ca în exemplele de mai jos:

```
<dimensiune> = /10/ mic | /2/ mediu | /1/ mare;
<culoare> = /0.5/ roșu | /0.1/ albastru închis | /0.2/ verde deschis;
<ațiune> = te rog (/20/salvează fișierele | /1/șterge toate fișierele);
<locație> = /20/ <oraș> | /5/ <țară>;
```

Ponderile trebuie să reflecte frecvențele relative de apariție ale variantelor unui set de alternative. În primul exemplu se indică faptul că cuvântul **mic** are o probabilitate de apariție de zece ori mai mare decât cuvântul **mare** și de cinci ori mai mare decât cuvântul **mediu**.

Atunci când se folosesc ponderi pentru variantele unei alternative, trebuie respectate următoarele condiții:

- dacă se specifică o pondere pentru o variantă a setului de alternative, atunci trebuie să se specifice ponderi pentru toate variantele lui.
- ponderile pot fi numere raționale, următoarele variante fiind toate corecte: 56, 0.056, 3.14e3, 8f.
- între caracterele slash nu poate apărea altceva în afară de pondere și spații.
- ponderile pot avea și valoarea zero. O pondere cu valoarea zero indică faptul că acea variantă nu poate fi pronunțată. Ponderile cu valoarea zero pot fi utilizate la dezvoltarea gramaticilor.
- într-un set de alternative trebuie să existe cel puțin o variantă cu pondere nenulă.

Valorile potrivite pentru ponderi sunt de obicei dificil de determinat, iar “ghicirea” acestor valori nu duce întotdeauna la îmbunătățirea performanțelor sistemului de recunoaștere a vorbirii. Pentru a obține valori potrivite pentru ponderile unui set de alternative se realizează studii statistice ale unor situații reale de vorbire sau text.

Gruparea regulilor

Orice expansiune legală de regulă poate fi grupată explicit utilizând paranteze rotunde. Gruparea are precedență în fața altor operații astfel încât poate fi folosită pentru a asigura interpretarea corectă a regulilor. De asemenea, gruparea este utilă și pentru îmbunătățirea clarității definiției unei reguli. De exemplu, din cauză că secvențele au o

precedență mai mare decât alternativele, este nevoie de paranteze pentru a crea expansiunile `te rog închide` și `te rog șterge` în definiția următoare:

`<acțiune> = te rog (deschide | închide | șterge);`

Următorul exemplu prezintă o secvență de trei item-uri, fiecare dintre ele fiind un set de alternative înconjurate de paranteze rotunde, asigurând astfel gruparea lor corectă:

`<comandă> = (deschide | închide) (geamurile | ușile) (imediat | mai târziu);`

Pentru a spune ceva care să se potrivească cu regula `<comandă>`, vorbitorul trebuie să spună câte un cuvânt din fiecare set de alternative. De exemplu: `deschide geamurile imediat` sau `închide ușile mai târziu`.

Nu se pot defini grupuri vide. Exemplul de mai jos nu este corect:

`() // grupare ilegală`

Parantezele pătrate pot fi folosite pentru a indica faptul că o expansiune de regulă (cea din interiorul lor) este opțională. În alte privințe, ele sunt echivalente cu parantezele rotunde și au aceeași precedență. De exemplu, gramatica de mai jos:

`<formulăDePolitețe> = te rog | te implor;`
`public <cerință> = ajută-mă [ < formulăDePolitețe > ];`

îi permite vorbitorului să spună `ajută-mă` și să adauge, opțional, una dintre formulele de politețe `te rog` sau `te implor`.

Nu se pot defini grupuri opționale vide. Exemplul de mai jos nu este corect:

`[] // grupare opțională ilegală`

Operatori unari

Există trei operatori unari în JSGF: operatorul steluță Kleene, operatorul plus și etichetele. Toți operatorii unari au următoarele proprietăți:

- un operator unar poate fi atașat unei expansiuni de regulă legale,
- au precedență mare: sunt atașați expansiunii de regulă pe care o preced,
- unei expansiuni de regulă i se poate asocia un singur operator,
- spațiile dintre expansiunea de regulă și operatorul unar sunt ignorate.

Datorită faptului că precedența operatorilor unari este mai mare decât cea a secvențelor și alternativelor, este necesară folosirea parantezelor (gruparea) pentru a atașa un operator unar unei secvențe sau unei alternative.

O expansiune de regulă urmată de caracterul asterisc indică faptul că acea expansiune poate fi pronunțată de zero sau mai multe ori. Simbolul asterisc este cunoscut sub denumirea de steluța Kleene (după numele lui Stephen Cole Kleene). De exemplu, construcția

```
public <cerință> = ajută-mă <formulăDePolitețe>*;
```

îi permite vorbitorului să rostească *ajută-mă te rog* sau *ajută-mă te rog te implor te rog*, etc. sau să nu folosească nicio rugămințe și să spună scurt *ajută-mă*.

Precedența operatorului steluță este mai mare decât cea a secvențierii, astfel că în următoarea construcție:

```
public <cerință> = ajută-mă te rog *;
```

operatorul se aplică terminalului *rog*, nu întregii secvențe. Astfel, pentru a spune ceva care să se potrivească cu regula *<cerință>*, vorbitorul ar putea spune *ajută-mă*, *ajută-mă te* sau *ajută-mă te rog rog rog*, dar nu *ajută-mă te rog te rog*. Parantezele pot fi utilizate pentru a modifica acest comportament. De exemplu,

```
public <cerință> = ajută-mă (te rog) *;
```

O expansiune de regulă urmată de caracterul plus indică faptul că acea expansiune poate fi pronunțată de una sau mai multe ori. De exemplu, construcția

```
public <cerință> = ajută-mă <formulăDePolitețe>+;
```

necesită pronunția a cel puțin o formulă de politețe.

4.3.4 Câteva exemple de gramatici pentru niște sarcini uzuale

Gramatica dates pentru limba română

#JSGF V1.0;

grammar dates;

<day> = întâi | unu | doi | trei | patru | cinci | șase | șapte | opt | nouă | zece | unsprezece | doisprezece | treisprezece | paisprezece | cincisprezece | șaisprezece | șaptesprezece | optsprezece | nouăsprezece | douăzeci | douăzeci și unu | douăzeci și doi | douăzeci și trei | douăzeci și patru | douăzeci și cinci | douăzeci și șase | douăzeci și șapte | douăzeci și opt | douăzeci și nouă | treizeci | treizeci și unu;

<month> = întâia | a doua | a treia | a patra | a cincea | a șasea | a șaptea | a opta | a noua | a zecea | a unsprezecea | a doisprezecea | a douăsprezecea | ianuarie | februarie | martie | aprilie | mai | iunie | iulie | august | septembrie | octombrie | noiembrie | decembrie;

<year> = /80/ o mie nouă sute <tens> | /20/ două mii <smallYears>;

public <date> = <day> <month> [<year>];

Gramatica numbers pentru limba română

#JSGF V1.0;

grammar numbers;

<unitsWO012> = trei | patru | cinci | șase | șapte | opt | nouă;

<unitsWO0> = unu | una | doi | două | <unitsWO012>;

<units> = zero | <unitsWO0>;

<elevens> = zece | unsprezece | doisprezece | douăsprezece | treisprezece | paisprezece | cincisprezece | șaisprezece | șaptesprezece | optsprezece | nouăsprezece;

<simpleTens> = douăzeci | treizeci | patruzeci | cincizeci | șaizeci | șaptezeci | optzeci | nouăzeci;

<max2digitNumbers> = <units> | <elevens> | <simpleTens> [și <unitsWO0>];

<3digitNumbers> = (o sută | (două | <unitsWO012>) sute) [<max2digitNumbers>];

<max3digitNumbers> = <max2digitNumbers> | <3digitNumbers>;

<max6digitNumbers> = <max3digitNumbers> | ((o mie | <max3digitNumbers> [de] mii) [<max3digitNumbers>]);

<max9digitNumbers> = <max6digitNumbers> | ((un milion | <max3digitNumbers> [de] milioane) [<max6digitNumbers>]);

public <rationalNumbers> = <max9digitNumbers> [virgulă <max3digitNumbers>];

5. EVALUAREA SISTEMELOR DE RECUNOAȘTERE A VORBIRII CONTINUE

5.1 Introducere

Evaluarea sistemului de recunoaștere a vorbirii este o etapă esențială în dezvoltarea și optimizarea acestuia. Această etapă își propune să determine cât de adecvat este sistemul de recunoaștere pentru a transcrie vorbire, în contextul unei sarcini de recunoaștere bine precizate.

Evaluarea unui sistem de recunoaștere a vorbirii necesită existența unei baze de date de înregistrări audio etichetate, numită în continuare bază de date de evaluare. Pentru o evaluare corectă, baza de date de evaluare trebuie să reprezinte cât mai fidel sarcina de recunoaștere la care va fi supus sistemul:

- domeniul vorbirii: general sau particular (de exemplu: vorbire liberă, pe orice subiect sau discuții pe teme bine precizate: științifice/culturale/istorice/etc.);
- tipologia vorbirii: cuvinte pronunțate izolat, vorbire continuă citită, vorbire continuă spontană, etc.;
- caracteristicile vorbitorului/vorbitorilor: un singur vorbitor sau mai mulți vorbitori, vorbitori nativi/non-nativi, persoane tinere/în vârstă, etc.;
- mediul acustic: zgomotos, nezmotos.

5.2 Evaluarea sistemului de recunoaștere a vorbirii din punctul de vedere al domeniului vorbirii

Dintre cele patru caracteristici ale unei sarcini de recunoaștere de vorbire prezentate mai sus, domeniul vorbirii este singura care nu depinde de elementele acustice ale înregistrărilor. În consecință, sistemul de recunoaștere a vorbirii poate fi evaluat din acest punct de vedere utilizând numai *transcrierile clipurilor audio* din baza de date de evaluare,

așa cum ilustrează **Figura 5.1**. Bineînțeles că acest tip de evaluare vizează identificarea potențialelor probleme sau deficiențe ale *modelului de limbă* integrat în sistemul de recunoaștere a vorbirii.

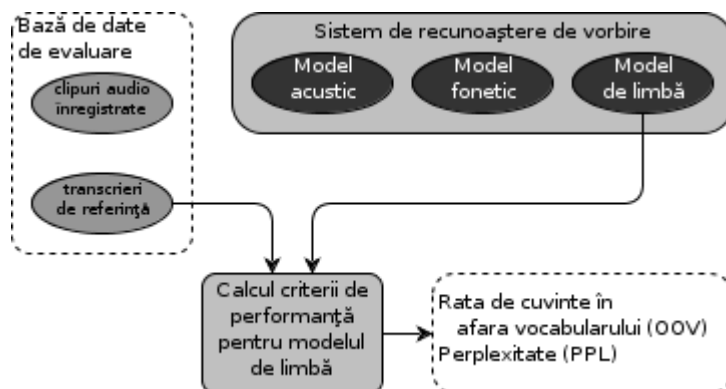


Figura 5.1. Evaluarea modelului de limbă din sistemul de recunoaștere

Într-un sistem de recunoaștere a vorbirii rolul modelului de limbă este de a estima probabilitățile secvențelor de cuvinte din spațiul de căutare. În consecință, modelul de limbă specifică ce secvențe de cuvinte sunt mai probabile într-un anumit context, pentru o anumită sarcină de recunoaștere, dar și ce secvențe de cuvinte nu sunt valide.

Capacitatea de predicție a modelului de limbă se poate evalua măsurând probabilitatea pe care el o asignează secvențelor de cuvinte din baza de date de evaluare. Un model de limbă adecvat, adaptat domeniului vorbirii, ar trebui să atribuie o probabilitate mare transcrierilor din baza de date de evaluare. Criteriul de performanță utilizat în acest caz este *perplexitatea*. Peplexitatea (*PPL*) este o mărime derivată din entropie ($H(p_{LM})$ – mărime ce măsoară incertitudinea unei distribuții de probabilitate) și se definește astfel:

$$H(p_{LM}) = -\frac{1}{n} \log p_{LM}(w_1, w_2, \dots, w_n) = -\frac{1}{n} \sum_{i=1}^n \log p_{LM}(w_i | w_1, w_2, \dots, w_{i-1}) \quad (5.1)$$

$$PPL = 2^{H(p_{LM})}$$

Probabilitățile $p_{LM}(w_1, w_2, \dots, w_n)$ și $p_{LM}(w_i | w_1, w_2, \dots, w_{i-1})$ sunt probabilitățile atribuite de modelul de limbă secvenței de cuvinte $w_1, w_2, \dots, w_n$, respectiv cuvântului w_i , în condițiile în care este precedat de secvența de cuvinte $w_1, w_2, \dots, w_{i-1}$. O perplexitate mai mică pentru o secvență de cuvinte particulară denotă o capacitate de predicție a secvenței mai mare pentru respectivul model de limbă.

Este posibil ca unele dintre cuvintele din textul de evaluare să nu se regăsească în vocabularul sistemului de recunoaștere a vorbirii și, în consecință, nici în vocabularul modelului de limbă. De aceea aceste cuvinte nu pot apărea în transcrierile rezultate în urma procesului de recunoaștere a vorbirii. Aceste cuvinte sunt numite *cuvinte în afara vocabularului* (*out of vocabulary* - *OOV*) și nu pot fi prezise de modelul de limbă

(probabilitatea lor de apariție este nulă). Cuvintele OOV fac evaluarea unui model de limbă mai dificilă: din cauză că perplexitatea lor este infinită, aceasta nu poate fi adunată la perplexitățile celorlalte cuvinte pentru a obține perplexitatea întregii secvențe de cuvinte. Astfel, pe lângă perplexitate, procentul de cuvinte OOV trebuie să fie și el specificat pentru ca evaluarea modelului de limbă să fie completă. Este evident că un sistem de recunoaștere a vorbirii bun trebuie să aibă un procent de cuvinte OOV cât mai mic.

5.2.1 Evaluarea perplexității și a ratei de cuvinte OOV: exemplu

Având la dispoziție transcrierile clipurilor audio din baza de date de evaluare și modelul de limbă (de tip n-gram) integrat în sistemul de recunoaștere a vorbirii, evaluarea perplexității și a ratei de cuvinte OOV poate fi realizată utilizând *pachetul de programe SRI-LM*, mai precis aplicația *ngram*. Mai jos este prezentată o selecție dintr-un raport de evaluare generat de această aplicație:

```
...
...
iaurturile ajută la prevenirea osteoporozei
  p( <unk> | <s> )    = [OOV] 0 [ -inf ]
  p( ajută | <unk> ...) = [1gram] 3.07228e-05 [ -4.51254 ]
  p( la | ajută ...)   = [2gram] 0.130606 [ -0.884037 ]
  p( prevenirea | la ...) = [3gram] 0.0334988 [ -1.47497 ]
  p( <unk> | prevenirea ...) = [OOV] 0 [ -inf ]
  p( </s> | <unk> ...) = [1gram] 0.0576355 [ -1.23931 ]
1 sentences, 5 words, 2 OOVs
0 zeroprobs, logprob= -8.11086 ppl= 106.59 ppl1= 505.381
...
...
opoziția e la fel de patetică
  p( opoziția | <s> )    = [2gram] 0.000168717 [ -3.77284 ]
  p( e | opoziția ...)   = [3gram] 0.013834 [ -1.85905 ]
  p( la | e ...)         = [3gram] 0.0211873 [ -1.67393 ]
  p( fel | la ...)       = [3gram] 0.15404 [ -0.812366 ]
  p( de | fel ...)       = [3gram] 0.384352 [ -0.415271 ]
  p( patetică | de ...)   = [2gram] 9.19136e-08 [ -7.03662 ]
  p( </s> | patetică ...) = [2gram] 0.222222 [ -0.653212 ]
1 sentences, 6 words, 0 OOVs
0 zeroprobs, logprob= -16.2233 ppl= 207.784 ppl1= 505.687
...
...
file transcripts.test: 100 sentences, 1703 words, 34 OOVs
0 zeroprobs, logprob= -4243.58 ppl= 251.315 ppl1= 350.038
```

În raportul de evaluare de mai sus se poate observa modul de calcul al perplexității pentru un fișier text numit *transcripts.test* ce conține 100 de propoziții de evaluare. În prima parte a exemplului este detaliată evaluarea făcută pentru două dintre cele 100 de propoziții. Utilizând modelul de limbă se estimează probabilitatea de apariție a fiecărui cuvânt din propoziție, în contextul respectiv. Se observă că primul cuvânt (“iaurturile”) este un OOV,

probabilitatea lui este 0, iar probabilitatea lui logaritmică este minus infinit. În consecință acest cuvânt nu va fi luat în considerare în calculul perplexității, dar va fi contorizat în cadrul ratei de cuvinte OOV. Al doilea cuvânt (“ajută”) face parte din vocabularul modelului de limbă (ML), iar probabilitatea lui de apariție (în afara contextului) este $3.07e-05$. Probabilitatea acestui cuvânt nu poate fi estimată ținând cont de contextul în care apare pentru că cuvântul ce îl precede (“iaurturile”) nu este modelat de ML. Probabilitatea de apariție a celui de-al treilea cuvânt (“la”) este estimată ținând cont de context (“la” precedat de “ajută”) și este aproximativ 0.13. Probabilitatea de apariție a celui de-al patrulea cuvânt (“prevenirea”) este estimată ținând cont de context (“prevenirea” precedat de “ajută la”) și este aproximativ 0.03. În fine, ultimul cuvânt din această propoziție (“osteoporozei”) este din nou un cuvânt OOV pentru care probabilitatea estimată de modelul de limbă este nulă. După detalierea acestor probabilități pentru cuvintele propoziției “iaurturile ajută la prevenirea osteoporozei”, raportul de evaluare sumarizează rezultatele obținute: propoziția conține 5 cuvinte, dintre care 2 cuvinte OOV. Probabilitatea logaritmică de apariție a acestei secvențe de cuvinte este -8.11, iar perplexitatea ei este 106.59.

Urmărind al doilea exemplu observăm că propoziția “opoziția e la fel de patetică” conține 6 cuvinte, dar niciun cuvânt OOV. Deși perplexitatea ei este mai mare decât în primul exemplu (modelul de limbă atribuie o probabilitate mai mică de apariție acestei secvențe de cuvinte), totuși faptul că numărul de cuvinte OOV este 0 este mai important în contextul recunoașterii vorbirii. Motivul este simplu: în primul caz sistemul de recunoaștere a vorbirii nu are nicio șansă să recunoască mai mult de 3 cuvinte din 5 (pentru că celelalte 2 nu se află în vocabularul sistemului). Chiar dacă în al doilea caz probabilitățile de apariție ale cuvintelor sunt mai mici, toate aceste cuvinte se află în vocabular și pot apărea în transcrierea generată de sistem. În concluzie, procentul de cuvinte OOV este cel mai important criteriu de performanță pentru modelul de limbă, iar atunci când această rată este similară pentru mai multe sisteme, perplexitatea este criteriul care îl identifică pe cel mai potrivit.

În finalul raportului de evaluare al modelului de limbă sunt sumarizate valorile pentru întregul text: 100 de propoziții, 1703 cuvinte, dintre care 34 cuvinte OOV, iar perplexitatea globală 251.3.

5.3 Evaluarea completă a sistemului de recunoaștere a vorbirii

În capitolul anterior am arătat cum *transcrierile clipurilor audio* din baza de date de evaluare pot fi utilizate pentru a determina cât de potrivit este un sistem de recunoaștere a vorbirii pentru a transcrie vorbirea dintr-un anumit domeniu. Evaluarea completă a sistemului de recunoaștere a vorbirii se face comparând în mod automat transcrierile clipurilor audio de evaluare (numite în continuare transcrieri de referință) cu transcrierile rezultate în urma

procesului de decodare a clipurilor audio de evaluare (numite în continuare transcrieri ipotetice), așa cum arată Figura 5.2.

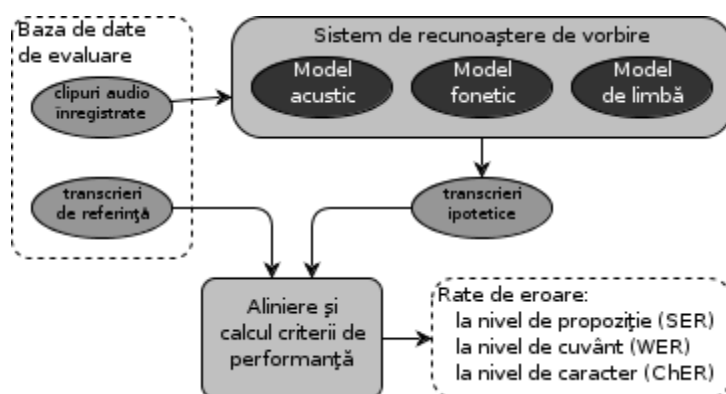


Figura 5.2. Evaluarea completă a sistemului de recunoaștere

Unul dintre criteriile de performanță utilizate în evaluarea sistemelor de recunoaștere a vorbirii este *rata de eroare la nivel de propoziție* (SER - sentence error rate). Rata de eroare la nivel de propoziție se calculează ca raport între numărul de propoziții eronate (propoziții ce conțin cel puțin o greșeală) și numărul total de propoziții din transcrierea de referință:

$$SER[\%] = \frac{\# \text{Propoziții eronate}}{\# \text{Propoziții în transcrierea de referință}} \times 100 \quad (5.2)$$

Acest criteriu de performanță este relevant numai în cazul în care transcrierea ipotetică a unei propoziții este inutilă dacă conține un cuvânt greșit (de exemplu într-un sistem de recunoaștere de comenzi de 1-2 cuvinte sau într-un sistem de recunoaștere al unui CNP sau IBAN, etc.). În cele mai multe cazuri acest criteriu de performanță nu este atât de relevant pentru că utilizatorul final poate înțelege mesajul transmis chiar și dacă transcrierea conține 10% cuvinte greșite (însă SER penalizează orice greșeală).

Având în vedere cele de mai sus, pentru evaluarea sistemelor de recunoaștere a vorbirii se folosește în mod uzual un alt criteriu de performanță: *rata de eroare la nivel de cuvânt* (WER – word error rate). Pentru calculul WER, analiza celor două texte se face frază cu frază, în două etape: a) fraza ipotetică și fraza de referință sunt aliniate printr-un algoritm care urmărește să minimizeze costul de editare al frazei ipotetice în vederea obținerii frazei de referință, după care b) se numără toate erorile de recunoaștere la nivel de cuvânt (cuvinte inserate, cuvinte substituite, respectiv cuvinte șterse). După aliniere, WER se calculează pe baza acestor trei tipuri de erori utilizând formula:

$$WER[\%] = \frac{\# \text{Erori de inserție} + \# \text{Erori de substituție} + \# \text{Erori de ștergere}}{\# \text{Cuvinte în transcrierea de referință}} \times 100 \quad (5.3)$$

Rata de eroare la nivel de cuvânt este mult mai adecvată pentru măsurarea performanțelor unui sistem de recunoaștere a vorbirii pentru că ea se calculează pe baza costului de editare al frazelor ipotetice. Totuși acest criteriu de performanță nu ține cont de

faptul că o eroare de substituție, de exemplu, poate fi reparată mult mai ușor dacă procentul de caractere eronate din cuvântul ipotetic este mic (ca în cazul “fotografie” “fotografieze”), decât dacă procentul de caractere eronate din cuvântul ipotetic este mare (ca în cazul “fotografie” “gravitez”). Din această cauză, în unele cazuri este mai adecvată folosirea ratei de eroare la nivel de caracter (ChER – character error rate) pentru evaluarea sistemelor de recunoaștere a vorbirii. Această rată de eroare se calculează în mod similar cu WER-ul, singura diferență fiind la nivelul unității de bază utilizate în comparație: caracterul, nu cuvântul:

$$ChER\% = \frac{\#Erori\ de\ insertie\ de\ caract + \#Erori\ de\ substituție\ de\ caract + \#Erori\ de\ ștergere\ de\ caract}{\#Caracter\ în\ transcrierea\ de\ referință} \times 100 \quad (5.4)$$

În concluzie, există trei criterii de performanță utilizate în evaluarea sistemelor de recunoaștere a vorbirii: rata de eroare la nivel de propoziție (SER), rata de eroare la nivel de cuvânt (WER) și rata de eroare la nivel de caracter (ChER). Dintre acestea, cel mai des folosit (criteriul standard) este WER-ul.

5.3.1 Evaluarea ratei de eroare la nivel de cuvânt: exemplu

Pachetul de programe standard utilizat în evaluarea sistemelor de recunoaștere a vorbirii este *NIST Scoring Toolkit (NIST SCTL)*. Din acest pachet de programe, aplicația *sclite* se poate folosi în mod explicit pentru alinierea transcrierilor ipotetice și transcrierilor de referință și calculul ratei de eroare la nivel de propoziție și de cuvânt. În continuare este prezentată o selecție dintr-un raport de evaluare generat de această aplicație pentru o sarcină uzuală de recunoaștere de vorbire continuă:

SYSTEM SUMMARY PERCENTAGES by SPEAKER

testProject.cd_cont_4000_16-1-1.match.spk.rec									
SPKR	# Snt	# Wrđ	Corr	Sub	Del	Ins	Err	S.Err	
0101	100	1703	75.9	20.8	3.3	5.7	29.8	88.0	
0201	100	1703	87.3	11.9	0.8	7.6	20.3	82.0	
0301	100	1703	86.3	11.7	2.1	6.5	20.3	74.0	
0401	100	1703	89.1	9.9	1.0	5.3	16.2	75.0	
0601	100	1703	89.8	9.3	0.9	3.2	13.4	68.0	
1001	100	1703	88.7	10.4	0.9	4.2	15.5	70.0	
1601	100	1703	87.7	10.5	1.8	2.7	15.0	78.0	
1701	100	1703	74.4	21.9	3.7	5.5	31.1	88.0	
1801	100	1703	87.3	11.2	1.5	3.3	16.0	74.0	
1901	100	1703	85.4	12.3	2.3	2.9	17.4	78.0	
2001	100	1703	87.1	10.6	2.3	3.2	16.0	71.0	
=====									
Sum/Avg	1100	18733	82.8	14.7	2.5	4.1	21.2	65.6	
=====									
Mean	141.7	1522.1	81.7	15.6	2.6	4.1	22.4	69.0	
S.D.	55.4	387.4	10.9	8.7	2.4	1.5	11.4	15.3	
Median	100.0	1703.0	86.3	11.9	2.0	3.6	19.7	74.0	

SENTENCE RECOGNITION PERFORMANCE

sentences		2976
with errors	65.6%	(1951)
with substitutions	62.7%	(1865)
with deletions	18.5%	(551)
with insertions	26.5%	(789)

WORD RECOGNITION PERFORMANCE

Percent Total Error	=	21.2%	(6785)
Percent Correct	=	82.8%	(26482)
Percent Substitution	=	14.7%	(4684)
Percent Deletions	=	2.5%	(798)
Percent Insertions	=	4.1%	(1303)
Percent Word Accuracy	=	78.8%	

Ref. words	=	(31964)
Hyp. words	=	(32469)
Aligned words	=	(33267)

CONFUSION PAIRS	Total	(3142)
	With >= 1 occurances	(3142)
1: 40 -> da ==> fda		
2: 21 -> bucata ==> bucată		
3: 18 -> am ==> a		

```

4: 18 -> da ==> la
5: 18 -> fotografiez ==> fotografieze
6: 18 -> returnez ==> returneze
7: 18 -> rezerv ==> rezerve
8: 16 -> da ==> dar
9: 15 -> cărămida ==> cărămidă
10: 14 -> o ==> rezervă
...
...
...

```

INSERTIONS

Total (496)
With >= 1 occurrences (496)

```

1: 67 -> de
2: 59 -> cu
3: 54 -> și
4: 39 -> a
5: 39 -> o
6: 37 -> în
7: 34 -> că
8: 30 -> nu
9: 29 -> să
10: 24 -> ce
...
...
...

```

SUBSTITUTIONS

Total (1022)
With >= 1 occurrences (1022)

```

1: 195 -> da
2: 97 -> o
3: 72 -> de
4: 72 -> să
5: 63 -> la
6: 63 -> rezerv
7: 51 -> e
8: 47 -> în
9: 46 -> a
10: 45 -> pe
...
...
...

```

FALSELY RECOGNIZED

Total (1849)
With >= 1 occurrences (1849)

```

1: 117 -> în
2: 99 -> și
3: 98 -> a
4: 78 -> o
5: 74 -> de
6: 56 -> la
7: 46 -> un
8: 44 -> să
9: 41 -> fda
10: 38 -> că
...
...
...

```

```

id: (0401_0510)
Scores: (#C #S #D #I) 4 1 0 3
REF:  iaurturile ajută la prevenirea **** ** ** OSTEOPOROZEI
HYP:  iaurturile ajută la prevenirea POST OP RO ZEI
Eval:                                     I   I   I   S
...
...
...
id: (0401_0537)
Scores: (#C #S #D #I) 8 3 0 1
REF:  am coborât din autocar ***** VAL VÂRTEJ și am început să FOTOGRAFIEZ
HYP:  am coborât din autocar BARBU TEI și și am început să FOTOGRAFIEZE
Eval:                                     I   S   S                               S
...
...
...
id: (0401_0556)
Scores: (#C #S #D #I) 11 1 1 0
REF:  aș vrea să afle cumva că GESTUL LUI          disperat nu ne-a fost
indiferent
HYP:  aș vrea să afle cumva că ***** GESTULUI disperat nu ne-a fost
indiferent
Eval:                                     D           S
...
...
...
id: (0602_0145)
Scores: (#C #S #D #I) 4 3 0 2
REF:  îMI PUTEȚI sugera niște opere interesante **** ** POSTBELICE
HYP:  îL PUTEM   sugera niște opere interesante POST DE ELICE
Eval: S      S                               I   I   S
...
...
...
id: (0603_0033)
Scores: (#C #S #D #I) 0 1 0 0
REF:  DA
HYP:  FDA
Eval: S
...
...
...
id: (1601_0534)
Scores: (#C #S #D #I) 10 2 0 0
REF:  omul a aruncat CĂRĂMIDA într-o fereastră a autocarului și a STRIGAT
ceva
HYP:  omul a aruncat CĂRĂMIDă într-o fereastră a autocarului și a STRIGA
ceva
Eval:                                     S                               S
id: (1601_0535)
Scores: (#C #S #D #I) 11 1 0 0
REF:  apoi a încercat să desfășoare BUCATA de pânză pe care scria ceva
HYP:  apoi a încercat să desfășoare BUCATĂ de pânză pe care scria ceva
Eval:                                     S
...
...
...

```

În prima parte a raportului de evaluare este prezentat un tabel ce conține informații sumariate privind evaluarea sistemului pentru fiecare dintre vorbitorii implicați în procesul de evaluare (câte un vorbitor pe fiecare dintre liniile tabelului). Pentru fiecare vorbitor sunt afișate numărul de clipuri audio de evaluare (coloana #Snt), numărul de cuvinte din aceste clipuri de evaluare (coloana #Wrd), procentul de cuvinte recunoscute corect (coloana Corr),

erorile procentuale de substituție, stergere și inserție (coloanele **Sub**, **Del** și **Ins**) și, în final, ratele de eroare la nivel de cuvânt și la nivel de propoziție (coloanele **Err** și **S.Err**). Ultima parte a tabelului prezintă valorile medii pentru toate aceste criterii de performanță.

Tabelul de la începutul raportului de evaluare prezintă o multitudine de criterii de performanță care pot fi utilizate pentru caracterizarea sistemului și optimizarea lui. De exemplu, dacă procentul de inserții este foarte mare comparativ cu procentul de cuvinte șterse, atunci ne putem aștepta să obținem rezultate mai bune dacă modificăm anumiți parametri specifici procesului de decodare: micșorarea **wordInsertionProbability** în acest caz. Informațiile din acest tabel pot fi utilizate și pentru a ne face o idee despre “duritatea” SER, comparativ cu WER: pentru vorbitorul 0401, de exemplu, avem doar 16% erori la nivel de cuvânt, însă acestea duc la un SER de 75%!

Chiar dacă toate informațiile prezentate în tabelul de evaluare sunt foarte utile, în momentul în care se cere o informație succintă despre performanța sistemului de recunoaștere de vorbire cel mai relevant criteriu de performanță (criteriul standard, de altfel) este rata de eroare (medie) la nivel de cuvânt. În consecință, utilizând un singur număr, sistemul evaluat în exemplul de mai sus este caracterizat de WER de 21.2%. Astfel de evaluări succinte sunt foarte utile atunci când ne propunem să optimizăm sistemul de recunoaștere de vorbire și dorim să prezentăm rezultatele obținute în urma procesului de optimizare. De exemplu, dacă variem doi dintre parametrii sistemului (să-i denumim alpha și beta) și dorim să punem în evidență dependența performanței sistemului de acești doi parametri, vom efectua o serie de

Tabelul 5.1. Exemplu de prezentare a rezultatelor unui proces de optimizare

WER [%]		alpha			
		100	200	500	1000
beta	0.8	18.3	14.9	11.4	13.2
	1.0	17.9	13.4	10.1	11.0
	1.2	16.2	13.5	8.6	12.0
	1.4	20.4	13.9	9.2	12.5

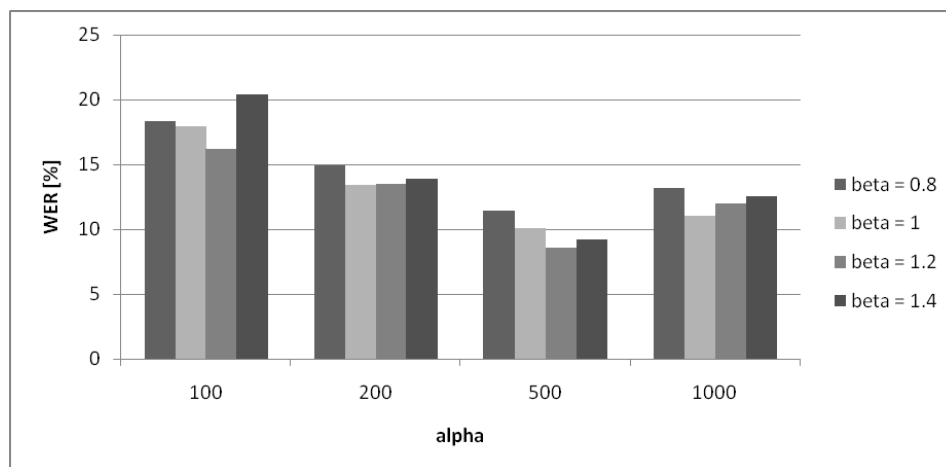


Figura 5.3. Exemplu de prezentare a rezultatelor unui proces de optimizare

experimente și vom prezenta numai WER-ul obținut pentru fiecare experiment într-un tabel similar cu Tabelul 5.1 și într-un grafic similar cu cel din Figura 5.3.

Evaluarea per vorbitor este și ea foarte interesantă întrucât ne arată dacă performanța sistemului de recunoaștere variază în funcție de caracteristicile vorbitorului. În cazul ideal sistemul ar trebui să funcționeze la fel de bine pentru orice vorbitor, însă în exemplul prezentat nu se întâmplă acest lucru. Pentru a facilita analiza acestor rezultate, ele ar trebui sumarizate succint într-un tabel separat (vezi Tabelul 5.2) și reprezentate grafic (vezi Figura 5.4). Din grafic se observă foarte ușor că există o serie de vorbitori pentru care sistemul de recunoaștere are performanțe bune (aproximativ 15% WER), doi vorbitori pentru care se obțin valori medii (aproximativ 20% WER) și doi vorbitori pe care sistemul îi recunoaște foarte slab (WER de aproximativ 30%).

Tabelul 5.2. Exemplu de evaluare per vorbitor pentru un sistem de recunoaștere de vorbire

WER [%]		Vorbitorul											
		01	02	03	04	06	10	16	17	18	19	20	medie
Setul de test	CS_01	29.8	20.3	20.3	16.2	13.4	15.5	15.0	31.1	16.0	17.4	16.0	21.2

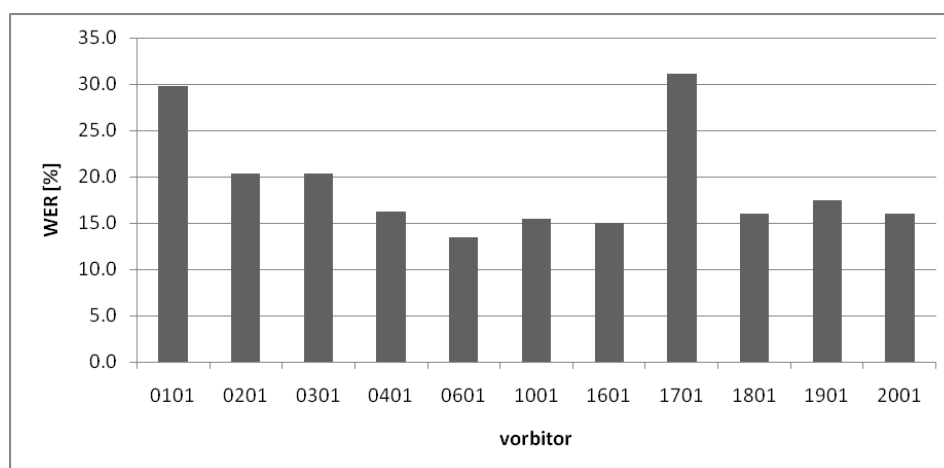


Figura 5.4. Exemplu de evaluare per vorbitor pentru un sistem de recunoaștere de vorbire

Raportul de evaluare continuă cu două secțiuni (*Sentence recognition performance* și *Word recognition performance*) care sumarizează criteriile de performanță medii prezentate în tabelul anterior.

Urmează alte cinci secțiuni foarte interesante care prezintă statistici despre perechile de cuvinte confundate cel mai frecvent (*Confusion Pairs*), cuvintele inserate cel mai frecvent (*Insertions*), cuvintele șterse cel mai frecvent (*Deletions*), cuvintele substituite cel mai frecvent (*Substitutions*) și cuvintele recunoscute fals cel mai frecvent (*Falsely Recognized*).

Analiza acestor cinci secțiuni poate furniza informații prețioase despre cele mai frecvente erori ale sistemului și poate reprezenta baza viitoarelor îmbunătățiri. Secțiunea

Confusion Pairs prezentată în exemplu ne arată faptul că un procent important din erori se datorează apariției cuvântului “fda” în locul cuvântului „da” în transcrierile ipotetice. Acest cuvânt nu este de fapt un cuvânt valid al limbii române. Probabil că el a apărut în vocabularul sistemului și în modelul de limbă ca urmare a unor erori statistice neidentificate în corpusurile de text de antrenare (pentru modelul de limbă). Exemplul evidențiază de asemenea că sunt confundate foarte frecvent formele de prezent la persoana I singular cu formele de prezent la persoana a doua singular („rezerv” -> „rezerve”, „fotografieez” -> „fotografieze”, etc.). Această eroare sistematică este tot un efect al utilizării unui model de limbă slab, care permite succesiuni de cuvinte de genul: “am început să fotografieze”. O altă confuzie frecventă ilustrată de această secțiune este cea între substantivele feminine articulate hotărât, respectiv nehotărât („bucata” -> „bucată”, „cărămidă” -> „cărămidă”, etc.). Ea este generată tot din cauza modelului de limbă, care permite succesiuni de cuvinte de genul: “să desfășoare bucată de pânză” sau “a aruncat cu cărămidă”. Iată cum secțiunea *Confusion Pairs* poate evidenția unele probleme (erori sistematice) ale sistemului de recunoaștere.

Secțiunea *Insertions* furnizează informații despre cele mai des întâlnite erori de inserție. Se poate observa că cuvintele inserate sistematic sunt cuvinte scurte de 1-2 foneme. Numărul erorilor de inserție poate fi scăzut prin mai multe metode:

- micșorarea parametrului de decodare [wordInsertionProbability](#),
- introducerea unor cuvinte OOV în vocabularul sistemului. Astfel se evită erori de genul: “opere interesante *postbelice*” -> „opere interesante *post de elice*” sau „prevenirea *osteoporozei*” -> „prevenirea *post op ro zei*”, apărute din cauza inexistenței cuvintelor “postbelice”, respectiv “osteoporozei” în vocabular,
- îmbunătățirea modelului de limbă prin tratarea corectă a tuturor cuvintelor cu cratimă, etc. Astfel se evită erori de genul “roluri care-și” -> „roluri care și” apărute datorită probabilității mult mai mici a secvenței “care-și”, comparativ cu secvența „care și”,
- îmbunătățirea modelului acustic prin introducerea unor modele speciale pentru sunetele ce nu reprezintă vorbire (respirații, plescăieli, zgomote scurte), evitând astfel inserții de cuvinte scurte în locul acestora.

Secțiunea *Substitutions* furnizează informații despre cele mai des întâlnite erori de substituție sau, altfel spus, despre cele mai frecvente cuvinte substituite incorect (cu un alt cuvânt). În exemplul dat mai sus în topul listei se află cuvinte precum „da”, „o”, „de”, „să”, „la”, etc. Acest tip de eroare apare în principal din cauza probabilității scăzute pe care o atribuie modelul de limbă respectivelor cuvinte sau secvențelor de cuvinte formate cu acestea.

În cele din urmă, secțiunea *Falsely Recognized* listează cele mai frecvente cuvinte introduse în mod eronat (în locul altora). Acest tip de eroare apare în principal din cauza probabilității mari pe care o atribuie modelul de limbă respectivelor cuvinte sau secvențelor

de cuvinte formate cu acestea. Remedierea acestui tip de erori poate fi abordată într-un mod similar cu cea pentru erorile de inserție.

Ultima parte a raportului de evaluare conține detalierea tuturor erorilor întâlnite în transcrierile ipotetice. În această secțiune analiza erorilor se face pentru fiecare frază în parte astfel:

- transcrierea de referință este aliniată cu transcrierea ipotetică marcându-se toate erorile cu una dintre literele S (substituție), D (ștergere) sau I (inserție),
- numărul de cuvinte recunoscute corect (#C) și numărul de erori de substituție (#S), ștergere (#D) și inserție (#I) sunt prezentate sub denumirea *Scores*.

Această ultimă secțiune a raportului de evaluare este foarte utilă pentru identificarea concretă a erorilor sistematice evidențiate în secțiunile anterioare ale raportului. Pe baza alinierii celor două transcrieri se pot trage concluzii importante ce pot fi folosite în continuare pentru a îmbunătăți sistemul de recunoaștere de vorbire.

5.4 Concluzii

Modelul de limbă al unui sistem de recunoaștere de vorbire poate fi evaluat din punctul de vedere al compatibilității lui cu domeniul vorbirii utilizând numai transcrierile de referință ale clipurilor audio din baza de date de evaluare. În acest caz criteriile standard de evaluare sunt rata de cuvinte în afara vocabularului (OOV) și perplexitatea (PPL).

Evaluarea completă a sistemului de recunoaștere a vorbirii se poate face numai decodând clipurile audio din baza de date de evaluare și comparând transcrierile ipotetice rezultate cu transcrierile de referință. În acest caz criteriile standard de evaluare sunt rata de eroare la nivel de propoziție (SER), rata de eroare la nivel de cuvânt (WER) și rata de eroare la nivel de caracter (ChER).

Analiza listei cuvintelor OOV și a erorilor aparute în procesul de decodare (substituții, ștergeri și inserții) este utilă în identificarea și caracterizarea erorilor sistematice și poate oferi informații prețioase pentru îmbunătățirea ulterioară a sistemului de recunoaștere a vorbirii.

6. CONSTRUCȚIA UNUI SISTEM SIMPLU DE RECUNOAȘTERE A VORBIRII CONTINUE

- recunoașterea secvențelor audio ce conțin cifre -

6.1 Introducere

Procesul recunoașterii vorbirii continue (RVC) vizează transformarea unui semnal audio ce conține vorbire într-o succesiune de cuvinte. Un sistem de recunoaștere a secvențelor audio de cifre este un sistem de recunoaștere a vorbirii continue cu vocabular redus, în sensul că singurele cuvinte ce pot fi recunoscute de către sistem sunt cele zece cifre ale sistemului zecimal: zero, unu, ..., nouă.

6.1.1 Arhitectura generală a unui sistem de recunoaștere a vorbirii continue - sumar

Arhitectura generală a unui sistem de recunoaștere automată a vorbirii (RAV) este prezentată în Figura 6.1. Din figură reies două lucruri esențiale în ceea ce privește procesul de recunoaștere a vorbirii: a) recunoașterea se face utilizând o serie de parametri vocali extrași din semnalul vocal (nu direct, utilizând mesajul vorbit) și b) recunoașterea se face pe baza unor modele (acustic, fonetic și lingvistic) dezvoltate în prealabil.

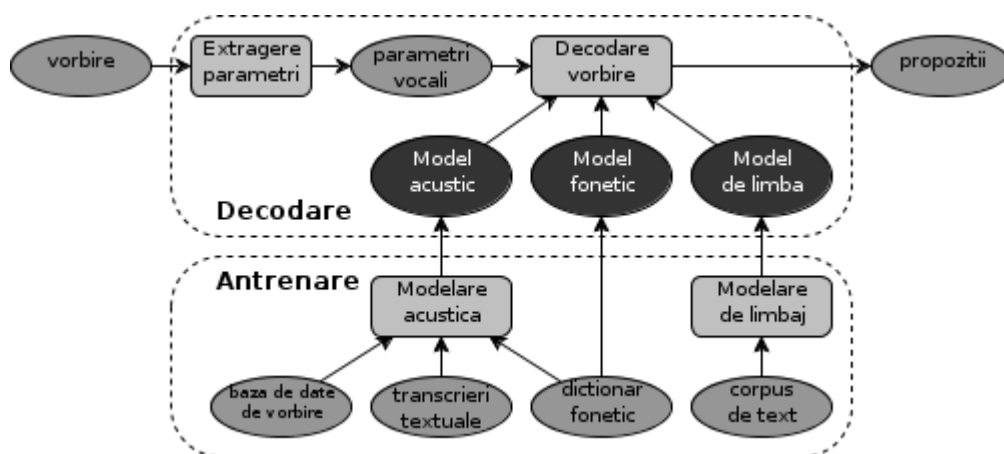


Figura 6.1. Arhitectura generală a unui sistem de recunoaștere a vorbirii

Modelul acustic are rolul de a estima probabilitatea unui mesaj vorbit, dată fiind o succesiune de cuvinte. În sistemele de RVC state-of-the-art modelul acustic nu folosește cuvinte ca unități acustice de bază pentru că: a) fiecare sarcină de RAV are un vocabular de cuvinte diferit pentru care nu există modele deja antrenate și nici date de antrenare disponibile și b) numărul de cuvinte diferite dintr-o limbă este prea mare. În locul cuvintelor, se utilizează unități acustice de bază sub-lexicale, de exemplu foneme, sau, mai recent, chiar unități sub-fonemice numite senone. Prin urmare, modelul acustic este format dintr-un set de modele pentru foneme (sau senone) care se conectează în timpul procesului de decodare pentru a forma modele pentru cuvinte și apoi modele pentru succesiuni de cuvinte. Acestea sunt utilizate în final pentru a estima probabilitatea ca mesajul rostit de vorbitor să fie format dintr-o succesiune sau alta de cuvinte. Experiența a arătat că această abordare generativă pentru modelele acustice poate fi foarte bine implementată utilizând modele Markov ascunse (hidden Markov models - HMM). Pentru mai multe informații despre modul cum funcționează modelele acustice consultați bibliografia îndrumarului de proiect.

Modelul de limbă este utilizat în timpul decodării pentru a estima probabilitățile tuturor secvențelor de cuvinte din spațiul de căutare. În general, rolul unui model de limbă este de a estima probabilitatea ca o secvență de cuvinte $W = w_1, w_2, \dots, w_n$, să fie o propoziție validă a limbii. Probabilitatea acestor secvențe de cuvinte ajută foarte mult modelul acustic în procesul de decizie. De exemplu, aceste două fraze: *casa ta e mare?* și *casat a Ema re?* sunt similare din punct de vedere acustic, însă cea de-a doua nu are niciun sens. Rolul modelului de limbă este acela de a asigura o probabilitate mult mai mare primei secvențe de cuvinte și, prin urmare, de a ajuta sistemul de RAV să decidă în favoarea acesteia.

Modelul fonetic are rolul de a conecta modelul acustic (care estimează probabilitățile acustice ale *fonemelor*) cu modelul de limbă (care estimează probabilitățile secvențelor de *cuvinte*). De cele mai multe ori modelul fonetic este un dicționar de pronunție care asociază fiecărui cuvânt din vocabular una sau mai multe secvențe de foneme adecvate, reprezentând modul în care se poate pronunța respectivul cuvânt.

Figura 6.1 ilustrează, de asemenea, procesele implicate în dezvoltarea unui sistem de RAV, dar și resursele necesare creării modelelor acustice, lingvistice și fonetice. Modelul acustic se construiește strict pe baza unui set de clipuri audio înregistrate asociate cu transcrierea textuală a mesajelor vorbite și a unui dicționar fonetic ce cuprinde toate cuvintele din respectiva transcriere textuală. În cazul sistemelor de RVC cu vocabular extins se folosesc modele de limbă statistice, care se construiesc utilizând corpusuri de text de dimensiuni cât mai mari și cât mai adaptate domeniului din care fac parte mesajele vorbite ce trebuie decodate. În sistemele de RVC cu vocabular redus (cum este și sistemul pe care îl veți dezvolta în continuare) se folosesc preponderent modele de limbă de tip gramatică cu stări finite, iar pentru construcția acestora nu este nevoie de resurse textuale.

6.2 Înregistrarea unei baze de date de clipuri audio

Așa cum am precizat mai sus, pentru a construi modelul acustic al unui sistem de recunoaștere a vorbirii este nevoie în primul rând de o bază de date de clipuri audio înregistrate. Pentru proiectul de recunoaștere de secvențe audio ce conțin cifre veți înregistra un set de 100 de clipuri audio a câte 12 cifre fiecare. Pentru înregistrarea clipurilor audio veți folosi o aplicație de înregistrări accesibilă on-line la adresa <http://speed.pub.ro/speech-recorder>. Pentru a accesa această aplicație aveți nevoie de numele de utilizator și de parola care v-au fost comunicate la prima ședință de proiect.

Lansați aplicația de înregistrări

1. Conectați un microfon la calculator și verificați ca volumul acestuia să fie la maxim.
Atenție: Nu folosiți pentru înregistrări microfonul încorporat al laptopului!
2. Porniți browser-ul și încărcați pagina <http://speed.pub.ro/speech-recorder>.
Notă: În cazul în care aplicația nu se încarcă va trebui să instalați plug-in-ul Java pentru browser.
3. Urmați instrucțiunile de pe primul ecran pentru a calibra microfonul și sistemul audio.
Notă: Etapa de calibrare are ca scop configurarea corectă a setărilor microfonului în condiții de liniște. Este foarte important ca înregistrările să se facă într-o cameră în care nu există surse de zgomot constant (de exemplu un ventilator sau un frigider), dar nici surse de zgomot intermitent (de exemplu un televizor, alte persoane care vorbesc, etc.).

În cazul în care aveți probleme la etapa de calibrare, urmăriți mesajele de eroare care apar și încercați să modificați volumul microfonului, să vorbiți mai tare/încet sau să utilizați setările speciale +10/20/30dB ale driverului audio. Dacă nu reușiți să realizați calibrarea contactați-l prin email pe îndrumătorul de proiect.

Conectați-vă la serverul de înregistrări

1. Conectați-vă la sistem folosind username-ul și parola furnizate de îndrumătorul de proiect.
2. Creați un vorbitor nou selectând **New Speaker** din câmpul **Select Speaker**. Completați apoi informațiile cerute și apăsați Save pentru a salva datele și *I agree* (pentru a accepta termenii de utilizare ai aplicației).
Notă: Această etapă este necesară numai dacă nu ați creat în prealabil un vorbitor cu numele vostru.
3. Selectați din câmpul **Select Speaker** numele vorbitorului nou creat.

Înregistrați clipurile audio

1. Selectați din câmpul **Phrase Group** setul de fraze specificat de îndrumătorul de proiect.

Notă: Aplicația va afișa fraza 1 a acestui grup de fraze și statusul ei (not recorded yet).

2. Apăsați butonul **Record**, așteptați o secundă și apoi **pronunțați fraza cât mai natural**.
3. Apăsați butonul **Stop** pentru a încheia înregistrarea.
Notă: Dacă verificarea semnalului audio se încheie cu succes fraza va fi încărcată pe server, iar aplicația va afișa automat fraza următoare.
4. Continuați cu înregistrările până când terminați de înregistrat toate frazele din acest set.
Notă: Dacă considerați că înregistrarea curentă nu a fost făcută corect (a apărut un zgomot neașteptat, v-ați bâlbâit, ați respirat zgomotos, ați pronunțat un cuvânt în plus sau în minus, ați ezitat, etc.) ascultați înregistrarea apăsând butonul **Play last**. Dacă înregistrarea nu a fost făcută corect reveniți la fraza anterioară apăsând butonul < și reînregistrați-o.
5. După ce ați terminat de înregistrat toate frazele utilizați butoanele <, >, <<, >> și **Play** pentru a parcurge și reasculta toate clipurile audio, reînregistrându-le pe cele care vi se par incorecte.
6. După ce ați efectuat toate înregistrările și ați verificat corectitudinea lor contactați-l prin email pe îndrumătorul de proiect.

6.3 Crearea dicționarului fonetic și a listei de foneme

Un dicționar fonetic este un instrument lingvistic care specifică modul în care se pronunță cuvintele unei limbi. Altfel spus, dicționarul fonetic face corespondența între forma scrisă și forma fonetică a cuvintelor unei limbi. Forma fonetică a unui cuvânt este o succesiune de foneme (fonemul fiind unitatea de sunet fundamentală a unei limbi). Fonemele limbii române, codificarea internă utilizată în continuare și câteva exemple de cuvinte transcrise fonetic au fost prezentate în Tabelul 2.1.

Într-un sistem de recunoaștere a vorbirii continue dicționarul fonetic are rolul de a face legătura între modelul acustic (care modelează modul de producere a *sunetelor specifice limbii*) și modelul de limbă (care modelează modul în care se succed *cuvintele limbii*). În conștință dicționarul fonetic trebuie să conțină toate cuvintele posibile pentru o anumită sarcină de recunoaștere a vorbirii alături, bineînțeles, de transcrierile fonetice ale acestor cuvinte. Pentru mai multe informații despre rolul dicționarului fonetic în cadrul unui sistem de recunoaștere a vorbirii continue consultați bibliografia îndrumarului de proiect.

Pentru sarcina curentă (recunoașterea secvențelor audio ce conțin cifre) dicționarul fonetic trebuie să conțină numai cele zece cifre ale sistemului zecimal: zero, unu, doi, trei,

patru, cinci, șase, șapte, opt și nouă. Sistemul de recunoaștere a vorbirii continue CMU Sphinx utilizează dicționare fonetice în format text, ce conțin câte un cuvânt pe fiecare linie a fișierului. Transcrierea fonetică a fiecărui cuvânt trebuie să apară pe aceeași linie. Fonemele trebuie separate de câte un spațiu. Crearea dicționarului fonetic în formatul agreeat este descrisă mai jos.

[Logați-vă la serverul de dezvoltare](#), apoi creați un director de lucru temporar:

```
mkdir phonetics
cd phonetics
```

Creați și editați un nou fișier numit `rodigits.dic` utilizând editorul vim:

```
vi rodigits.dic
```

Utilizând Tabelul 1 editați acest fișier astfel încât să aveți pe fiecare linie câte o cifră alături de transcrierea ei fonetică. Câteva instrucțiuni despre modul de editare al unui fișier text utilizând editorul vim sunt prezentate [aici](#). În final fișierul ar trebui să arate în felul următor:

```
zero z e r o
unu u n u
...
nouă n o1 w a1
```

Notă: în cazul în care nu puteți tasta caractere cu diacritice utilizând sistemul de operare și tastatura calculatorului pe care lucrați, puteți copia denumirile cifrelor ce conțin diacritice din lista următoare: șase, șapte, nouă. În consolă se poate da paste utilizând combinația de taste ctrl+alt+V.

În cazul sistemelor de recunoaștere a vorbirii cu vocabular redus (cum este și cel pe care tocmai îl construiți) manualul CMU Sphinx recomandă utilizarea de foneme specifice fiecărui cuvânt. Altfel spus ni se recomandă să modelăm diferit un același fonem X în funcție de cuvântul (și de poziția în cuvânt) în care apare acesta. În consecință veți modifica în continuare fișierul `rodigits.dic` adaugând în numele fiecărui fonem numele cuvântului din care face parte, cât și poziția fonemului în cuvânt (numai acolo unde este cazul). După această operație dicționarul fonetic ar trebui să arate astfel:

```
zero z_zero e_zero r_zero o_zero
unu u_unu1 n_unu u_unu2
...
nouă n_nouă o1_nouă w_nouă a1_nouă
```

Notă: în cazul în care nu puteți tasta caractere cu diacritice utilizând sistemul de operare și tastatura calculatorului pe care lucrați, puteți copia denumirile cifrelor ce conțin

diacritice din lista următoare: șase, șapte, nouă. În consolă se poate da paste utilizând combinația de taste ctrl+alt+V.

În continuare trebuie să creăm lista tuturor fonemelor ce vor fi modelate de modelul acustic. Este necesar să scriem un fișier text care să conțină câte un fonem pe fiecare linie, iar liniile trebuie să fie sortate în ordine alfabetică. Crearea acestui fișier se poate face manual (parcurgând dicționarul fonetic creat anterior) sau automat (urmând instrucțiunile de mai jos).

Creați un nou fișier (`rodigits.phones.temp`) care să conțină toate cuvintele și toate fonemele din dicționarul fonetic (câte un cuvânt/fonem pe linie):

```
sed 's/ /\n/g' rodigits.dic > rodigits.phones.temp
```

Din această listă selectați numai fonemele (profitând de faptul că ele conțin caracterul “_”) și listați-le într-un nou fișier numit `rodigits.phones.temp.onlyPhones`:

```
grep "_" rodigits.phones.temp > rodigits.phones.temp.onlyPhones
```

Adăugați la sfârșitul acestui fișier “fonemul” SIL (acest “fonem” este utilizat pentru a modela zonele de liniște):

```
echo "SIL" >> rodigits.phones.temp.onlyPhones
```

Sortați fonemele în ordine alfabetică, redenumiți fișierul rezultat și ștergeți toate celelalte fișiere temporare:

```
sort -u rodigits.phones.temp.onlyPhones > rodigits.phones.temp.onlyPhones.sorted  
mv rodigits.phones.temp.onlyPhones.sorted rodigits.phones  
rm rodigits.phones.temp*
```

6.4 Antrenarea modelului acustic

Antrenarea modelului acustic se va face în cadrul unui proiect CMU Sphinx și, așa cum am precizat în prima secțiune, presupune existența următoarelor resurse:

- un set de clipuri audio ce conțin vorbire (set pe care l-ați înregistrat în prealabil),
- transcrierea textuală corespunzătoare cuvintelor pronunțate în clipurile audio (deja existentă),
- un dicționar fonetic cuprinzând toate cuvintele (dicționar pe care l-ați creat în prealabil),
- un dicționar cu elemente acustice care nu sunt foneme (liniște, tuse, râs, muzică, etc.);

6.4.1 Crearea unui proiect CMU Sphinx numit rodigits

Logați-vă pe serverul de dezvoltare și creați un director numit `projects` în care veți crea ulterior toate proiectele:

```
cd ~  
mkdir projects
```

Intrați în directorul `projects` și creați un nou director numit `rodigits` pentru proiectul curent. Apoi intrați în directorul `rodigits`:

```
cd ~/projects  
mkdir rodigits  
cd ~/projects/rodigits
```

Configurați directorul `rodigits` pentru a putea fi utilizat ca proiect CMU Sphinx:

```
sphinxtrain_biosinf -t rodigits setup  
mkdir logdir
```

6.4.2 Adăugarea resurselor necesare în directorul proiectului

Logați-vă pe serverul de dezvoltare și intrați în directorul proiectului:

```
cd ~/projects/rodigits
```

Copiați clipurile audio înregistrate în prealabil în directorul `wav` (important: modificați comanda următoare pentru a include propriul `SPEAKER_ID`, așa cum vi l-a specificat îndrumătorul de proiect):

```
mkdir wav  
cp -r /home/biosinfShare/resources/wav/SPEAKER_ID/ ~/projects/rodigits/wav/
```

Copiați transcrierea textuală a cuvintelor pronunțate în clipurile audio în directorul `etc`:

```
cp /home/biosinfShare/resources/transcription/rodigits.data ~/projects/rodigits/etc/
```

Copiați dicționarul fonetic și lista de foneme create în prealabil în directorul `etc`:

```
cp ~/phonetics/rodigits.dic ~/projects/rodigits/etc/  
cp ~/phonetics/rodigits.phones ~/projects/rodigits/etc/
```

Copiați dicționarul cu elemente acustice care nu sunt foneme în directorul `etc`:

```
cp /home/biosinfShare/resources/phonetic/rodigits.filler ~/projects/rodigits/etc/
```

Copiați scripturile `createFileIds.sh` și `createTranscriptions.sh` în directorul `etc`:

```
cp -r /home/biosinfShare/scripts/createFileIds.sh ~/projects/rodigits/etc/  
cp -r /home/biosinfShare/scripts/createTranscriptions.sh ~/projects/rodigits/etc/
```

6.4.3 Transformarea fișierelor necesare procesului de antrenare în formatul corespunzător

Înainte de a putea începe antrenarea modelului acustic, trebuie creat un fișier cu lista clipurilor audio ce vor fi folosite la antrenare și un alt fișier cu transcrierea textuală a respectivelor clipuri audio. Din motive de eficiență vom crea tot acum și lista clipurilor audio ce vor fi folosite la evaluare, cât și fișierul cu transcrierea textuală a acestora.

Logați-vă pe serverul de dezvoltare și intrați în directorul proiectului, în subdirectorul `etc`:

```
cd ~/projects/rodigits/etc/
```

Vizualizați scriptul `createFileIds.sh` utilizând editorul vim:

```
vi createFileIds.sh
```

Intrați în modul de editare (tastând “i”) și schimbați valoarea variabilei `SPEAKER_ID` pentru a utiliza propriul `SPEAKER_ID`, așa cum vi l-a specificat îndrumătorul de proiect. Ieșiți din modul de editare (tastând Esc), salvați modificările (tastând “:w” și apoi Enter) și părăsiți editorul (tastând “:q” și apoi Enter).

Creați lista tuturor clipurilor audio înregistrate:

```
./createFileIds.sh > rodigits.fileids.all
```

Creați lista clipurilor audio ce vor fi utilizate pentru antrenare selectând primele 50 și ultimele 30 de clipuri audio din lista tuturor clipurilor (`rodigits.fileids.all`):

```
head -50 rodigits.fileids.all > rodigits.fileids.train  
tail -30 rodigits.fileids.all >> rodigits.fileids.train
```

Celelalte clipuri audio vor fi utilizate pentru evaluare. Creați un fișier cu lista acestora:

```
diff rodigits.fileids.train rodigits.fileids.all | grep '>' | sed 's/> //' > rodigits.fileids.test
```

Vizualizați scriptul `createTranscriptions.sh` utilizând editorul vim:

```
vi createTranscriptions.sh
```

Intrați în modul de editare (tastând “i”) și schimbați valoarea variabilei `SPEAKER_ID` pentru a utiliza propriul `SPEAKER_ID`, așa cum vi l-a specificat îndrumătorul de proiect. Ieșiți din modul de editare (tastând Esc), salvați modificările (tastând “:w” și apoi Enter) și părăsiți editorul (tastând “:q” și apoi Enter).

Creați transcrierea textuală (în formatul corespunzător cerut de CMU Sphinx) pentru toate clipurile audio:

```
./createTranscriptions.sh > rodigits.transcription.all
```

Creați un fișier care să conțină numai transcrierea textuală pentru clipurile audio ce vor utilizate pentru antrenare selectând primele 50 și ultimele 30 de linii din transcrierea textuală:

```
head -50 rodigits.transcription.all > rodigits.transcription.train
tail -30 rodigits.transcription.all >> rodigits.transcription.train
```

Celelalte linii din transcrierea textuală vor fi utilizate pentru evaluare. Creați un fișier cu lista acestora:

```
diff rodigits.transcription.train rodigits.transcription.all | grep '>' | sed 's/> //' >
rodigits.transcription.test
```

6.4.4 Configurarea procesului de antrenare

Orice proiect de recunoaștere de vorbire CMU Sphinx are un fișier de configurare se specifică parametrii modelului acustic ce va fi antrenat și locația resurselor necesare în procesul de antrenare. Pentru primul proiect pe care îl veți face nu veți pastra valorile implicite ale parametrilor modelului acustic, specificând numai locația resurselor.

Logați-vă pe serverul de dezvoltare, intrați în directorul proiectului, în subdirectorul *etc* și vizualizați fișierul de configurare al proiectului cu ajutorul editorului vim:

```
cd ~/projects/rodigits/etc/
vi sphinx_train.cfg
```

Identificați în acest fișier liniile care specifică locația resurselor utilizate în procesul de antrenare (dicționarul fonetic, lista de foneme, dicționarul de sunete care nu sunt foneme, lista de clipuri audio de antrenare și transcrierea textuală a acestor clipuri audio):

```
$CFG_DICTIONARY = "$CFG_LIST_DIR/$CFG_DB_NAME.dic";
$CFG_RAWPHONEFILE = "$CFG_LIST_DIR/$CFG_DB_NAME.phone";
$CFG_FILLERDICT = "$CFG_LIST_DIR/$CFG_DB_NAME.filler";
$CFG_LISTOFFILES = "$CFG_LIST_DIR/${CFG_DB_NAME}_train.fileids";
$CFG_TRANSCRIPTFILE = "$CFG_LIST_DIR/${CFG_DB_NAME}_train.transcription";
```

Intrați în modul de editare (tastând “i”) și modificați aceste linii astfel:

```
$CFG_DICTIONARY = "$CFG_LIST_DIR/rodigits.dic";
$CFG_RAWPHONEFILE = "$CFG_LIST_DIR/rodigits.phones";
$CFG_FILLERDICT = "$CFG_LIST_DIR/rodigits.filler";
$CFG_LISTOFFILES = "$CFG_LIST_DIR/rodigits.fileids.train";
$CFG_TRANSCRIPTFILE = "$CFG_LIST_DIR/rodigits.transcription.train";
```

Ieșiți din modul de editare (tastând Esc), salvați modificările (tastând “:w” și apoi Enter) și apoi închideți editorul (tastând “:q” și apoi Enter).

6.4.5 Crearea fișierelor cu coeficienți MFCC corespunzătoare clipurilor audio de antrenare

Logați-vă pe serverul de dezvoltare, intrați în directorul proiectului și creați fișierele cu coeficienți cepstrali executând următoarea comandă:

```
cd ~/projects/rodigits/  
/usr/local/sphinx/lib/sphinxtrain/scripts/000.comp_feat/slave_feat.pl
```

Notă: este normal ca în urma execuției comenzii de mai sus să apară eroarea `Failed to open control ... : No such file or directory at .../make_feats.pl line 89`. Ea ne atenționează asupra faptului că lista de fișiere audio de evaluare nu a fost găsită (acest lucru este normal, pentru că această listă nu a fost încă generată).

Execuția acestei comenzi are ca efect crearea a 80 de fișiere cu coeficienți cepstrali corespunzătoare celor 80 de clipuri audio de antrenare. Verificați faptul că fișierele cu coeficienți cepstrali au fost create corect (important: modificați comenzile următoare pentru a include propriul `SPEAKER_ID`, așa cum vi l-a specificat îndrumătorul de proiect):

```
ls feat/SPEAKER_ID/*  
ls feat/SPEAKER_ID/* | wc -l  
ls wav/SPEAKER_ID/* | wc -l
```

Prima comandă din grupul celor de mai sus va afișa toate fișierele din subdirectorul `feat/SPEAKER_ID`. A doua comandă va afișa numărul de fișiere din respectivul subdirector (rezultatul ar trebui să fie 80 pentru că doar 80 de clipuri audio se folosesc pentru antrenare). A treia comandă va afișa numărul de fișiere din subdirectorul `wav/SPEAKER_ID` (rezultatul ar trebui să fie 100 pentru că în total ar trebui să aveți 100 de clipuri audio înregistrate).

6.4.6 Antrenarea modelului acustic

Logați-vă pe serverul de dezvoltare, intrați în directorul proiectului și porniți antrenarea modelului acustic executând următoarea comandă:

```
cd ~/projects/rodigits/  
sphinxtrain_biosinf run
```

Așteptați terminarea procesului de antrenare (ar trebui să dureze aproximativ 2 minute). În cazul în care apar erori verificați faptul că fișierele cu resurse sunt în formatul corespunzător și că nu ați modificat greșit fișierul de configurare. În cazul în care nu reușiți să identificați și să rezolvați problema discutați cu îndrumătorul de proiect.

6.5 Crearea modelului de limbă (gramaticii)

Sistemele de recunoaștere de vorbire cu vocabular mare de ultimă generație utilizează modele de limbă statistice de tip n-gram. Aceste modele de limbă se construiesc pe baza unor corpusuri mari de text specific sarcinii de recunoaștere, estimând probabilitățile de apariție ale cuvintelor și secvențelor de cuvinte specifice respectivei sarcini. Modelele de limbă de tip n-gram se utilizează apoi în procesul de decodare (recunoaștere a vorbirii) pentru a calcula probabilitățile secvențelor de cuvinte propuse de modelul acustic. Pentru mai multe informații despre modele de limbă statistice și despre rolul modelului de limbă în cadrul unui sistem de recunoaștere a vorbirii continue consultați bibliografia îndrumarului de proiect.

Sarcina de recunoaștere a secvențelor audio ce conțin cifre este o sarcină de recunoaștere cu vocabular foarte redus pentru care nu se pretează modelele de limbă statistice. Mai mult, cifrele și succesiunile de cifre din cadrul clipurilor audio înregistrate apar cu probabilități aproximativ egale (nu se poate afirma că o cifră este pronunțată sistematic mai des decât alta). În aceste condiții, modelele de limbă de tip gramatică cu stări finite (FSG – finite state grammar) sunt mult mai potrivite. O gramatică cu stări finite este un model de tip graf în care nodurile reprezintă cuvinte ale limbii, iar tranzițiile între cuvinte sunt reprezentate de arcele grafului. Un astfel de model de limbă specifică în mod explicit toate secvențele de cuvinte permise de gramatica sarcinii de recunoaștere. Mai mult, fiecărui arc în parte îi poate fi asignat un cost specificând astfel probabilitatea ca un cuvânt să fie precedat de un altul (sau altfel spus probabilitatea respectivei secvențe de două cuvinte).

În **Figura 6.2** este reprezentată grafic gramatica cu stări finite a sarcinii noastre de recunoaștere. Se poate observa că modelul este format din 14 noduri, dintre care numai zece reprezintă cuvinte ale limbii (cifrele de la zero la nouă), iar celelalte patru sunt utilizate pentru a realiza „intrarea”, respectiv „ieșirea” din graf și pentru tranziția înapoi. Tranzițiile arată modul în care se poate parcurge această gramatică și implicit secvențele de cuvinte permise: în fiecare clip audio se pot pronunța una sau mai multe cifre.

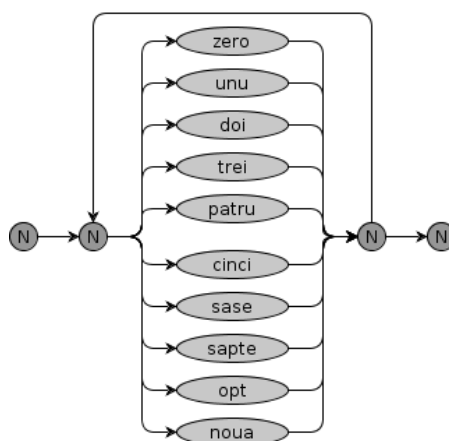


Figura 6.2. Gramatica cu stări finite a sarcinii de recunoaștere rodigits

Implementarea unei astfel de gramatici FSG poate fi făcută foarte simplu utilizând formatul Java Speech Grammar (JSGF) așa cum vom vedea în continuare.

Logați-vă pe serverul de dezvoltare, intrați în directorul proiectului, în subdirectorul `etc` și creați fișierul `rodigits.jsgf` utilizând editorul vim:

```
cd ~/projects/rodigits/etc/  
vi rodigits.jsgf
```

Intrați în modul de editare (tastând “i”) și scrieți următoarele trei linii:

```
#JSGF V1.0;  
grammar rodigits;  
public <numbers> = (zero | unu | doi | trei | patru | cinci | șase | șapte | opt | nouă) * ;
```

Ieșiți din modul de editare (tastând Esc), salvați modificările (tastând “:w” și apoi Enter) și apoi închideți editorul (tastând “:q” și apoi Enter). Prima linie pe care ați scris-o în fișier specifică numele gramaticii, iar cea de-a doua linie specifică cuvintele și succesiunile de cuvinte permise astfel: `numbers` este *zero* sau *unu* sau ... *nouă* de oricâte ori. Sau-ul este specificat prin caracterul |, iar faptul că cifrele se pot repeta de oricâte ori este specificat de caracterul *. Pentru mai multe detalii despre formatul Java Speech Grammar și modul în care puteți crea și alte tipuri de gramatici cu stări finite consultați site-ul JSGF precizat în bibliografia îndrumarului de proiect.

Transformați (cu ajutorul utilitarului `sphinx_jsgf2fsg`) gramatica sarcinii de recunoaștere din format JSGF (Java Speech Grammar Format) în formatul intern FSG (Finite State Grammar) utilizat de CMU Sphinx:

```
sphinx_jsgf2fsg -jsgf rodigits.jsgf -fsg rodigits.fsg
```

Vizualizați fișierul rezultat (`rodigits.fsg`) utilizând editorul vim:

```
vi rodigits.fsg
```

Observați modul în care se specifică nodurile, tranzițiile și probabilitățile de tranziție ale gramaticii cu stări finite în formatul intern utilizat CMU Sphinx. Ca exercițiu, desenați pe o bucată de hârtie graful definit în acest fișier și comparați-l cu graful prezentat în **Figura 6.2**.

6.6 Evaluarea sistemului de recunoaștere și interpretarea rezultatelor

Cele trei componente fundamentale ale unui sistem de recunoaștere a vorbirii (modelul acustic, modelul de limbă și modelul fonetic) sunt în acest moment disponibile. În consecință se poate trece la decodarea clipurilor audio de evaluare și apoi la compararea

transcrierii textuale rezultate cu transcrierea textuală de referință în vederea evaluării sistemului de recunoaștere de secvențe audio ce conțin cifre.

6.6.1 Configurarea procesului de decodare (și evaluare)

Pentru a configura procesul de decodare vom edita același fișier de configurare care a fost modificat și pentru antrenare (`sphinx_train.cfg`). De această dată nu ne vom concentra asupra parametrilor de antrenare, ci asupra parametrilor de decodare.

Logați-vă pe serverul de dezvoltare, intrați în directorul proiectului, în subdirectorul `etc` și vizualizați fișierul de configurare al proiectului cu ajutorul editorului `vim`:

```
cd ~/projects/rodigits/etc/  
vi sphinx_train.cfg
```

Identificați în acest fișier linia care specifică numele scriptului utilizat pentru decodare:

```
$DEC_CFG_SCRIPT = 'psdecode.pl';
```

Intrați în modul de editare (tastând “i”) și modificați această linie astfel:

```
$DEC_CFG_SCRIPT = 'psdecode_fsg.pl';
```

Identificați apoi liniile care specifică numele și locația resurselor utilizate în procesul de decodare (dicționarul fonetic, dicționarul de sunete care nu sunt foneme, lista de clipuri audio de evaluare și transcrierea textuală a acestor clipuri audio) și numele subdirectorului în care se vor salva rezultatele:

```
$DEC_CFG_DICTIONARY   = "$CFG_BASE_DIR/etc/$CFG_DB_NAME.dic";  
$DEC_CFG_FILLERDICT   = "$CFG_BASE_DIR/etc/$CFG_DB_NAME.filler";  
$DEC_CFG_LISTOFFILES  = "$CFG_BASE_DIR/etc/${CFG_DB_NAME}_test.fileids";  
$DEC_CFG_TRANSCRIPTFILE = "$CFG_BASE_DIR/etc/${CFG_DB_NAME}_test.transcription";  
$DEC_CFG_RESULT_DIR   = "$CFG_BASE_DIR/result";
```

Modificați aceste linii astfel:

```
$DEC_CFG_DICTIONARY   = "$CFG_BASE_DIR/etc/rodigits.dic";  
$DEC_CFG_FILLERDICT   = "$CFG_BASE_DIR/etc/rodigits.filler";  
$DEC_CFG_LISTOFFILES  = "$CFG_BASE_DIR/etc/rodigits.fileids.test";  
$DEC_CFG_TRANSCRIPTFILE = "$CFG_BASE_DIR/etc/rodigits.transcription.test";  
$DEC_CFG_RESULT_DIR   = "$CFG_BASE_DIR/result.cd_cont_200_8";
```

Identificați apoi linia care specifică numele modelului de limbă ce va fi utilizat la decodare:

```
$DEC_CFG_LANGUAGEMODEL = "$CFG_BASE_DIR/etc/${CFG_DB_NAME}.lm.DMP";
```

Și modificați-o astfel:

```
$DEC_CFG_LANGUAGEMODEL = "$CFG_BASE_DIR/etc/rodigits.fsg";
```

Identificați apoi linia care specifică aplicația cu care se va face alinierea textului rezultat în urma recunoașterii cu textul de referință (în vederea evaluării ratei de eroare la nivel de cuvânt):

```
$DEC_CFG_ALIGN = "builtin";
```

Și modificați-o astfel:

```
$DEC_CFG_ALIGN = "sclite";
```

Ieșiți din modul de editare (tastând Esc), salvați modificările (tastând “:w” și apoi Enter) și apoi închideți editorul (tastând “:q” și apoi Enter).

6.6.2 Crearea fișierelor cu coeficienți MFCC corespunzătoare clipurilor audio de evaluare

Logați-vă pe serverul de dezvoltare, intrați în directorul proiectului și creați fișierele cu coeficienți cepstrali executând următoarea comandă:

```
cd ~/projects/rodigits/  
/usr/local/sphinx/lib/sphinxtrain/scripts/000.comp_feat/slave_feat.pl
```

Execuția acestui script are ca efect crearea a 100 de fișiere cu coeficienți cepstrali corespunzătoare tuturor clipurilor audio înregistrate (atât cele utilizate pentru antrenare, cât și cele utilizate pentru evaluare). Verificați faptul că fișierele cu coeficienți cepstrali au fost create corect (important: modificați comenzile următoare pentru a include propriul `SPEAKER_ID`, așa cum vi l-a specificat îndrumătorul de proiect):

```
ls feat/SPEAKER_ID/*  
ls feat/SPEAKER_ID/* | wc -l  
ls wav/SPEAKER_ID/* | wc -l
```

Prima comandă din grupul celor de mai sus va afișa toate fișierele din subdirectorul `feat/SPEAKER_ID`. A doua comandă va afișa numărul de fișiere din respectivul subdirector (rezultatul ar trebui să fie 100 pentru că în total sunt 100 de clipuri audio înregistrate). A treia comandă va afișa numărul de fișiere din subdirectorul `wav/SPEAKER_ID` (rezultatul ar trebui să fie tot 100).

6.6.3 Decodarea clipurilor audio de evaluare

Logați-vă pe serverul de dezvoltare, intrați în directorul proiectului și porniți procesul de evaluare executând următoarea comandă:

```
cd ~/projects/rodigits/  
/usr/local/sphinx/lib/sphinxtrain/scripts/decode/slave.pl
```


Așteptați terminarea procesului de decodare (ar trebui să dureze aproximativ 1 minut). În cazul în care apar erori încercați să identificați problema vizualizând fișierul de log al procesului de decodare:

```
vi logdir/decode/rodigits-1-1.log
```

În cazul în care nu reușiți să identificați și să rezolvați problema discutați cu îndrumătorul de proiect. În cazul în care procesul de decodare se încheie cu succes ar trebui ca pe ecran să vă apară un mesaj similar cu acesta:

```
MODULE: DECODE Decoding using models previously trained (decode script:
psdecode_fsg.pl)
  Decoding 20 segments starting at 0 (part 1 of 1)
  0%
  Aligning results to find error rate
  SENTENCE ERROR: 50.0% (10/20) WORD ERROR RATE: 12.9% (30/240)
```

6.6.4 Vizualizarea și interpretarea rezultatelor evaluării

Evaluarea unui sistem de recunoaștere a vorbirii se face comparând în mod automat cele două transcrieri textuale ale clipurilor audio de evaluare: transcrierea textuală de referință (cea despre care se știe a priori că este corectă) și transcrierea textuală ipotetică (cea rezultată în urma procesului de decodare). Analiza se face frază cu frază, în două etape: a) fraza ipotetică și fraza de referință sunt aliniate printr-un algoritm care urmărește să minimizeze numărul de erori de transcriere, după care b) se numără toate erorile de recunoaștere la nivel de cuvânt (cuvinte inserate, cuvinte substituie, respectiv cuvinte șterse). Se disting astfel două criterii de performanță standard folosite în evaluarea sistemelor de recunoaștere a vorbirii continue: rata de eroare la nivel de propoziție (SER - sentence error rate) și rata de eroare la nivel de cuvânt (WER – word error rate). Rata de eroare la nivel de propoziție se calculează ca raport între numărul de propoziții corecte (propoziții ce nu conțin nicio greșeală) și numărul total de propoziții. Rata de eroare la nivel de cuvânt se calculează cu formula:

$$WER\% = \frac{\#Erori\ de\ insertie + \#Erori\ de\ substituie + \#Erori\ de\ ștergere}{\#Cuvinte\ în\ transcrierea\ de\ referință} \times 100 \quad (6.1)$$

Așa cum se observă din mesajul afișat în urma execuției procesului de decodare (dat ca și exemplu mai sus), în urma evaluării transcrierii ipotetice s-a obținut o rată de eroare la nivel de propoziție de 50% și o rată de eroare la nivel de cuvânt de 12.9%. Mai multe informații despre rezultatele evaluării sistemului de recunoaștere puteți afla vizualizând întreg raportul de evaluare:

```
cd ~/projects/rodigits/
vi result.cd_cont_200_8/rodigits.align
```

Primul tabel din raportul de evaluare afișează pe câte o linie sumarul rezultatelor pentru diverșii vorbitori din baza de date de evaluare. În cazul de față aveți un singur vorbitor în baza de date, deci o singură linie în tabel. Pe coloane aveți mai multe informații: id-ul vorbitorului, numărul de propoziții de evaluare, numărul de cuvinte de evaluare, procentul de cuvinte recunoscute corecte, procentul de cuvinte substituite/șterse/inserate, rata de eroare la nivel de cuvânt și rata de eroare la nivel de propoziție.

La secțiunea *Sentence Recognition Performance* se oferă mai multe detalii cu privire la erorile la nivel de propoziție, iar la secțiunea *Word Recognition Performance* se oferă mai multe detalii cu privire la erorile la nivel de cuvânt.

Urmează alte cinci secțiuni foarte interesante care prezintă statistici despre perechile de cuvinte confundate cel mai frecvent (*Confusion Pairs*), cuvintele inserate cel mai frecvent (*Insertions*), cuvintele șterse cel mai frecvent (*Deletions*), cuvintele substituite cel mai frecvent (*Substitutions*) și cuvintele recunoscute fals cel mai frecvent (*Falsely Recognized*).

În final, raportul de evaluare prezintă toate frazele ipotetice aliniate la frazele de referință corespunzătoare specificând, de asemenea, și tipurile de erori apărute la fiecare frază în parte.

Toate aceste informații și rezultate din raportul de evaluare al sistemului de recunoaștere vor fi comparate cu rezultatele ce vor fi obținute în secțiunea următoare și vor fi în final prezentate și discutate în raportul final al proiectului de cercetare-dezvoltare.

6.7 Optimizarea modelului acustic

Înainte de a putea face câteva optimizări simple ale modelului acustic trebuie să studiați și să înțelegeți câteva aspecte importante privind modul de construcție al unui astfel de model:

- arhitectura generală a modelului acustic (platforma HMM-GMM). Informații despre acest subiect găsiți în bibliografia îndrumarului de proiect: [3] pag. 31, [4] pag. 23, [5] pag. 8, [6] pag. 306.
- alegerea unităților de vorbire (foneme, trifoneme, senone) și dependența/independența de context. Informații despre aceste subiecte găsiți în bibliografia îndrumarului de proiect: [3] pag. 33, [5] pag. 52, [6] pag. 310.

Utilizând deja noțiunile amintite mai sus (și explicate pe larg în bibliografie), vom sumariza în continuare modul particular în care toolkitul de recunoaștere a vorbirii pe care îl folosim (CMU Sphinx) realizează antrenarea modelului acustic. În mod implicit modelul acustic creat de CMU Sphinx este construit cu unități de vorbire dependente de context de tip trifonem (fonem căruia i se precizează vecinii stânga-dreapta). Tot în mod implicit aceste trifoneme sunt implementate ca HMM-uri cu trei stări emiseive, fiecare având o distribuție de

probabilitate de ieșire implementată cu un GMM. Fiecare astfel de stare-model se numește *senone* și ea poate fi comună mai multor trifoneme (în funcție de similaritățile dintre ele: de exemplu trifonemele p-a-r și t-a-v pot avea senone-ul inițial comun). În mod implicit numărul de senone este configurat la 200, iar numărul de densități de probabilitate pentru fiecare GMM este configurat la 8.

În consecință modelul acustic antrenat în prealabil modelează cele câteva foneme specificate în dicționarul fonetic utilizând modele fonetice dependente de context având în total 200 de senone, fiecare senone folosind câte 8 densități Gaussiene de probabilitate pentru a modela parametri acustici de ieșire. Procesul de antrenare implementat de CMU Sphinx conduce la obținerea acestui model acustic astfel:

1. inițializarea modelelor fonetice independente de context (câte un model cu câte trei senone folosind câte o singură densitate Gaussiană pentru fiecare *fonem* în parte)
2. antrenarea modelelor fonetice (independente de context) rezultate la pasul anterior
3. transformarea modelelor fonetice independente de context în modele fonetice dependente de context (se obține câte un model cu câte trei senone folosind câte o singură densitate Gaussiană pentru fiecare *trifonem* în parte; în total modelul acustic poate avea *oricâte senone!*)
4. antrenarea modelelor fonetice (dependente de context) rezultate la pasul anterior
5. legarea stărilor (senonelor) modelelor fonetice dependente de context pe baza similarităților dintre ele (se obține câte un model cu câte trei senone folosind câte o singură densitate Gaussiană pentru fiecare trifonem în parte; în total modelul acustic va avea *200 senone!*)
6. antrenarea modelelor fonetice (dependente de context) rezultate la pasul anterior
7. dublarea numărului de densități Gaussiene în cadrul fiecărei senone (se obține câte un model cu câte trei senone folosind câte *două densități Gaussiane* pentru fiecare trifonem în parte; în total modelul acustic va avea *200 senone!*)
8. antrenarea modelelor fonetice (dependente de context) rezultate la pasul anterior
9. dublarea numărului de densități Gaussiene în cadrul fiecărei senone (se obține câte un model cu câte trei senone folosind câte *patru densități Gaussiane* pentru fiecare trifonem în parte; în total modelul acustic va avea *200 senone!*)
10. antrenarea modelelor fonetice (dependente de context) rezultate la pasul anterior
11. dublarea numărului de densități Gaussiene în cadrul fiecărei senone (se obține câte un model cu câte trei senone folosind câte *opt densități Gaussiane* pentru fiecare trifonem în parte; în total modelul acustic va avea *200 senone!*)

În consecință, pentru a obține modelul acustic cu modele fonetice dependente de context având 200 de senone și 8 densități Gaussiene per senone s-a trecut și prin etapele în care modelul acustic era format din:

- modele fonetice independente de context (după pasul 2)
- modele fonetice dependente de context cu o singură densitate Gaussiană per senone (după pasul 6)

- modele fonetice dependente de context cu 2 densități Gaussiane per senone (după pasul 8)
- modele fonetice dependente de context cu 4 densități Gaussiane per senone (după pasul 10)

Toate aceste modele se află în directorul proiectului, în subdirectorul `model_parameters`. Puteți verifica acest lucru astfel:

```
cd ~/projects/rodigits  
ls -all model_parameters
```

Lista subdirectoarelor afișate ar trebui să fie următoarea:

```
rodigits.cd_cont_200  
rodigits.cd_cont_200_1  
rodigits.cd_cont_200_2  
rodigits.cd_cont_200_4  
rodigits.cd_cont_initial  
rodigits.cd_cont_untied  
rodigits.ci_cont  
rodigits.ci_cont_flatinitial
```

În mod implicit decodarea realizată în capitolul 6.6 a utilizat modelele din subdirectorul `rodigits.cd_cont_200` (modelele fonetice dependente de context cu 8 densități Gaussiane per senone). În continuare ne propunem să configurăm procesul de decodare pentru a utiliza pe rând modelele din subdirectoarele:

- `rodigits.ci_cont` (modelele fonetice independente de context)
- `rodigits.cd_cont_200_1` (modelele fonetice dependente de context cu o densitate Gaussiană per senone)
- `rodigits.cd_cont_200_2` (modelele fonetice dependente de context cu 2 densități Gaussiane per senone)
- `rodigits.cd_cont_200_4` (modelele fonetice dependente de context cu 4 densități Gaussiane per senone)

6.7.1 Decodarea clipurilor audio de evaluare folosind modele fonetice independente de context

Logați-vă pe serverul de dezvoltare, intrați în directorul proiectului, în subdirectorul `etc` și vizualizați fișierul de configurare al proiectului cu ajutorul editorului vim:

```
cd ~/projects/rodigits/etc/  
vi sphinx_train.cfg
```

Identificați în acest fișier linia care specifică modelele acustice utilizate pentru decodare:

```
# Models to use.
$DEC_CFG_MODEL_NAME =
"$CFG_EXPTNAME.cd_${CFG_DIRLABEL}_${CFG_N_TIED_STATES}";
```

Intrați în modul de editare (tastând “i”), copiați linia respectivă și comentați una dintre copii (pentru a păstra un back-up) și modificați cealaltă copie astfel încât să specificați locația modelelor fonetice independente de context. Această zonă a fișierului ar trebui să arate în final ca mai jos:

```
# Models to use.
#$DEC_CFG_MODEL_NAME =
"$CFG_EXPTNAME.cd_${CFG_DIRLABEL}_${CFG_N_TIED_STATES}";
$DEC_CFG_MODEL_NAME = "$CFG_EXPTNAME.ci_${CFG_DIRLABEL}";
```

Identificați apoi linia care specifică numele subdirectorului în care se vor salva rezultatele:

```
$DEC_CFG_RESULT_DIR = "$CFG_BASE_DIR/result.cd_cont_200_8";
```

Și modificați-o astfel încât numele subdirectorului în care se vor salva rezultatele să fie sugestiv:

```
$DEC_CFG_RESULT_DIR = "$CFG_BASE_DIR/result.ci_cont";
```

Ieșiți din modul de editare (tastând Esc), salvați modificările (tastând “:w” și apoi Enter) și apoi închideți editorul (tastând “:q” și apoi Enter).

Intrați apoi în directorul proiectului și porniți procesul de evaluare executând următoarea comandă:

```
cd ~/projects/rodigits
/usr/local/sphinx/lib/sphinxtrain/scripts/decode/slave.pl
```

Așteptați terminarea procesului de decodare (ar trebui să dureze aproximativ 1 minut). În cazul în care apar erori încercați să identificați problema vizualizând fișierul de log al procesului de decodare:

```
vi logdir/decode/rodigits-1-1.log
```

În cazul în care nu reușiți să identificați și să rezolvați problema discutați cu îndrumătorul de proiect. În cazul în care procesul de decodare se încheie cu succes ar trebui ca pe ecran să vă apară un mesaj similar cu acesta:

```
MODULE: DECODE Decoding using models previously trained (decode script:
psdecode_fsg.pl)
  Decoding 20 segments starting at 0 (part 1 of 1)
  0%
  Aligning results to find error rate
  SENTENCE ERROR: 55.0% (11/20)  WORD ERROR RATE: 6.3% (15/240)
```

Aceste rezultate, precum și informațiile și rezultatele detaliate din raportul de evaluare generat la acest pas (`result.ci_cont/rodigits.align`) vor fi comparate cu rezultatele obținute în prealabil și cu rezultatele ce vor fi obținute în continuare și vor fi în final prezentate și discutate în raportul final al proiectului de cercetare-dezvoltare.

6.7.2 Decodarea clipurilor audio de evaluare folosind modele fonetice dependente de context cu mai puține densități Gaussiene per senone

Logați-vă pe serverul de dezvoltare, intrați în directorul proiectului, în subdirectorul `etc` și vizualizați fișierul de configurare al proiectului cu ajutorul editorului `vim`:

```
cd ~/projects/rodigits/etc/
vi sphinx_train.cfg
```

Identificați în acest fișier linia care specifică modelele acustice utilizate pentru decodare:

```
# Models to use.
#$DEC_CFG_MODEL_NAME =
"$CFG_EXPTNAME.cd_${CFG_DIRLABEL}_${CFG_N_TIED_STATES}";
$DEC_CFG_MODEL_NAME = "$CFG_EXPTNAME.ci_${CFG_DIRLABEL}";
```

Intrați în modul de editare (tastând “i”), ștergeți linia decommentată (cea care specifică locația modelelor fonetice independente de context), apoi copiați linia comentată (pentru a păstra un back-up) și decommentați una dintre copii. Modificați copia decommentată astfel încât să specificați locația modelelor fonetice dependente de context ce folosesc o singură densitate Gaussiană per senone. Această zonă a fișierului ar trebui să arate în final ca mai jos:

```
# Models to use.
#$DEC_CFG_MODEL_NAME =
"$CFG_EXPTNAME.cd_${CFG_DIRLABEL}_${CFG_N_TIED_STATES}";
$DEC_CFG_MODEL_NAME =
"$CFG_EXPTNAME.cd_${CFG_DIRLABEL}_${CFG_N_TIED_STATES}_1";
```

Identificați apoi linia care specifică numele subdirectorului în care se vor salva rezultatele:

```
$DEC_CFG_RESULT_DIR = "$CFG_BASE_DIR/result.cd_cont_200_8";
```

Și modificați-o astfel încât numele subdirectorului în care se vor salva rezultatele să fie sugestiv:

```
$DEC_CFG_RESULT_DIR = "$CFG_BASE_DIR/result.cd_cont_200_1";
```

Ieșiți din modul de editare (tastând `Esc`), salvați modificările (tastând “:w” și apoi `Enter`) și apoi închideți editorul (tastând “:q” și apoi `Enter`).

Copiați definiția modelului acustic din subdirectorul ce conține modelele fonetice dependente de context ce folosesc 8 densități Gaussianne per senone (`rodigits.cd_cont_200`) în subdirectorul ce conține modelele fonetice dependente de context ce folosesc o singură densitate Gaussiană per senone (`rodigits.cd_cont_200_1`):

```
cd ~/projects/rodigits
cp model_parameters/rodigits.cd_cont_200/mdef
model_parameters/rodigits.cd_cont_200_1/
```

Porniți apoi procesul de evaluare executând următoarea comandă:

```
/usr/local/sphinx/lib/sphinxtrain/scripts/decode/slave.pl
```

Așteptați terminarea procesului de decodare (ar trebui să dureze aproximativ 1 minut). În cazul în care apar erori încercați să identificați problema vizualizând fișierul de log al procesului de decodare:

```
vi logdir/decode/rodigits-1-1.log
```

În cazul în care nu reușiți să identificați și să rezolvați problema discutați cu îndrumătorul de proiect. În cazul în care procesul de decodare se încheie cu succes ar trebui ca pe ecran să vă apară un mesaj similar cu acesta:

```
MODULE: DECODE Decoding using models previously trained (decode script:
psdecode_fsg.pl)
  Decoding 20 segments starting at 0 (part 1 of 1)
  0%
  Aligning results to find error rate
  SENTENCE ERROR: 40.0% (8/20)  WORD ERROR RATE: 9.6% (23/240)
```

Aceste rezultate, precum și informațiile și rezultatele detaliate din raportul de evaluare generat la acest pas (`result.cd_cont_200_1/rodigits.align`) vor fi comparate cu rezultatele obținute în prealabil și cu rezultatele ce vor fi obținute în continuare și vor fi în final prezentate și discutate în raportul final al proiectului de cercetare-dezvoltare.

Repeți pașii de mai sus realizând decodarea clipurilor audio de evaluare cu modele fonetice dependente de context cu 2 și respectiv 4 densități Gaussianne per senone.

6.7.3 Antrenarea unui model acustic având în total 100 de senone

În continuare ne propunem să verificăm dacă numărul de 200 de senone (configurat implicit într-un proiect CMU Sphinx) este optim sau nu pentru sarcina de recunoaștere a vorbirii propusă în acest îndrumar (recunoașterea secvențelor audio ce conțin cifre).

Vom începe investigația antrenând un nou model acustic având în total 100 de senone disponibile pentru modelarea fonemelor ce formează cuvintele specifice sarcinii de recunoaștere. Pentru a realiza acest lucru ar trebui urmăriți în principiu pașii descriși în

capitolul 6.4. Însă, pentru această sarcină nu vom crea un nou proiect CMU Sphinx, în consecință nu va trebui nici să adăugăm alte fișiere de resurse și nici să transformăm fișierele cu resurse în formatul cerut de CMU Sphinx. Vom modifica numai fișierul de configurare al proiectului `rodigits` astfel încât numărul de senone să fie 100 (nu 200, așa cum era în mod implicit):

```
cd ~/projects/rodigits/etc/  
vi sphinx_train.cfg
```

Identificați în acest fișier linia care specifică numărul de senone:

```
# Number of tied states (senones) to create in decision-tree clustering  
$CFG_N_TIED_STATES = 200;
```

Intrați în modul de editare (tastând “i”) și modificați-o astfel:

```
# Number of tied states (senones) to create in decision-tree clustering  
$CFG_N_TIED_STATES = 100;
```

Ieșiți din modul de editare (tastând Esc), salvați modificările (tastând “:w” și apoi Enter) și apoi închideți editorul (tastând “:q” și apoi Enter).

Intrați în directorul proiectului și porniți antrenarea modelului acustic executând următoarea comandă:

```
cd ~/projects/rodigits  
sphinxtrain_biosinf run
```

Așteptați terminarea procesului de antrenare (ar trebui să dureze aproximativ 2 minute). În cazul în care apar erori verificați faptul că fișierele cu resurse sunt în formatul corespunzător și că nu ați modificat greșit fișierul de configurare. În cazul în care nu reușiți să identificați și să rezolvați problema discutați cu îndrumătorul de proiect.

După încheierea cu succes a procesului de antrenare verificați faptul că au fost generate modelele acustice corespunzătoare listând subdirectoarele din directorul proiectului, subdirectorul `model_parameters`:

```
cd ~/projects/rodigits  
ls -all model_parameters
```

Lista subdirectoarelor afișate ar trebui să fie următoarea:


```

rodigits.cd_cont_100
rodigits.cd_cont_100_1
rodigits.cd_cont_100_2
rodigits.cd_cont_100_4
rodigits.cd_cont_200
rodigits.cd_cont_200_1
rodigits.cd_cont_200_2
rodigits.cd_cont_200_4
rodigits.cd_cont_initial
rodigits.cd_cont_untied
rodigits.ci_cont
rodigits.ci_cont_flatinitial

```

Observați că, pe lângă subdirectoarele existente anterior (cele care conțin modele acustice cu 200 de senone și număr variabil de densități Gaussiene), în această listă apar acum și subdirectoarele `rodigits.cd_cont_100*` ce conțin modele acustice cu 100 de senone și număr variabil de densități Gaussiene.

6.7.4 Decodarea clipurilor audio de evaluare folosind modele acustice cu 100 de senone

Urmează ca acum să evaluăm aceste 4 noi modele acustice. Vom începe cu modelul acustic cu 100 de senone și 8 densități Gaussiene per senone. Logați-vă pe serverul de dezvoltare, intrați în directorul proiectului, în subdirectorul `etc` și vizualizați fișierul de configurare al proiectului cu ajutorul editorului vim:

```

cd ~/projects/rodigits/etc/
vi sphinx_train.cfg

```

Identificați în acest fișier linia care specifică modelele acustice utilizate pentru decodare:

```

# Models to use.
#$DEC_CFG_MODEL_NAME =
"$CFG_EXPTNAME.cd_${CFG_DIRLABEL}_${CFG_N_TIED_STATES}";
$DEC_CFG_MODEL_NAME =
"$CFG_EXPTNAME.cd_${CFG_DIRLABEL}_${CFG_N_TIED_STATES}_4";

```

Intrați în modul de editare (tastând “i”), ștergeți linia decommentată (cea care specifică locația modelelor fonetice dependente de context cu 4 densități Gaussiene), apoi decommentați linia comentată (ce specifică locația modelelor fonetice implicite). Această zonă a fișierului ar trebui să arate în final ca mai jos:

```

# Models to use.
$DEC_CFG_MODEL_NAME =
"$CFG_EXPTNAME.cd_${CFG_DIRLABEL}_${CFG_N_TIED_STATES}";

```

Identificați apoi linia care specifică numele subdirectorului în care se vor salva rezultatele:

```
$DEC_CFG_RESULT_DIR = "$CFG_BASE_DIR/result.cd_cont_200_4";
```

Și modificați-o astfel încât numele subdirectorului în care se vor salva rezultatele să fie sugestiv:

```
$DEC_CFG_RESULT_DIR = "$CFG_BASE_DIR/result.cd_cont_100_8";
```

Ieșiți din modul de editare (tastând Esc), salvați modificările (tastând “:w” și apoi Enter) și apoi închideți editorul (tastând “:q” și apoi Enter).

Intrați apoi în directorul proiectului și porniți procesul de evaluare executând următoarea comandă:

```
cd ~/projects/rodigits  
/usr/local/sphinx/lib/sphinxtrain/scripts/decode/slave.pl
```

Așteptați terminarea procesului de decodare (ar trebui să dureze aproximativ 1 minut). În cazul în care apar erori încercați să identificați problema vizualizând fișierul de log al procesului de decodare:

```
vi logdir/decode/rodigits-1-1.log
```

În cazul în care nu reușiți să identificați și să rezolvați problema discutați cu îndrumătorul de proiect. În cazul în care procesul de decodare se încheie cu succes ar trebui ca pe ecran să vă apară un mesaj similar cu acesta:

```
MODULE: DECODE Decoding using models previously trained (decode script:  
psdecode_fsg.pl)  
  Decoding 20 segments starting at 0 (part 1 of 1)  
  0%  
  Aligning results to find error rate  
  SENTENCE ERROR: 5.0% (1/20)  WORD ERROR RATE: 0.4% (0/240)
```

Aceste rezultate, precum și informațiile și rezultatele detaliate din raportul de evaluare generat la acest pas (`result.cd_cont_100_8/rodigits.align`) vor fi comparate cu rezultatele obținute în prealabil și cu rezultatele ce vor fi obținute în continuare și vor fi în final prezentate și discutate în raportul final al proiectului de cercetare-dezvoltare.

Repetăți pașii de mai sus (ținând cont și de capitolul care prezintă decodarea folosind modele cu mai puține densități Gaussiene per senone) realizând decodarea clipurilor audio de evaluare cu modele fonetice dependente de context cu 100 de senone și 1, 2 și respectiv 4 densități Gaussiene per senone.

6.8 Construcția unui sistem de recunoaștere independent de vorbitor

Acest capitol vă propune realizarea unui al doilea proiect de recunoaștere a secvențelor audio ce conțin cifre, de această dată *independent de vorbitor* (îl vom numi `rodigitsIndep`). Pentru a justifica necesitatea unui sistem independent de vorbitor ar trebui să evaluați sistemul `rodigits` creat anterior pe fișierele de evaluare ale celorlalți colegi de grupă (fișiere ce se află pe serverul de dezvoltare în directorul `/home/biosinfShare/resources/wav/`). În urma evaluării veți constata că rezultatele de recunoaștere a vorbirii sunt mai slabe pentru ceilalți vorbitori decât pentru voi înșivă. Explicația este simplă: modelul acustic din proiectul `rodigits` a fost antrenat doar cu vocea voastră și, în consecință, vă recunoaște bine numai pe voi!

Crearea unui proiect independent de vorbitor presupune reluarea tuturor etapelor prezentate în acest îndrumar, începând cu capitolul 6.4, cu o singură modificare: clipurile audio de antrenare și evaluare nu vor mai fi doar clipurile înregistrate de voi, ci toate clipurile înregistrate de toți colegii de grupă. Bineînțeles că această modificare poate duce și la alte modificări (de exemplu adaptarea scripturilor `createTranscriptions.sh` și `createFileIds.sh`), însă pe toate acestea le veți aborda și eventual discuta cu îndrumătorul de proiect la momentul respectiv.

6.9 Redactarea raportului final al proiectului de cercetare-dezvoltare

După efectuarea experimentelor prezentate în acest îndrumar veți sumariza toate rezultatele obținute și veți face o discuție asupra lor într-un raport final aferent proiectului de cercetare-dezvoltare.

Raportul de cercetare-dezvoltare trebuie să înceapă cu o scurtă secțiune care să cuprindă noțiuni introductive cu privire la arhitectura sistemelor de recunoaștere a vorbirii, sub-sistemele implicate, etc. Această primă secțiune are rolul de a evidenția nivelul de înțelegere (asupra sistemelor de recunoaștere a vorbirii) pe care l-ați dobândit pe parcursul realizării proiectului.

O a doua secțiune a raportului de cercetare-dezvoltare va prezenta pe scurt cerința (tema) proiectului și va enumera pașii ce au fost urmați pentru a dezvolta sistemul de recunoaștere a secvențelor audio ce conțin cifre.

În final, partea cea mai consistentă a raportului trebuie să prezinte **toate** rezultatele obținute și să le discute în mod comparativ. În funcție de stadiul în care ați ajuns cu experimentele veți avea rezultate de recunoaștere pentru:

- modelul acustic dependent de vorbitor:
 - cu 200 de senone:
 - independente de context
 - dependente de context, cu 1/2/4/8 densități Gaussiene per senone
 - cu 100 de senone:
 - independente de context
 - dependente de context, cu 1/2/4/8 densități Gaussiene per senone
- modelul acustic independent de vorbitor (cu aceleași variații ale parametrilor)

Prezentarea individuală a rezultatelor (pentru fiecare experiment în parte) va trebui să se facă punând accentul pe rata de eroare la nivel de cuvânt (Percent Total Error = Word Error Rate), dar se pot discuta și alte metrici adiacente cum sunt rata de cuvinte inserate (Percent Insertions), rata de cuvinte substituite (Percent Substitutions), rata de cuvinte șterse (Percent Deletions), rata de eroare la nivel de propoziție (Sentences with errors), etc. De asemenea, în funcție de rezultatele obținute, pot fi prezentate perechile de cuvinte confundate foarte frecvent, cuvintele inserate/substituite/șterse cel mai frecvent și eventual 1-2 exemple de fraze (evidențiind varianta corectă a frazei și varianta obținută la ieșirea sistemului).

Compararea rezultatelor obținute în diversele experimente se va face a) numai în funcție de Word Error Rate (WER), dacă ați efectuat multe experimente sau b) în funcție de mai multe metrici, dacă aveți puține rezultate. În primul caz rezultatele vor fi prezentate folosind tabele și grafice de variație a ratei de eroare la nivel de cuvânt (WER). Rezultatele se vor discuta pe baza acestor tabele și grafice și, în final, se vor trage concluzii cu privire la dependența/independența de context, numărul optim de senone, numărul optim de densități Gaussiene, etc. pentru modelul acustic dependent de vorbitor, respectiv independent de vorbitor.

Notă: redactarea raportului de cercetare-dezvoltare este obligatorie, iar transmiterea lui (prin e-mail) către îndrumătorul de proiect trebuie să se facă până în săptămâna a 12-a a semestrului.

Notă: raportul final trebuie să fie redactat și transmis îndrumătorului de proiect indiferent de stadiul în care ați ajuns cu experimentele.

Notă: toate rezultatele prezentate în raportul de cercetare-dezvoltare vor fi însoțite de referințe către rapoartele de evaluare de pe serverul de dezvoltare (calea către fișierele corespunzătoare).

6.10 Desfășurarea activității pe server-ul de dezvoltare

Proiectul de recunoaștere a secvențelor audio ce conțin cifre va fi realizat în întregime pe un server de dezvoltare aflat în cadrul laboratorului de cercetare Speed (<http://speed.pub.ro>). Studentul va trebui să dispună de un calculator conectat la Internet care va fi folosit pentru conectarea de la distanță (prin intermediul Internetului) la serverul de dezvoltare.

Serverul de dezvoltare utilizează un sistem de operare de tip unix (Ubuntu 12.04). În consecință studentul trebuie să fie familiarizat sau să se familiarizeze cu modul de operare pe un astfel de sistem. În continuare sunt date câteva informații despre modul de conectare la server și despre utilizarea editorului vim.

Conectarea la serverul de dezvoltare

Conectarea la serverul de dezvoltare se face extrem de simplu în cazul în care calculatorul de care dispuneți utilizează un sistem de operare de tip unix/linux. În acest caz conectarea la server se face utilizând un terminal în care executați următoarea comandă (important: modificați comanda următoare pentru a utiliza:

- propriul nume de utilizator `USER_NAME`
- portul `PORT` și numele serverul `SERVER_NAME`, comunicate de îndrumătorul de proiect.

```
ssh -p PORT USER_NAME@SERVER_NAME
```

Serverul va raspunde cerându-vă să tastați parola contului vostru. Pentru a continua va trebui să tastați parola contului vostru, așa cum v-a fost specificată îndrumătorul de proiect. În cazul conectării cu succes pe ecran vă va apărea următorul prompt (unde veți executa comenzile precizate în acest îndrumar):

```
USER_NAME@alpha:~$
```

În cazul în care calculatorul de care dispuneți utilizează un sistemul de operare Windows atunci conectarea la serverul de dezvoltare se va face cu ajutorul unei aplicații intermediare numită *putty*. Această aplicație poate fi descărcată gratuit de pe siteul: <http://putty.en.softonic.com/>. După descărcare, porniți aplicația și completați în primul ecran *Host Name:* `SERVER_NAME` și *Port:* `PORT`. Atenție: `SERVER_NAME` și `PORT` sunt date ce v-au fost comunicate de îndrumătorul de proiect. Pentru a salva aceste informații și a nu fi nevoiți să le completați de fiecare dată completați *Saved Sessions:* `biosinfServer` și apăsați butonul *Save*. Pentru ca această aplicație să poată afișa corect caracterele diacritice, intrați în meniul *Window*, submeniul *Translation* și selectați *Remote Character Set:* `UTF-8`. Apoi reveniți la ecranul inițial, selectând meniul *Session* și apăsați butonul *Open* pentru a stabili conexiunea. Serverul vă va cere numele de utilizator și apoi parola (informații pe care vi le-a

specificat îndrumătorul de proiect). În cazul conectării cu succes pe ecran vă va apărea următorul prompt (unde veți executa comenzile precizate în acest îndrumar):

```
USER_NAME@alpha:~$
```

Recomandare: primul lucru pe care ar trebui să îl faceți după ce vă conectați pe server pentru prima dată este să vă schimbați parola. Pentru aceasta executați următoarea comandă și introduceți, pe rând, parola actuală, noua parola și încă o dată noua parolă:

```
passwd
```

Notă: indiferent de modul de conectare la serverul de dezvoltare (terminal linux sau putty) aveți posibilitatea și de multe ori este chiar foarte util să deschideți mai multe conexiuni simultane către server.

Utilizarea editorului în mod text vim

În cadrul acestui proiect veți lucra extensiv (creare, vizualizare, editare, modificare, etc.) cu fișiere text. Cel mai simplu mod de a manipula fișierele text aflate pe un server linux este cu ajutorul editorului vim.

Pentru a vizualiza un fișier deja existent sau a crea un fișier nou (cu numele filename) tastați comanda:

```
vi filename
```

În acest moment vizualizați prima parte a fișierului, iar cu ajutorul tastelor cu săgeți și a tastelor PgUp/PgDn puteți vizualiza și alte zone ale fișierului. Pentru a căuta un cuvânt (sau orice alt șir de caractere) tastați “/”, apoi cuvântul respectiv și în final Enter. Pentru a căuta același cuvânt din nou tastați “/” și apoi Enter. Pentru a intra în modul de editare (în vederea modificării fișierului) tastați “i”. Acum puteți șterge și scrie în acest fișier ca în oricare alt editor de text. Pentru a ieși din modul de editare tastați Esc. Pentru a salva modificările făcute tastați “:w” și apoi Enter. Pentru a părăsi editorul tastați “:q” și apoi Enter. Pentru mai multe informații de utilizare și alte comenzi de editare, copiere, etc. consultați site-ul vim precizat în bibliografia îndrumarului de proiect [Vi].

BIBLIOGRAFIE

[Bahl, 1991] Bahl, L.R., et al., "Multitonic Markov Word Models for Large Vocabulary Continuous Speech Recognition," *IEEE Transactions on Speech and Audio Processing*, 1993, vol. 1, no. 3, pp. 334-344, 1991.

[Baker, 1975] Baker, J., "The DRAGON system – an overview," *IEEE Transactions on Acoustics, Speech, and Signal Processing*, vol. 23, no. 1, pp. 24-29, 1975.

[Baum, 1972] Baum, L.E., "An inequality and associated maximization technique in statistical estimation for probabilistic functions of Markov processes," *Inequalities*, vol. 3, no. 1, pp. 1-8, 1972.

[Buzo, 2011] Buzo, A., "Automatic Speech Recognition over Mobile Communication Networks," Teză de doctorat, Universitatea Politehnica din București, România, 2011, <http://speed.pub.ro/academic/research-and-development-project-in-spoken-language-technology>.

[Cucu, 2011] Cucu, H., "Towards a speaker-independent, large-vocabulary continuous speech recognition system for Romanian," Teză de doctorat, Universitatea Politehnica din București, România, 2011, <http://speed.pub.ro/academic/research-and-development-project-in-spoken-language-technology>.

[Dempster, 1977] Dempster, A., Laird, N., Rubin, D., "Maximum likelihood from incomplete data via the EM algorithm," *Journal of the Royal Statistical Society, Series B*, vol. 39, pp. 1-38, 1977.

[Huang, 2001] Huang, X., Acero, A., Wuen-Hon, H., *Spoken Language Processing – A Guide to Theory, Algorithm, and System Development*, Prentice Hall, 2001.

[Hwang, 1991] Hwang, M.Y., X.D. Huang, "Acoustic Classification of Phonetic Hidden Markov Models," *Eurospeech 1991*, 1991.

[Hwang, 1993] Hwang, M.Y., Huang, X., Alleva, F., "Predicting Unseen Triphones with Senones," *ICASSP 1993*, Minneapolis, pp. 311-314, 1993.

[JSGF] Java Speech Grammar Format, <http://java.sun.com/products/java-media/speech/forDevelopers/JSGF/index.html>.

[Jelinek, 1998] Jelinek, F., *Statistical Methods for Speech Recognition*, The MIT Press, Cambridge, MA, 1998.

[Jurafsky, 2009] Jurafsky, D., Martin, J., "Automatic Speech Recognition," Chapter 9 in *Speech and Language Processing: An introduction to natural language processing, computational linguistics, and speech recognition (2nd Ed.)*, Pearson Education, 2009.

[Poritz, 1988] Poritz, A., "Hidden Markov models: a guided tour," *ICASSP 1988*, pp. 7–13.

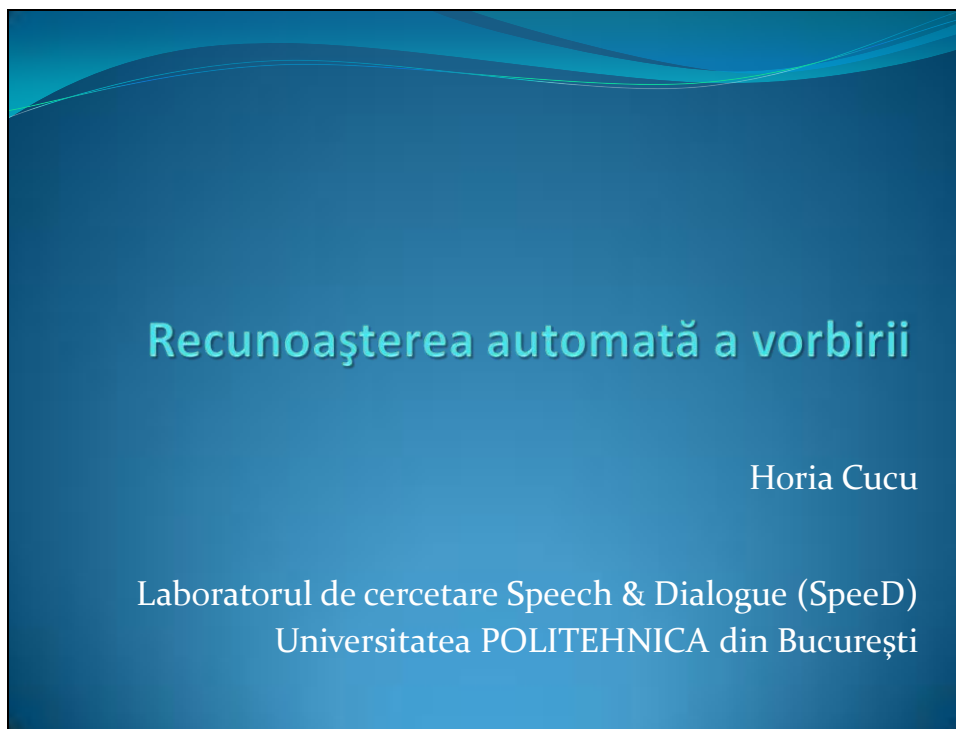
[Rabiner, 1989] Rabiner, L.R., "A tutorial on hidden Markov models and selected applications in speech recognition," *Proceedings of the IEEE*, vol. 77, no. 2, pp. 257-286, 1989.

[Renalds, 2010] Renalds S., Hain, T., "Speech Recognition," *Handbook of Computational Linguistics and Natural Language Processing*, 2010.

[Sphinx] Toolkitul CMU Sphinx și tutorialele asociate, <http://cmusphinx.sourceforge.net/wiki/>

[Vi] Manualul editorului de text vi, <http://www.cs.fsu.edu/general/vimanual.html>.

ANEXA 1. DIAPOZITIVE DE PREZENTARE

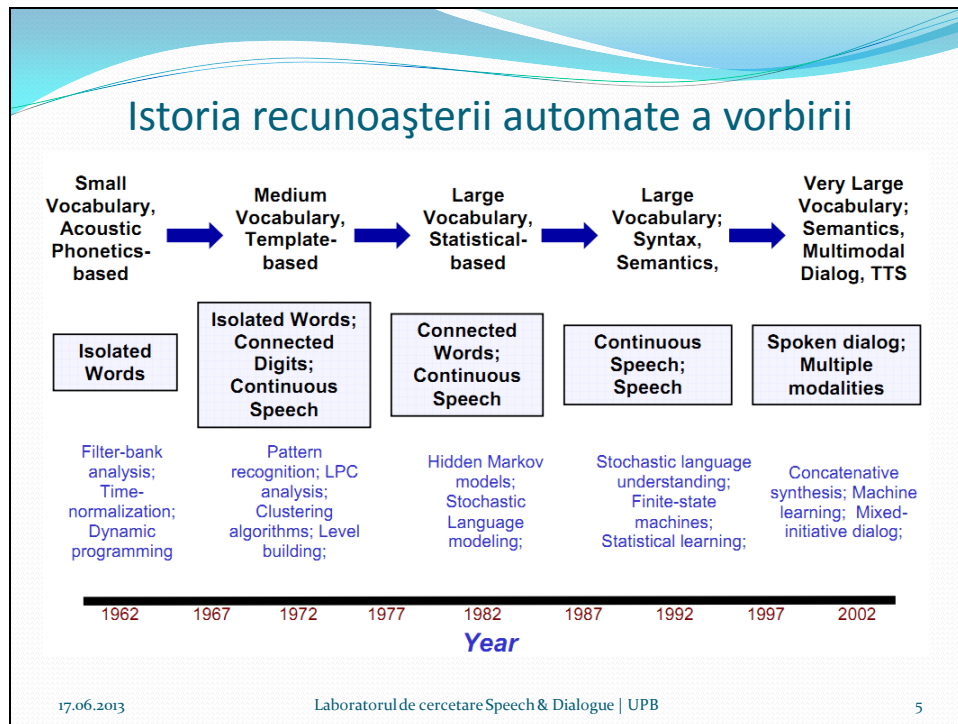


Cuprins

- Recunoașterea automată a vorbirii (RAV): introducere
 - Introducere în recunoașterea automată a vorbirii
 - Aplicații ale recunoașterii automate a vorbirii
- Recunoașterea automată a vorbirii: principii generale
- Resurse fonetice, acustice și lingvistice
- Modelarea acustică în sisteme de RAV
- Modelarea limbajului în sisteme de RAV
- Concluzii, discuții și demonstrație practică

Introducere în recunoașterea automată a vorbirii

- Recunoașterea automată a vorbirii (RAV) -> procesul de transformare în text a unui semnal audio ce conține de vorbire
- Face parte din domeniul mai larg de recunoașterea formelor
- Abordarea din starea artei este în mod principal statistică



Aplicații de transcriere speech-to-text (I)

- Completarea documentelor structurate (rapoarte medicale)
 - Sarcină de transcriere simplă (cuvinte izolate și vorbire continuă)
 - Un singur utilizator per aplicație (se poate face adaptare la vorbitor)
 - Utilizatorul își dorește să fie recunoscut de aplicație!
 - Limbajul este restricționat (depinde puternic de domeniul aplicației)
- Transcrierea vorbirii continue citite (știri, conferințe de presă)
 - Recunoașterea vorbirii continue citite
 - Mai mulți vorbitori per aplicație
 - Limbajul nu este restricționat
 - Vocabularul de cuvinte conține practic toate cuvintele limbii

17.06.2013 Laboratorul de cercetare Speech & Dialogue | UPB 6

Aplicații de transcriere speech-to-text (II)

- Dictarea documentelor (SMS-uri, email-uri, documente word)
 - Recunoaștere de vorbire continuă spontană
 - Un singur utilizator per aplicație (se poate face adaptare la vorbitor)
 - Utilizatorul își dorește să fie recunoscut!
 - Limbajul și vocabularul nu sunt restricționate (toate cuvintele limbii)
- Transcrierea vorbirii continue spontane (interviuri, talk-shows)
 - Cea mai dificilă sarcină de transcriere (vorbire continuă spontană)
 - Mai mulți vorbitori per aplicație
 - Limbajul și vocabularul nu sunt restricționate (toate cuvintele limbii)
 - Dificultăți suplimentare: calitatea înregistrărilor, zgomotul de fundal

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

7

Aplicații de căutare și indexare

- Indexarea materialelor audio ce conțin vorbire
 - Nu este necesar ca recunoașterea să fie perfectă
 - Dificultăți: procesarea unei cantități mari de date
- Identificarea cuvintelor cheie (Keyword spotting)
 - Query-ul (cuvântul cheie) este introdus sub formă de text
 - Căutarea se face în materiale audio indexate sau direct în transcriere
- Detectarea termenilor vorbiți (Spoken term detection)
 - Query-ul este vorbit! Conținutul este material vorbit!
 - Interpretarea termenului (transcriere sau clasificare) este esențială

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

8

Aplicații de comandă și control

- Controlul telefonului/calculatorului/televizorului/casei inteligente
 - Sarcină foarte simplă: recunoașterea câtorva comenzi fixe
 - Singura dificultate: separarea comenzilor de vorbirea uzuală
 - Aplicație adaptată vorbitorului/vorbitorilor (3-4)
 - Utilizatorul își dorește să fie recunoscut de aplicație!
- Controlul avioanelor și al elicopterelor
 - Condiții extreme de zgomot!
 - Respirațiile vorbitorului trebuie tratate special
 - Precizia recunoașterii este esențială

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

9

Aplicații de dialog om-mașină

- Presupun trei blocuri de procesare:
 - Recunoașterea vorbirii continue
 - Înțelegerea vorbirii și generarea răspunsurilor adecvate
 - Sinteza vorbirii
- Dialog în call-centers (rezervări de bilete, suport tehnic)
 - Ghidarea inteligentă a utilizatorului prin meniuri ajută procesele de recunoaștere și înțelegere a vorbirii
- Dialog cu smart-phones (informații despre vreme, locații, etc.)
 - Puterea de procesare determină viteza și performanța aplicației
 - De cele mai multe ori procesarea se face în cloud

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

10

Aplicații speech-to-speech

- Presupun trei blocuri de procesare:
 - Recunoașterea vorbirii continue în limba sursă
 - Traducerea automată din limba sursă în limba destinație
 - Sinteza vorbirii în limba destinație
- Toate cele trei blocuri introduc erori
- Status:
 - Demonstrații în cadrul manifestărilor științifice
 - Nu există încă aplicații comerciale

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

11

Cuprins

- Recunoașterea automată a vorbirii (RAV): introducere
- Recunoașterea automată a vorbirii: principii generale
 - Surse de variabilitate în recunoașterea automată a vorbirii
 - Arhitectura unui sistem de RAV
 - Performanțe uzuale ale sistemelor de RAV
- Resurse fonetice, acustice și lingvistice
- Modelarea acustică în sisteme de RAV
- Modelarea limbajului în sisteme de RAV
- Concluzii, discuții și demonstrație practică

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

12

Surse de variabilitate în RAV

- Domeniul/Sarcina de recunoaștere
 - Limba și dimensiunea vocabularului
 - Incertitudinea lingvistică a discursului
- Caracteristicile vorbitorului
 - Accent, dialect, nativ/non-nativ, rapiditatea pronunției, vârsta
 - Sisteme de recunoaștere dependente/independente de vorbitor
- Stilul vorbirii
 - Cuvinte izolate, vorbire continuă, vorbire spontană (conversațională)
- Mediul acustic
 - Alți vorbitori, zgomot de fundal, caracteristicile microfonului

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

13

Formalismul actual al RAV

- Care este cea mai probabilă secvență de cuvinte W^* dat fiind mesajul vorbit X ?

$$W^* = \arg \max_W p(W | X) = \arg \max_W \frac{p(X | W)p(W)}{p(X)} = \arg \max_W p(X | W)p(W)$$

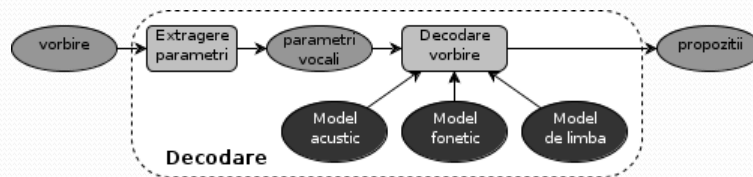
- W – toate secvențele posibile de cuvinte
- $p(X|W)$ – probabilitatea mesajului vorbit, dată fiind secvența de cuvinte W ; este estimată utilizând un model acustic!
- $p(W)$ – probabilitatea de apariție a secvenței de cuvinte W ; este estimată strict pe baza unui model de limbă!

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

14

Arhitectura unui sistem de recunoaștere a vorbirii



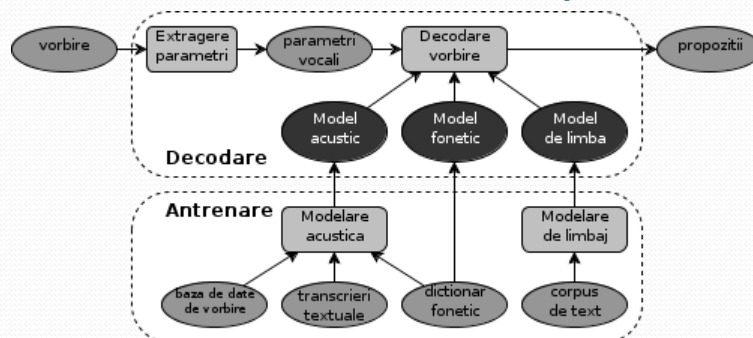
- Modelarea acustică (state-of-the-art):
 - Sisteme statistice de tip HMM-GMM
 - Unitățile de vorbire de bază: foneme și unități sub-fonetice (senone) dependente de context
 - Parametri acustici perceptuali (MFCC, PLP, etc.)
 - Toolkit-uri: HTK, CMU Sphinx, etc.
- Modelarea lingvistică (state-of-the-art):
 - Sisteme statistice de tip n-gram
 - Metode de netezire: Good-Turing, Witten-Bell, Kneser-Ney, etc.
 - Rolul modelului de limbă: “casa ta e mare” sau “casat a Ema re”?
 - Toolkit-uri: SRI-LM, CMU-SLM, etc.

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

15

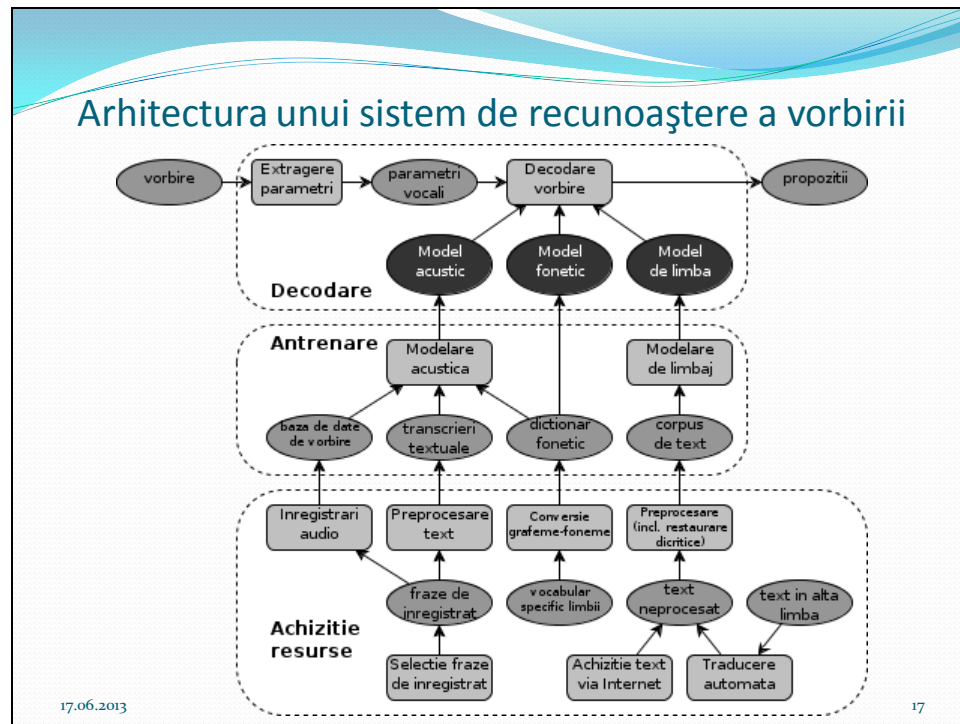
Arhitectura unui sistem de recunoaștere a vorbirii



17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

16



Evaluarea sistemelor de recunoaștere a vorbirii

- Compararea transcrierii de referință cu transcrierea ipotetică produsă de sistemul de RAV
- Rata de eroare la nivel de cuvânt (WER) – echivalentul unui cost de editare;
 - Alinierea celor două secvențe de cuvinte (DTW)
 - Calcularea costului pe baza numărului de erori de inserție, ștergere, substituție:
- Rata de eroare la nivel de propoziție (SER)

$$WER\% = \frac{\#Erori\ de\ inserție + \#Erori\ de\ substituție + \#Erori\ de\ ștergere}{\#Cuvinte\ în\ transcrierea\ de\ referință} \times 100$$

$$SER\% = \frac{\#Propoziții\ eronate}{\#Propoziții\ în\ transcrierea\ de\ referință} \times 100$$

17.06.2013 18

Laboratorul de cercetare Speech & Dialogue | UPB

Exemple de evaluare WER pentru RAV

REF: iaurturile ajută la prevenirea **** * * * OSTEOPOROZEI

HYP: iaurturile ajută la prevenirea POST OP RO ZEI

Eval: I I I S

Counts: (#C #S #D #I) 4 1 0 3

WER: 4/5 = 80.0%

REF: am coborât din autocar ***** VAL VÂRTEJ și am început să FOTOGRAFIEZ

HYP: am coborât din autocar BARBU TEI și și am început să FOTOGRAFIEZE

Eval: I S S S

Counts: (#C #S #D #I) 8 3 0 1

WER: 4/11 = 36.3%

REF: aş vrea să aflu cumva că GESTUL LUI disperat nu ne-a fost indiferent

HYP: aş vrea să aflu cumva că ***** GESTULUI disperat nu ne-a fost indiferent

Eval: D S

Counts: (#C #S #D #I) 11 1 1 0

WER: 2/13 = 15.3%

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

19

Performanțele sistemelor de RAV actuale

- Sarcina de recunoaștere și dimensiunea vocabularului influențează semnificativ performanțele sistemului de RAV

Sarcina de RAV	Limba	Vocabular	WER
TI Digits	Engleză	11 words	0.55%
Wall Street Journal read speech		5k words	3.0%
Wall Street Journal read speech		20k words	<6.6%
Broadcast News		64k+ words	9.9%
Conversational Telephone Speech		64k+ words	20.7%

* rezultate raportate în 2005 pentru diverse sarcini de RAV în limba engleză

- Cea mai dificilă sarcină de RAV: large vocabulary conversational speech

Sistem dezvoltat de	Limba	Vocabular	WER
RWTH Aachen University	Engleză	150k words	17.9%
	Germană	300k words	17.3%
	Franceză	200k words	20.9%
Karlsruhe Institute of Technology	Rusă	500k words	19.3%

* rezultate raportate în 2011 pentru broadcast conversational speech în diverse limbi

17.06.2013

20

Cuprins

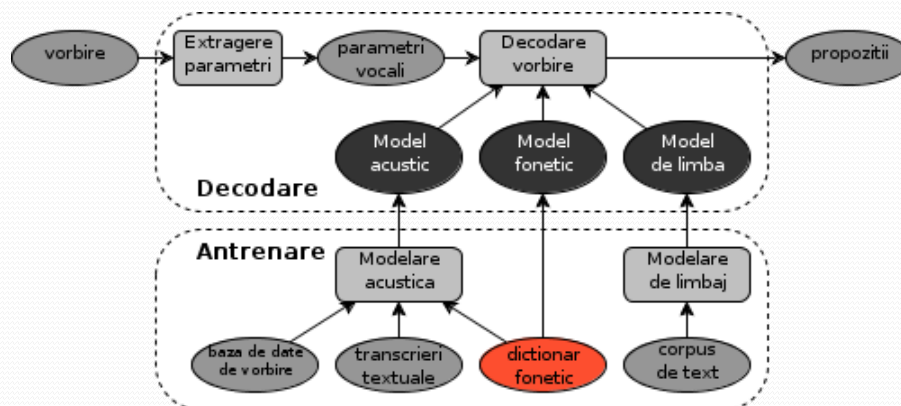
- Recunoașterea automată a vorbirii (RAV): introducere
- Recunoașterea automată a vorbirii: principii generale
- Resurse fonetice, acustice și lingvistice
 - Dicționare fonetice și fonetizarea automată
 - Baze de date de vorbire: metode de achiziție
 - Corpusuri de text: achiziție, normalizare și restaurare
- Modelarea acustică în sisteme de RAV
- Modelarea limbajului în sisteme de RAV
- Concluzii, discuții și demonstrație practică

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

21

Dicționarul fonetic într-un sistem de RAV



17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

22

Noțiuni generale de fonetică

- Fonemul – unitatea de sunet fundamentală din limba vorbită
 - Succesiune de foneme -> cuvinte
- Dicționarul fonetic: corespondență între forma scrisă și forma fonetică a cuvintelor
- Limba română folosește
 - 7 vocale de bază + 2 vocale împrumutate
 - 4 semivocale
 - 22 de consoane
 - Un simbol special pentru a arăta palatalizarea consoanei anterioare

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

23

Vocalele și semivocalele limbii române

Fonemul				Exemple de cuvinte	
Tip	Nr.	Simbol IPA	Simbol intern	Forma scrisă	Forma fonetică
vocale	1	a	a	sat	s a t
	2	e	e	mare	m a r e
	3	i	i	lift	l i f t
	4	o	o	loc	l o c
	5	u	u	șut	s 1 u t
	6	ə	a1	gură	g u r a1
	7	ɨ	i2	între	i2 n t r e
vocale împrumutate	8	y	y	ecru	e c r y
	9	ø	o2	bleu	b l o2
semivocale	10	ɛ̃	e1	deal	d e1 a l
	11	j	i3	fiară	f i3 a r a1
	12	ɔ̃	o1	oase	o1 a s e
	13	w	w	sau	s a w

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

24

Consoanele limbii române (I)

Fonemul			Exemple de cuvinte	
Nr.	Simbol IPA	Simbol intern	Forma scrisă	Forma fonetică
14	j	i1	tari	t a r i1
15	c	k2	chem	k2 e m
16	b	b	bar	b a r
17	p	p	par	p a r
18	k	k	acum	a k u m
19	tʃ	k1	cenușă	k1 e n u s1 a1
20	g	g	galben	g a l b e n
21	ɕ	g1	girafă	g1 i r a f a1
22	ʒ	g2	unghi	u n g2
23	d	d	dar	d a r
24	t	t	tot	t o t
25	f	f	fața	f a t1 a

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

25

Consoanele limbii române (II)

Fonemul			Exemple de cuvinte	
Nr.	Simbol IPA	Simbol intern	Forma scrisă	Forma fonetică
26	v	v	vapor	v a p o r
27	h	h	harta	h a r t a
28	ʒ	j	ajutor	a j u t o r
29	ʃ	s1	coș	k o s1
30	l	l	lac	l a c
31	m	m	măr	m a1 r
32	n	n	nas	n a s
33	s	s	sare	s a r e
34	z	z	zar	z a r
35	r	r	risc	r i s k
36	ʦ	t1	țaran	t1 a1 r a n

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

26

Construcția dicționarului fonetic

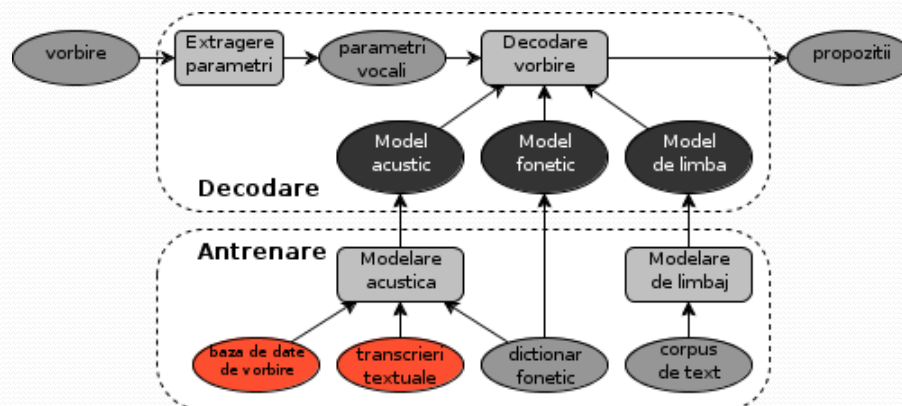
- Fonetizarea cuvintelor din vocabularul sistemului de RAV
 - Manuală, în cazul în care vocabularul este redus
 - Automată, pentru RAV cu vocabulare mari
- Fonetizarea automată: sisteme bazate pe reguli
 - Formalizarea regulilor și excepțiilor de fonetizare
 - Formalizarea regulilor și excepțiilor de despărțire în silabe
- Fonetizarea automată: sisteme de învățare automată
 - Bazate pe rețele neurale, traducere automată statistică, etc.
- Performanțe uzuale: WER între 3% și 5%

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

27

Baza de date de vorbire într-un sistem de RAV



17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

28

Caracteristicile bazei de date de vorbire

- Dimensiunea bazei de date:

- numărul de ore de vorbire, numărul de vorbitori

Sarcina RAV	Sistem dependent de vorbitor	Sistem independent de vorbitor
comandă și control (RVC-VR)	1 oră de înregistrări, 1 vorbitor	5 ore de înregistrări, 200 de vorbitori
dictare (RVC-VE)	10 ore de înregistrări, 1 vorbitor	50 ore de înregistrări, 200 de vorbitori

- Stilul vorbirii

- cuvinte izolate, vorbire continuă citită, vorbire conversațională

- Variabilitatea

- calitatea înregistrărilor, zgomotul de fundal, variabilitatea vorbitorilor

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

29

Metode de achiziție a bazei de date de vorbire

- Înregistrarea unor texte predefinite

- Control asupra textului, vorbitorilor, etc.
 - Vorbire continuă-citită
 - Eficiență dpdv al timpului de obținere
 - Erori de pronunție

- Etichetarea unor clipuri audio ce conțin vorbire

- Orice tip de vorbire
 - Erori de etichetare

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

30

Înregistrarea textelor predefinite

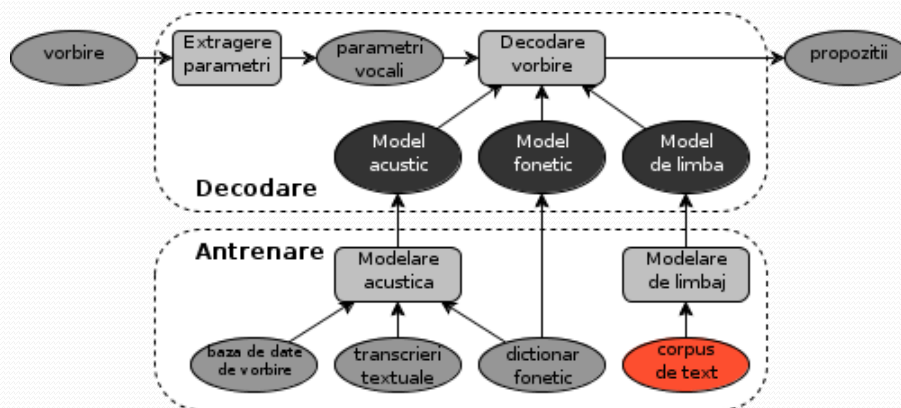
- Obținerea unei fluente în pronunție (fără ezitări/respirații)
 - fraze corecte dpdv lexical și gramatical, cu topică uzuală
 - fraze scurte (maxim 15-20 de cuvinte)
 - fraze conectate semantic
 - utilizarea cuvintelor frecvente în limbajul de zi cu zi
- Evitarea ambiguității pronunției
 - utilizarea cuvintelor frecvente în limbajul de zi cu zi
 - evitarea cuvintelor împrumutate din limbi străine
 - evitarea interjecțiilor
 - scrierea numerelor cu litere

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

31

Corpusul de text într-un sistem de RAV



17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

32

Utilizarea corpusului de text

- Sisteme RAV de comandă și control
 - Utilizează gramatici cu reguli
 - Construcția gramaticilor cu reguli nu necesită un corpus de text
- Sisteme RVC cu vocabular mare
 - Utilizează modele de limbă statistice
 - Un corpus de text este absolut necesar pentru extragerea de statistici
 - Calitatea modelului de limbă depinde de dimensiunea corpusului

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

33

Achiziția corpusului de text

- Surse de text utilizabil în modelarea lingvistică statistică
 - Cărți în format electronic
 - Știri de pe site-urile ziarelor on-line
 - Subtitrări de filme
 - Articole personale (blog-uri) on-line
- Dificultatea normalizării textelor depinde de sursa de achiziție
 - Fiecare carte are formatul ei specific (antet, cuprins, etc.)
 - Subtitrările conțin ștampile de timp (în diverse formate)
 - Site-urile de știri conțin meniuri, reclame, etc.

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

34

Normalizarea corpusului de text

- Eliminarea etichetelor HTML
- Expandarea abrevierilor
- Transcrierea numerelor scrise cu cifre în numere scrise cu litere
- Tratarea parantezelor și a caracterelor speciale
- Tratarea semnelor de punctuație
- Conversie uppercase-to-lowercase

Ieri, dl. Stan a obținut 6,5% din voturi.
Felicitări!

Ieri, domnul Stan a obținut 6,5% din
voturi. Felicitări!

Ieri, domnul Stan a obținut șase virgulă
cinci % din voturi. Felicitări!

Ieri, domnul Stan a obținut șase virgulă
cinci la sută din voturi. Felicitări!

ieri domnul stan a obținut șase virgulă
cinci la sută din voturi
felicitări

Restaurarea diacriticelor în corpusul de text

- 30 – 40% dintre cuvintele limbii române conțin diacritice!
 - Cuvinte neambigue dpdv al diacriticelor:
 - alb, astfel, fabricând, științific
 - Cuvinte ambigue:
 - sârmana/sărmană, pana/pană/până, tari / țări / țari / târî
- Textul de la ieșirea sistemului de RAV este de obicei scurt și eliptic, iar lipsa diacriticelor îl poate face de neînțeles!
- Restaurarea diacriticelor este obligatorie pentru textele colectate de pe Internet
- Restaurarea diacriticelor se face automat, prin metode statistice

diacriticele limbii române: ă, â, î, ș, ț
sunt câteodată înlocuite cu: a, a, i, s, t

Cuprins

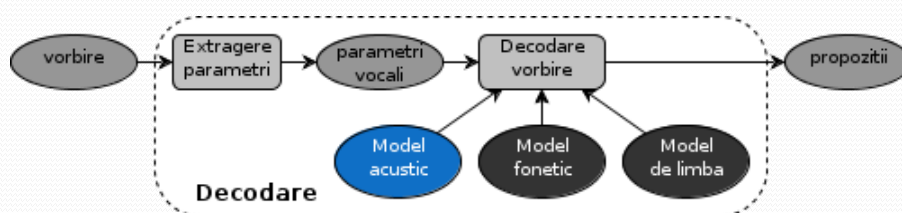
- Recunoașterea automată a vorbirii (RAV): introducere
- Recunoașterea automată a vorbirii: principii generale
- Resurse fonetice, acustice și lingvistice
- Modelarea acustică în sisteme de RAV
 - Semnalul de vorbire
 - Platforma HMM-GMM
 - Algoritmi de antrenare și decodare
- Modelarea limbajului în sisteme de RAV
- Concluzii, discuții și demonstrație practică

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

37

Rolul modelului acustic într-un sistem de RAV



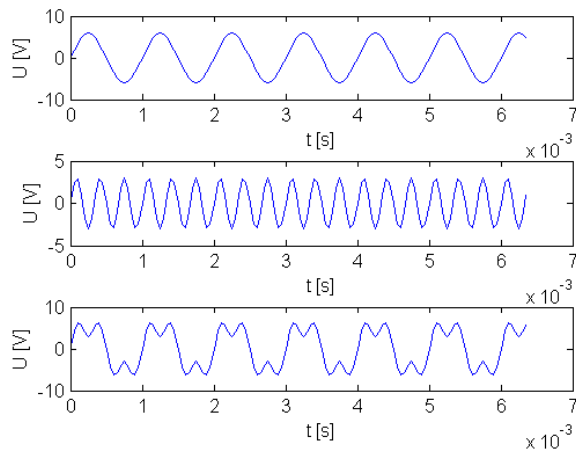
$$W^* = \arg \max_W p(W | X) = \arg \max_W \frac{p(X | W)p(W)}{p(X)} = \arg \max_W p(X | W)p(W)$$

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

38

Semnale deterministe

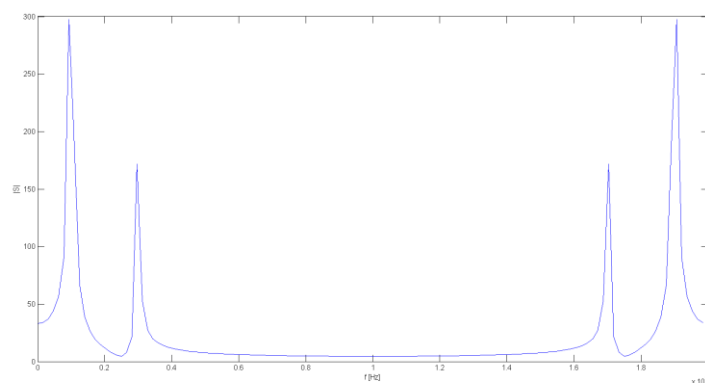


17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

39

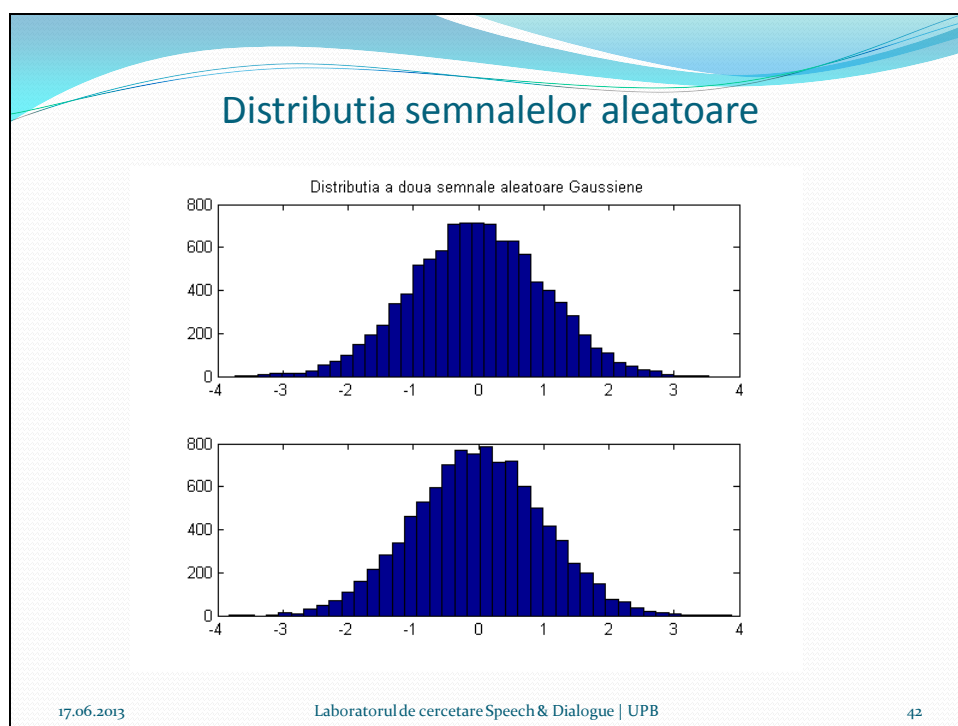
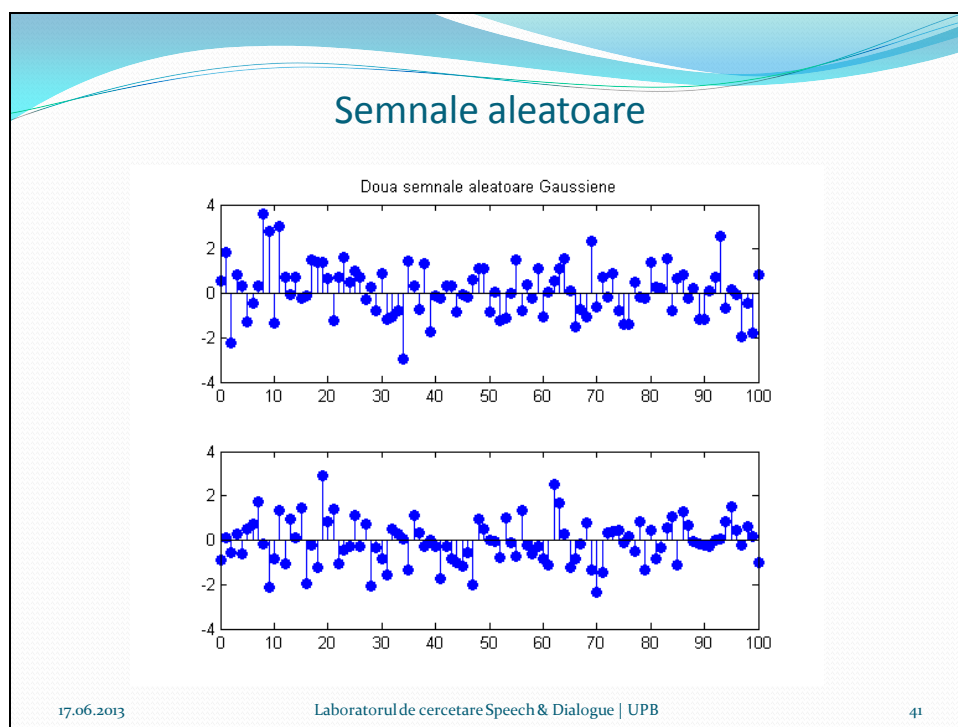
Spectrul semnalului determinist



17.06.2013

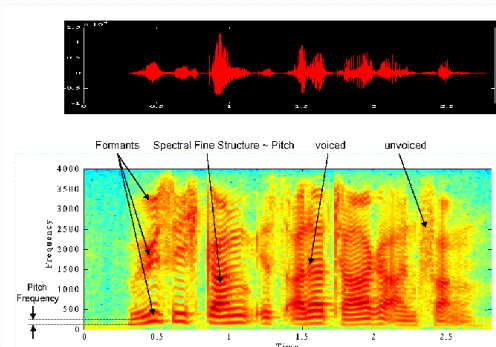
Laboratorul de cercetare Speech & Dialogue | UPB

40



Spectrul semnalului vocal

- Segmentarea semnalului vocal în ferestre (10-30 ms)
- Ferestrele sunt suprapuse
- Transformata Fourier rapidă (FFT)

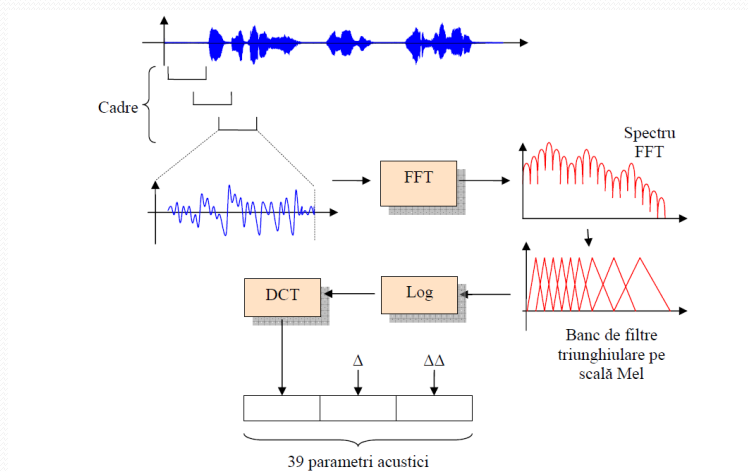


17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

43

Parametrii semnalului vocal - MFCC



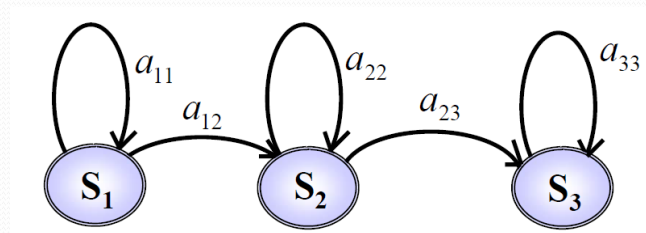
17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

44

HMM de foneme

- Fiecare fonem este reprezentat de un HMM cu 3 stări cu
- S_1 modelează începutul unui fonem, S_2 mijlocul, iar S_3 ultima parte a sunetului



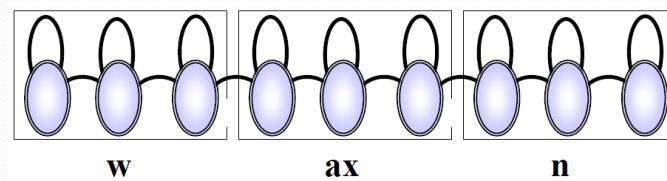
17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

45

Modelarea cuvintelor și a frazelor

- Se construiește HMM la nivel de cuvânt sau frază din concatenarea modelelor pentru foneme. De exemplu cuvântul în engleză “ONE” cu pronunția “W AX N”:



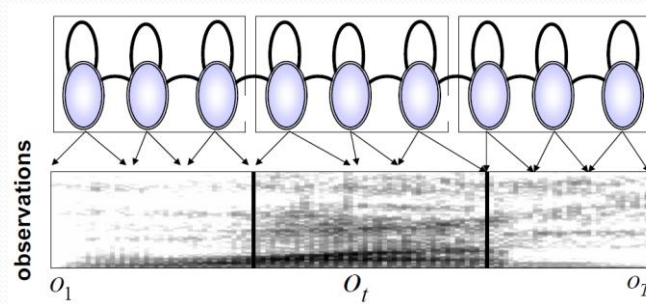
17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

46

Generarea vorbirii

- Un **semnal vocal** este generat (produs) de o **secvență** de HMM cu o anumită **probabilitate**



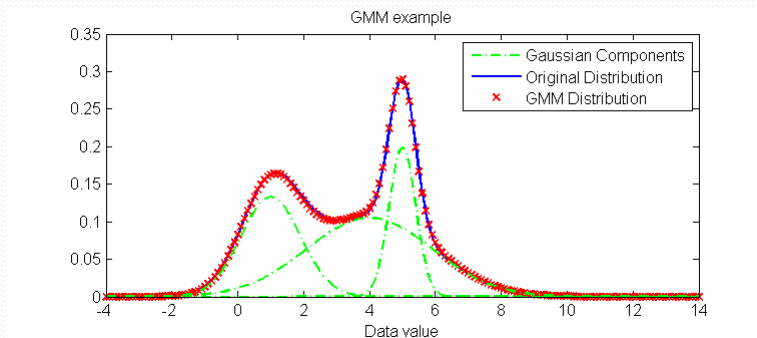
17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

47

Gaussian Mixture Models

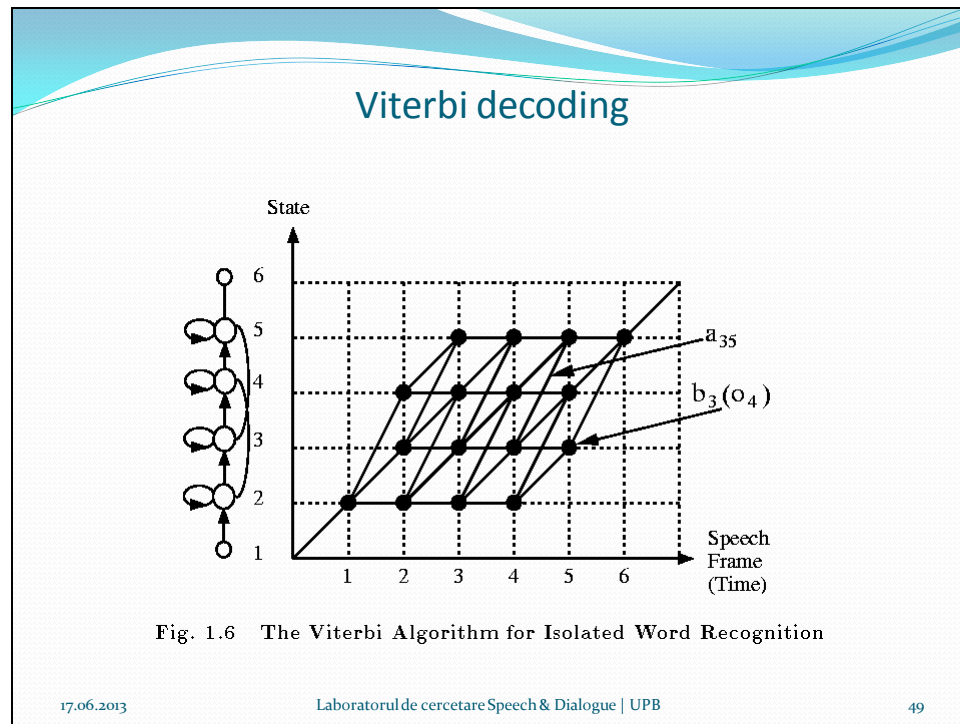
- Every distribution can be approximated with a linear combination of Gaussian distributions



17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

48



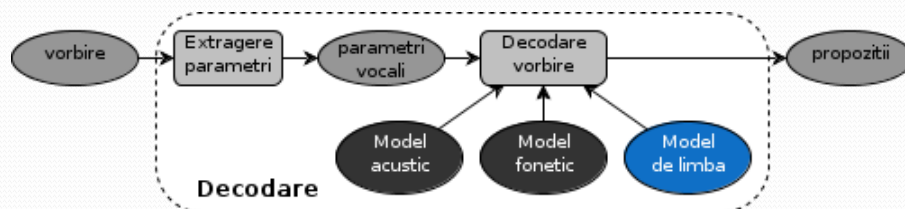
Cuprins

- Recunoașterea automată a vorbirii (RAV): introducere
- Recunoașterea automată a vorbirii: principii generale
- Resurse fonetice, acustice și lingvistice
- Modelarea acustică în sisteme de RAV
- Modelarea limbajului în sisteme de RAV
 - Modele de limbă statistice
 - Gramatici cu reguli
- Concluzii, discuții și demonstrație practică

17.06.2013
Laboratorul de cercetare Speech & Dialogue | UPB
50

Rolul modelului de limbă într-un sistem de RAV

Casa ta e mare? sau Casat a Ema re?



$$W^* = \arg \max_W p(W | X) = \arg \max_W \frac{p(X | W)p(W)}{p(X)} = \arg \max_W p(X | W)p(W)$$

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

51

Modele de limbă statistice de tip n-gram

- Se folosesc în sistemele de recunoaștere automată a vorbirii cu vocabular mare de peste 20 de ani
- Estimează probabilitatea unei secvențe de cuvinte W astfel:

$$p(W) = p(w_1, w_2, \dots, w_n) = p(w_1)p(w_2 | w_1) \dots p(w_n | w_1, w_2, \dots, w_{n-1})$$
- Istoria de cuvinte precedente - limitată la m cuvinte
 - m depinde de cantitatea de text de antrenare
 - m influențează eficiența computațională a sistemului
 - În mod tipic $m=1$ (bi-gram) sau $m=2$ (trigrame)

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

52

Modele de limbă de tip n-gram - construcție

- Exemplu pentru un model de tip bi-gram (secvențe de 2 cuvinte)
- Probabilitățile $p(w_j | w_i)$ pentru fiecare pereche (w_i, w_j) se estimează utilizând un corpus de text de antrenare, astfel:

$$p(w_j | w_i) = \frac{\text{count}(w_i, w_j)}{\sum_w \text{count}(w_i, w)}$$

- Dimensiuni tipice pentru corpusul de text de antrenare:
 - Zeci de mii de cuvinte pentru sisteme RAV adaptate la un domeniu restrâns
 - Sute de milioane de cuvinte pentru sisteme RAV fără constrângeri
- Metode de evitare a data sparseness: metode de netezire și back-off

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

53

Modele de limbă de tip n-gram - evaluare

- În contextul unui sistem de RAV
 - WER (rata de eroare la nivel de cuvânt)
- Independent, fără un sistem de RAV
 - Perplexitatea (măsoară capacitatea de predicție a unui model)
 - Rata de cuvinte în afara vocabularului (OOV-rate)

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

54

Gramatici cu reguli

- Aplicații de recunoaștere a vorbirii continue cu vocabular mare -> modele de limbă de tip n-gram
- Aplicații RAV de comandă și control -> gramatici cu reguli
- Gramaticile cu reguli specifică explicit:
 - Vocabularul de cuvinte posibile
 - Secvențele de cuvinte permise
 - Probabilitățile de apariție a cuvintelor în anumite contexte

17.06.2013

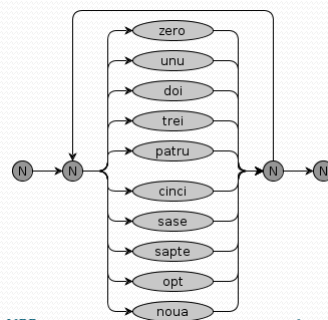
Laboratorul de cercetare Speech & Dialogue | UPB

55

Gramatici cu reguli – moduri de reprezentare

- Exemplu: sarcina de recunoaștere a cifrelor conectate
 - Vocabular de 10 cuvinte: zero, unu, doi, ..., nouă
 - Reguli: orice cuvânt poate apărea de oricâte ori, în orice context
- JSGF – Java Speech Grammar Format(<http://java.sun.com/products/java-media/speech/forDevelopers/JSGF/index.html>)


```
#JSGF V1.0;
grammar rodigits;
public <numbers> = (zero | unu | doi | trei |
    patru | cinci | șase | șapte | opt | nouă) * ;
```
- FSG – Finite State Grammar
- altele



17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

56

Cuprins

- Recunoașterea automată a vorbirii (RAV): introducere
- Recunoașterea automată a vorbirii: principii generale
- Resurse fonetice, acustice și lingvistice
- Modelarea acustică în sisteme de RAV
- Modelarea limbajului în sisteme de RAV
- Concluzii, discuții și demonstrație practică
 - De reținut despre RAV
 - Limitările și direcțiile viitoare în RAV
 - Primul sistem RAV cu vocabular extins pentru limba română

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

57

Concluzii

- Recunoașterea automată a vorbirii este abordată statistic!
- Componentele principale ale unui sistem de RAV sunt:
 - Modelul acustic (“urechea”)
 - Modelul de limbă (“creierul”)
 - Modelul fonetic (“vocabularul”)
- Complexitatea RAV depinde de:
 - Dimensiunea vocabularului
 - Calitatea înregistrării
 - Variabilitatea semnalului vocal

17.06.2013

Laboratorul de cercetare Speech & Dialogue | UPB

58

Limitări și direcții de cercetare

- Nu există sisteme de RAV cu acuratețe 100%
- Direcții de cercetare
 - Creșterea robusteții față de zgomot
 - RAV cu vocabular deschis
 - RAV pentru limbi cu resurse puține
 - RAV pentru vorbitori non-nativi
- Alte domenii de cercetare care includ RAV:
 - Înțelegerea limbajului vorbit
 - Traducerea limbajului vorbit
 - Căutarea în baze de date de vorbire
 - RAV multi-limbă

Demo & Discuții