

Universitatea ”POLITEHNICA” din București
Facultatea de Electronică, Telecomunicații și Tehnologia Informației

Recunoașterea genului vorbitorului pe baza vocii și detecția termenilor vorbiți în discursuri multi-lingve (Speech-based gender detection and multilingual spoken term detection)

PROIECT DE DIPLOMĂ

prezentat ca cerință parțială pentru obținerea titlului de
Inginer în domeniul *Electronică și Telecomunicații*
programul de studii de licență *Rețele și Software de Telecomunicații*

Conducător științific

Ș.l. Dr.Ing. Horia CUCU

Absolvent ,

Ancuța ARDELEANU

ANUL 2014

Anexa 1

Universitatea "Politehnica" din București
Facultatea de Electronică, Telecomunicații și Tehnologia Informației
Departamentul * Telecomunicații

Aprobat Director de Departament *:

Prof. Dr. Ing. Silviu Ciochina

TEMA PROIECTULUI DE DIPLOMĂ
a studentei Ardeleanu Gh. Ancuța, grupa 442D

1. Titlul temei: Recunoașterea genului vorbitorului pe baza vocii și detecția termenilor vorbiți în discursuri multi-lingve (Speech-based gender detection and multilingual spoken term detection)

2. Contribuția practică, originală a studentului va consta în (în afara părții de documentare):

Identificarea termenilor vorbiți:

- generarea și clusterizarea coeficienților mffc ai query-urilor și conținutului;
- calculul costului de similitudine între ferestrele fiecărui query și conținut pe baza distanței euclidiene;
- căutarea căii optime pe baza algoritmului DTW (Dynamic time warping);

Detecția genului vorbitorului:

- crearea unui model acustic cu înregistrări de voci masculine și feminine;
- evaluarea modelului cu alți vorbitori;
- realizarea unei aplicații client-server;

3. Proiectul se bazează pe cunoștințe dobândite în principal la următoarele 3-4 discipline: Tehnici de programare în Internet, Microcontrolere, Prelucrarea digitală semnalelor

4. Realizarea practică/ proiectul rămân în proprietatea: Laboratorul de Cercetare Speed, student Ancuța Ardeleanu

5. Proprietatea intelectuală asupra proiectului aparține: Laboratorul de Cercetare Speed, student Ancuța Ardeleanu

6. Locul de desfășurare a activității: UPB

7. Data eliberării temei: iunie 2013

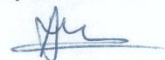
CONDUCĂTOR LUCRARE:

STUDENT:

Lect. Horia Cucu



Ancuța Ardeleanu



Copyright © 2014, ARDELEANU Ancuța

Toate drepturile rezervate Autorul acordă UPB dreptul de a reproduce și de a distribui public copii pe hârtie sau electronice ale acestei lucrări, în formă integrală sau parțială.

Declarație de onestitate academică

Prin prezenta declar că lucrarea cu titlul “*Algoritmi adaptivi cu convergență variabilă*”, prezentată în cadrul Facultății de Electronică, Telecomunicații și Tehnologia Informației a Universității “Politehnica” din București ca cerință parțială pentru obținerea titlului de *Inginer/ Master* în domeniul *Electronică și Telecomunicații* programul de studii *Tehnologii și Sisteme de Telecomunicații (ETC – TST)* este scrisă de mine și nu a mai fost prezentată niciodată la o facultate sau instituție de învățământ superior din țară sau străinătate.

Declar că toate sursele utilizate, inclusiv cele de pe Internet, sunt indicate în lucrare, ca referințe bibliografice. Fragmentele de text din alte surse, reproduse exact, chiar și în traducere proprie din altă limbă, sunt scrise între ghilimele și fac referință la sursă. Reformularea în cuvinte proprii a textelor scrise de către alți autori face referință la sursă. Înțeleg că plagiatul constituie infracțiune și se sancționează conform legilor în vigoare.

Declar că toate rezultatele simulărilor, experimentelor și măsurărilor pe care le prezint ca fiind făcute de mine, precum și metodele prin care au fost obținute, sunt reale și provin din respectivele simulări, experimente și măsurători. Înțeleg că falsificarea datelor și rezultatelor constituie fraudă și se sancționează conform regulamentelor în vigoare.

București, 1 iulie 2014

Absolvent Ancuța ARDELEANU

CUPRINS

Cuprins	9
Lista Figurilor	11
Lista cu tabele	13
Lista acronime	15
CAPITOLUL 1.INTRODUCERE	17
1.1. Motivarea lucrării	17
1.2. Obiectivele lucrării	18
CAPITOLUL 2. TEHNOLOGIA LIMBAJULUI VORBIT	21
2.1. Recunoașterea automată a vorbirii	21
2.1.1. Resursele fonetice, acustice și lingvistice necesare	23
2.1.2. Principiile de bază ale modelării acustice	24
2.1.2.1. Parametrii acustici	24
2.1.2.2. Platforma HMM/GMM	27
2.1.3. Modelarea fonetică	29
2.1.4. Principiile de bază ale modelării limbajului natural.....	30
2.1.4.1. Modele de limbă statistice.....	30
2.1.4.2. Construcția modelelor de tip n-gram	30
2.1.5. Gramatici cu stări finite FSG (Finite State Grammar).....	31
2.1.6.Evaluarea sistemelor	31
2.1.7 Unelte folosite în recunoașterea vorbirii	32
2.2 Spoken Term Detection	33
2.2.1 Introducere	33
2.2.2 Sarcina procesului și arhitectura sistemului	34
2.2.3. Termenii utilizați. Fișierele de intrare și ieșire ale sistemului	34
2.2.4. Metodologia evaluării	35
2.2.4.1 Curba erorii de detectare a compromisului	35
2.2.4.2 Performanțele de detectare ale sistemului	36
2.2.5 Procesul de căutare	37
2.2.5.1 Etapele parcurse în sistemul de căutare	37
2.3 Detecția Genului	41
2.3.1 Introducere	41
2.3.2 Uneltele folosite	42
2.3.3 Descrierea sistemului	42
2.3.4 Modelarea vorbitorului	44
2.3.5 Decodarea	46
2.3.6 Evaluarea performanțelor	48
CAPITOLUL 3. EXPERIMENTE	49
3.1. Recunoașterea secvențelor audio ce conțin cifre	49

3.1.1. Construcția unui sistem de recunoaștere a vorbirii continue	49
3.1.2 Construcția unui sistem de recunoaștere independent de vorbitor	54
3.1.3 Concluzii.....	57
3.2 MediaEval 2013: Sarcina de căutare a cuvintelor vorbite (Spoken Web Search)	58
3.2.1 Abordarea sarcinii	58
3.2.2 Sistemul de recunoaștere automată a vorbirii (ASR) pentru limba română	59
3.2.3 Metodologia căutării	59
3.2.4.Rezultate experimentale	60
3.2.5 Concluzii	61
CAPITOLUL 4. APLICAȚII DEMONSTRATIVE	63
4.1 Recunoaștere de cifre.....	63
4.1.1 Descriere.....	63
4.1.2 Interfața utilizatorului	64
4.1.3 Implementare	64
4.1.4 Concluzii	65
4.2 Detecția de gen	65
4.2.1 Descriere	65
4.2.2 Interfața utilizatorului	66
4.2.3 Implementarea.....	66
4.2.3.1 Clientul	68
4.2.3.2 Serverul	68
4.2.4 Dezvoltări ulterioare.....	69
4.2.5 Concluzii	69
CAPITOLUL 5. Concluzii	71
5.1 Concluzii generale	71
5.2 Contribuții personale	72
BIBLIOGRAFIE.....	74

Lista Figurilor

Figura 2.1	Arhitectura generală a unui sistem de recunoaștere automată a vorbirii	22
Figura 2.2	Arhitectura unui sistem de recunoaștere a vorbirii statistic	22
Figura 2.3	Spectograma propoziției ‘Ea a avut doar un copil’	25
Figura 2.4	Bancul de filtre pe frecvențe în scara MEL	26
Figura 2.5	Scara Mel în percepția auditivă	26
Figura 2.6	Schema bloc pentru analiza coeficienților MFC	27
Figura 2.7	Reprezentarea unui HMM ca un automat probabilist cu număr finit de stări	27
Figura 2.8.	Însumarea a șase gaussiene ponderate	28
Figura 2.9	Evaluarea unui sistem de recunoaștere	31
Figura 2.10	Arhitectura universală a sistemului STD	34
Figura 2.11	Gruparea punctelor în patru cluster	38
Figure 2.12	Calculul centroizilor fiecărui cluster până în momentul convergenței algoritmului	38
Figura 2.13	Formele de unda a doua semnale diferite	39
Figura 2.14	Rezultatul rulării algoritmului DTW	40
Figura 2.15	Aliniera cuvintelor cu algoritmul DTW	40
Figura 2.16	Cuvântul ‘minge’ rostit de doi vorbitori diferiți: a) femeie și b) bărbat	41
Figura 2.17	Arhitectura sistemului de diarization	42
Figura 2.18	Abordările pentru crearea unui model UBM	45
Figura 2.19	Algoritmul Viterbi	48
Figura 3.1	Gramatica cu stări finite a sarcinii de recunoaștere de cifre	51
Figura 3.2	Reprezentarea grafică a ratelor WER obținute	

	pentru un sistem dependent de vorbitor	53
Figura 3.3	Reprezentarea grafică a ratelor SER obținute pentru un sistem dependent de vorbitor	54
Figura 3.4	Reprezentarea grafică a ratelor WER obținute pentru un sistem independent de vorbitor	55
Figura 3.5	Reprezentarea grafică a ratelor SER obținute pentru un sistem independent de vorbitor	56
Figura 3.6	Reprezentarea grafică a ratelor WER obținute Pentru un sistem independent de vorbitor, pentru 100 de senone	56
Figura 3.7	Reprezentarea grafică a ratelor SER obținute pentru un sistem independent de vorbitor, pentru 100 de senone	57
Figura 3.8	Reprezentarea grafică a ratelor WER pentru diverși vorbitori, obținute pentru un sistem independent de vorbitor	57
Figura 3.9	Reprezentarea grafică a ratelor SER pentru diverși vorbitori, obținute pentru un sistem independent de vorbitor	58
Figura 3.10	Rezultatele rulării principale cu baza de evaluare	61
Figura 4.1	Interfața aplicației demonstrative	64
Figura 4.2	Interfața aplicației demonstrative	66
Figura 4.3	Diagrama de secvență a aplicației demonstrative	67
Figura 4.4	Diagrama UML (de clase) a aplicației demonstrative	67

Lista cu Tabele

Tabel 2.1	Clase de ambiguitate în fonetizarea cuvintelor limbii române	23
Tabel 3.1	Ratele WER obținute pentru un sistem dependent de vorbitor	53
Tabel 3.2	Ratele SER obținute pentru un sistem dependent de vorbitor	54
Tabel 3.3	Ratele WER obținute pentru un sistem independent de vorbitor	55
Tabel 3.4	Ratele SER obținute pentru un sistem independent de vorbitor	55
Tabel 3.5	Ratele WER obținute pentru un sistem independent de vorbitor, 100 de senone și un număr mai mare de densități gaussiene	56
Tabel 3.6	Ratele SER obținute pentru un sistem independent de vorbitor, 100 de senone și un număr mai mare de densități gaussiene	57
Tabel 3.7	Rezultatele algoritmului DTWSS	60

Lista acronime

ASR	Automatic Speech Recognition
ATWV	Actual Term Weighted Value
AVCSR	Audio-Video Continuous Speech Recognition
BIC	Bayesian Information Criterion
ChER	Character Error Rate
CLR	Cross Likelihood Ratio
CMU	Carnegie Mellon University
DER	Diarization Error Rate
DET	Detection Error Tradeoff
DCT	Discrete Cosinus Transform
DFT	Discrete Fourier Transform
DTW	Dynamic Time Warping
DTWSS	Dynamic Time Warping String Search
EDT	Event Dispatch Thread
EM	Expectation Maximization
FAP	False Alarm Probability
FFT	Transformata Fourier Discretă
FST	Finite-state Transducer
FSG	Finite State Grammar
GMM	Gaussian Mixture Model
HMM	Hidden Markov Model
HP	Hewlett Packard
HTK	Hidden Markov Toolkit
JSGF	Java Speech Grammar Format
LIUM	The Laboratoire d'Informatique de l'Université du Maine
LMG	Literal movement grammars
MAP	Maximum a posteriori
MERL	Laboratoarele Mitsubishi Electric Research Labs
MIT	Massachusetts Institute of Technology
MFCC	Mel-frequency cepstrum coefficients
ML	Maximum Likelihood
MP-	Miss Probability

MTWV	Maximum Term Weighted Value
NIST	National Institute of Standards and Technology
OOV	Out Of Vocabulary
PhER	Phone Error Rate
PLP	Perceptual Linear Prediction
PPL	Perplexity
RAV	Recunoașterea Automată a Vorbirii
RTTM	Real-Time Trade Matching
SER	Sentence Error Rate
SNR	Signal to Noise Ratio
ST	Sausage Technique
STD	Spoken Term Detection
TCP	Transmission Control Protocol
TWV	Term Weighted Value
UBM	Universal Background Model
UCSC	University Of California, Santa Cruz
UML	Unified Modeling Language
UTF	Universal Character Set Transformation Format
WER	Word Error Rate
XML	Extensible Markup Language

CAPITOLUL 1

INTRODUCERE

1.1 Motivarea lucrării

Lucrarea oferă o privire de ansamblu asupra tehnologiilor în domeniul recunoașterii vorbirii, accentul punându-se pe recunoașterea independentă de vorbitor. Progresele recente în tehnologie au crescut așteptările utilizatorilor în legătură cu diversitatea aplicațiilor de vorbire. Una dintre aceste schimbări principale o reprezintă nevoia de a susține ca date de intrare înregistrările în mai multe limbi. Astăzi există o diversitate de informații text disponibile pe Internet în cadrul carora se pot efectua căutări, dar cele în format audio sau video reprezintă o sursă mai comodă de informare. Procesarea informațiilor a devenit o activitate economică principală în lume, informațiile rostite reprezentând sursa majoră a acestor informații.

Lucrarea este motivată de toți acești factori și a fost creată pentru a ilustra mai mult forme de recunoaștere bazate pe calitățile vocii, ce pot servi la diverse aplicații utile în lumea reală. Procesul de recunoaștere automată a vorbirii (ASR) se ocupă de maparea unui semnal acustic cu o secvență de cuvinte.

Scopul detecției de termeni vorbiți este acela de a găsi un anumit termen într-o bază de înregistrări audio și reprezintă o oportunitate de a găsi informații în arhivele de vorbire înregistrate.

Vocea umană nu furnizează doar semantica cuvintelor rostite, ci reprezintă un semnal cu o infinită informație, caracteristice dependente de vorbitor precum gen, emoții, etnie, dialect, vârstă, identitate. În cadrul apelurilor telefonice de zi cu zi, ne bazăm pe aceste caracteristici ale vocii interlocutorului și încercăm să ne adaptăm propriul stil de vorbire în concordanță cu al său.

Recunoașterea genului cu ajutorul vocii se bazează pe forma tractului vocal, accent, ritm, stilul intonației și este utilă în simplificarea sarcinilor în domeniul vorbirii și îmbunătățirea serviciilor destinate clienților.

Cea de-a treia parte abordată în lucrare, detectarea cifrelor rostite, prezintă importanța în diverse aplicații precum: apelarea telefonică vocală, sisteme automate bancare, intrarea automată a datelor, intrarea PIN-ului. Este o sarcină de recunoaștere cu un vocabular foarte redus și fiind pentru limba română, limba țintă a acestui subcapitol, avem puține resurse text, audio sau video disponibile. Experimentele făcute în cadrul acestei abordări, variind numărul de senone și de densități gaussiene ajută la determinarea modelului acustic optim pentru a obține o rată de detecție greșită cât mai mică.

1.2 Obiectivele lucrării

Având în vedere observațiile făcute în motivarea lucrării, obiectivul principal al acestei lucrări este de a arăta diversele forme pe care le poate lua un sistem de recunoaștere a vorbirii. Sistemul de recunoaștere a cifrelor, o parte a lucrării, trebuie să fie capabil să recunoască cifrele rostite de un vorbitor și să le afișeze în forma textuală, iar sistemul de recunoaștere a genului. Pentru a realiza aceste lucruri, au fost folosite următoarele obiective specifice:

- a) Crearea modelelor (acustic, fonetic și lingvistic) pentru sistem de recunoaștere automată a vorbirii în limba română;
- b) Realizarea unor experimente folosind diverse modele de gen și configurații ale procesului de decodare pentru a evalua cele mai bune rezultate;
- c) Efectuarea unor experimente pentru a găsi un termen rostit într-o bază de clipuri audio și optimizarea procesului de căutare.
- d) Dezvoltarea unor aplicații care să realizeze aceste detecții de cifre, respectiv de gen și returnează rezultatele.

Lucrarea este organizată în 5 capitolele, precum urmează:

Capitolul 1 prezintă introducerea și conține motivarea care stă la baza realizării sale. Începe cu prezentarea tehnologiilor din domeniul vorbirii, importanța lor în diversitatea actuală a informațiilor audio și domeniile lor de aplicabilitate și continuă cu obiectivele lucrării.

Capitolul 2 dezvoltă aspectele teoretice ale sistemelor de recunoaștere automată a vorbirii, a detecției de gen și a detecției de termeni vorbiți. Sunt descrise, de asemenea, arhitecturile sistemelor, algoritmi de căutare, unelele folosite și metodele de evaluare.

Capitolul 3 subliniază experimentele făcute cu scopul de găsi modelul acustic optim pentru recunoașterea de cifre și cele realizate în cadrul competiției MediaEval 2013. Astfel, se începe cu construcția unui sistem de recunoaștere a vorbirii continue dependent de vorbitor, se continuă cu un sistem independent de vorbitor, pentru ambele variindu-se parametrii (numărul de senone și numărul de densități gaussiene). Capitolul se concentrează apoi, pe algoritmi de căutare și metodele utilizate în sistemul de căutare a termenilor vorbiți, precum și pe rezultatele obținute prin varierea parametrilor α și β .

În capitolul 4 sunt descrise cele două aplicații demonstrative împreună cu principiile lor de funcționare. Aplicația de detecție de cifre leagă modelul acustic optim obținut în urma experimentelor cu modelul

lingvistic și dictionarul. Acest penultim capitol este încheiat cu o scurtă descriere a dezvoltărilor ulterioare ce pot fi efectuate asupra aplicației de detecție a genului.

Capitolul 5 este dedicat concluziilor și contribuțiilor personale în realizarea lucrării.

CAPITOLUL 2

TEHNOLOGIA LIMBAJULUI VORBIT

2.1 Recunoașterea automată a vorbirii [ASR]

Recunoașterea automată a vorbirii reprezintă procesul de transcriere a vorbirii în text, adică transformarea semnalului acustic într-o secvență de cuvinte. Până de curând limbajul vorbit era o modalitate de interacțiune între subiecții umani, însă, o dată cu dezvoltarea microelectronicii, comunicarea prin voce a trecut la stadiul de conectare subiect uman-mașină de calcul, având în momentul de față o gamă largă de aplicabilitate. Cel mai important domeniu ar fi cel al interfețelor hands-free și eyes-free. Este util de asemenea în aplicațiile de dictare automată, în sistemele de dialog pentru call center, sistemele de traducere speech-to-speech, sau când dorim transcrierea unui monolog al unui vorbitor anume (cazul transcrierii discuțiilor din timpul proceselor de judecată, editarea unor documente de către jurnaliști sau scriitori). Arhitectura generală a unui sistem de recunoaștere automată a vorbirii este prezentată în figura 2.1.

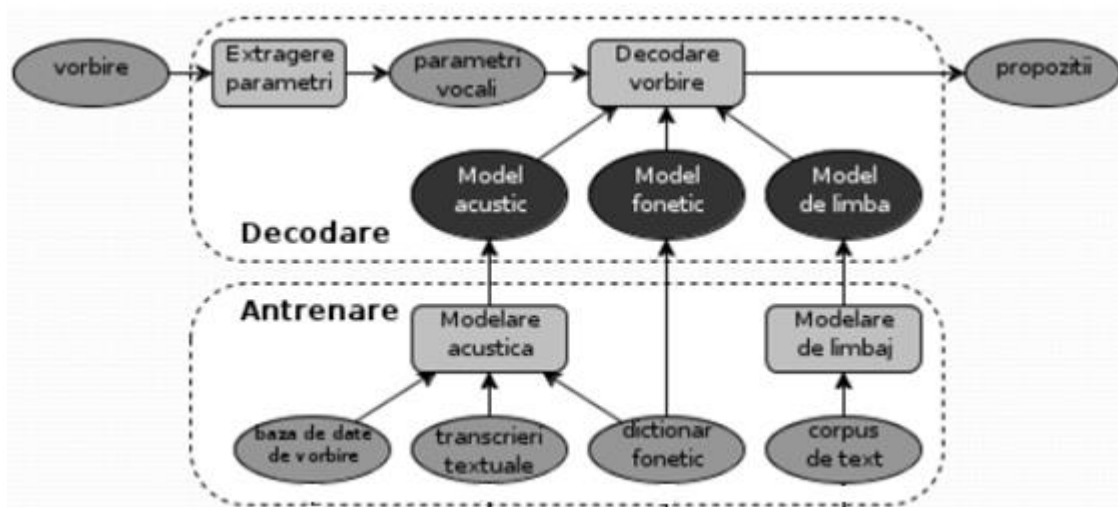


Figura 2.1 Arhitectura generală a unui sistem de recunoaștere automată a vorbirii

Modelul unui asemenea sistem statistic poate fi simplificat astfel, figura 2.2 :

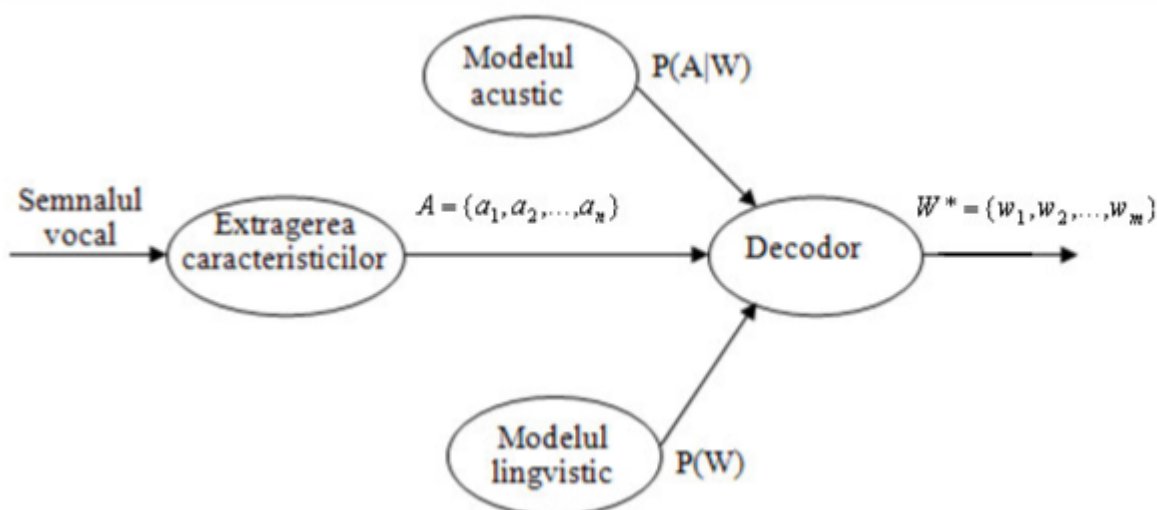


Figura 2.2 Arhitectura unui sistem de recunoaștere a vorbirii statistic

Având după prelucrarea și extragerea parametrilor semnalului vocal un set de vectori acustici $A = \{a_1, a_2, \dots, a_n\}$, căutăm o secvență ce se potrivește cel mai bine $W^* = \{w_1, w_2, \dots, w_m\}$:

$$W^* = \arg \max_W p(W|A) \quad (2.1)$$

Și folosind regula lui Bayes, din relația 2.1. rezultă:

$$W^* = \arg \max_W \frac{p(A|W)p(W)}{p(A)} = \arg \max_W (p(A|W)p(W)) \quad (2.2)$$

unde s-a ținut cont de faptul că probabilitatea mesajului vorbit $p(A)$, este independentă de secvența de cuvinte W și poate fi ignorată, $p(A|W)$ fiind probabilitatea mesajului vorbit și $P(W)$ probabilitatea de apariție a secvenței de cuvinte W . Problema găsirii unei secvențe de cuvinte pe baza mesajului vorbit a fost astfel împărțită în două sub-probleme mai simple:

- estimarea probabilității apriorice $p(W)$, utilizând un model de limbă;
- estimarea probabilității $p(A|W)$ cu ajutorul unui model acustic.

2.1.1. Resursele fonetice, acustice și lingvistice necesare

Precum se observă din schema arhitecturii sistemului de recunoaștere a vorbirii, pentru decodare avem nevoie de un model acustic, unul fonetic și unul de limbă, pe care trebuie să le pregătim în prealabil.

- *Modelul acustic*

Este construit pe baza setului de înregistrări audio și a transcrierilor textuale ale acestora.

- *Modelul fonetic*

Conține un model de pronunție pentru fiecare cuvânt, chiar dacă există mai multe pronunții, acestea sunt comasate.

- *Modelul de limbă*

Se bazează pe un corpus de text de dimensiuni mari.

1) Dicționarul fonetic conține transcrierea cuvintelor din forma lor literară în forma fonetică (precum se pronunță o succesiune de foneme). Acesta face legătura între modelul acustic și cel de limbă și trebuie să conțină cuvintele ce sunt posibile într-o anumită sarcină de recunoaștere a vorbirii. Fiecare intrare în dicționar conține atât cuvântul cât și forma lui fonetică și se poate realiza manual, automat sau semi-automat. În limba română, de foarte multe ori unei litere îi corespunde același fonem, însă există și excepții, când o literă se pronunță diferit în funcție de contextul în care apare.

Tabel 2.1 Clase de ambiguitate în fonetizarea cuvintelor limbii române

Litera	Fonemul	Exemple
e	e	cer, rest, rece
	e1	deal, te-am, ceainic, gheorghe
	i3e	eu, el, eram, erai
o	o	gol, potrivit
	o1	voal, soare

2) Baza de date de vorbire constituie într-o mulțime de înregistrări salvate în format wav și se folosește pentru antrenarea modelului acustic, cuprinzând 3 părți:

- un set de clipuri audio;
- un set ce corespunde transcrierilor textelor rostite în clipuri audio, precum și delimitarea fonemelor;
- în formații suplimentare despre domeniul vorbirii, identitatea vorbitorului;

Ea este un element important de care depinde atât rata de eroare cât și viteza de recunoaștere. Pentru fiecare bază de date se păstrează informații legate de tipul de înregistrare, numărul de cuvinte, dimensiunea bazei de date, calitatea etc. Indiferent de algoritmul de antrenare, cu cât baza de date este mai mare cu atât rezultatele obținute vor fi mai bune.

Înregistrarea acestor clipuri presupune parcurgerea unor pași precum: alegerea locului de înregistrare (laborator, studio), alegerea microfonului, stației de înregistrare, precum și alegerea textelor ce trebuiesc înregistrate. Acest ultim pas ține cont de mai multe criterii:

- frazele să fie corecte din punct de vedere lexical și gramatical;
- frazele să fie scurte (maxim 15-20 de cuvinte);
- frazele să fie conectate semantic pentru a-i fi mai ușor vorbitorului să înțeleagă sensul și să obțină o pronunție fluentă;

În cazul în care vorbitorul cunoaște aprioric textul ce va fi înregistrat, erorile de pronunție, bâlbele sau lipsa de fluentă sunt mult mai rare.

3) Corpusul de text este folosit la antrenarea modelului lingvistic, iar achiziția sa este realizată în mod automat, prin folosirea unui sistem de achiziție a textelor de pe Internet (pe principiul Web-as-Corpus).

Deoarece sistemele de recunoaștere a vorbirii transcriu vorbirea și scot la ieșire text simplu (fără simboluri speciale, litere mari etc.) textele trebuiesc normalizate și restaurate astfel:

- transformarea literelor mari în litere mici;
- transcrierea numerelor formate din cifre în numere din litere;
- înlăturarea caracterelor speciale și a semnelor de punctuație;
- scoaterea etichetelor HTML în cazul în care textele sunt de pe internet;

2.1.2 Principiile de bază ale modelării acustice [J01]

Sistemele de recunoaștere a vorbirii continue nu estimează probabilitatea mesajului, ci se bazează pe unități de vorbire mai mici (foneme) și calculează probabilitățile acestora. Astfel modelele pentru cuvinte stau la baza modelelor pentru secvențele de cuvinte și se bazează pe modelele pentru unitățile sub-lexicale. Ele sunt utilizate în final pentru a estima probabilitatea ca mesajul rostit de vorbitor să fie format dintr-o succesiune sau alta de cuvinte.

2.1.2.1 Parametrii acustici

Modelul acustic este implementat utilizând modelele Markov ascunse (HMM). Un astfel de HMM este un automat probabilistic cu stări finite, ce constă într-un set de stări conectate prin tranziții, dar înșiruirea stărilor este ascunsă. Aceste foneme legate formează modelele cuvintelor și apoi, secvențe de cuvinte ce sunt utilizate în estimarea lui $p(A|W)$.

Definite HMM:

Un HMM [R01] este complet definit prin:

- spațiul de stări: $Q = \{1, 2, 3, \dots N\}$;

- numărul de evenimente: $E = \{e_1, e_2, e_3, \dots, e_M\}$;
- probabilitatea ca stare inițială să fie i : $\pi_j = P[q_1 = j]$;
- $A = \{a_{ij}\}$ matricea cu probabilitățile de tranziție, unde $a_{ij} = P[q_t = j \mid q_{t-1} = i]$ este probabilitatea de a trece din starea i în starea j ;
- $B = \{b_j(k)\}$ matricea de ieșire, unde $b_j(k)$ este probabilitatea să se emită simbolul o_k la starea i ;
- $O = \{o_1, o_2, o_3, \dots, o_M\}$ - mulțimea observațiilor de ieșire, în cazul în care această mulțime este discretă;

Un model întreg se poate nota cu : $\lambda = (A, B, \pi)$;

Cum parametrii de voce iau valori într-un spațiu continuu trebuie ca și vectorul de ieșire al HMM-ului să aibă valori continue. Din acest motiv se folosesc mixturile gaussiene pentru modelarea ieșirii stărilor HMM-ului.

Sistemul de recunoaștere a vorbirii bazat pe HMM este format din HMM-urile unităților ‘subcuvintelor’ numite foneme. Numărul stărilor HMM distincte într-un sistem depinde de mărimea setului de unități ale subcuvintelor [H02]. În figura 2.3 este prezentată spectrograma unei propoziții.

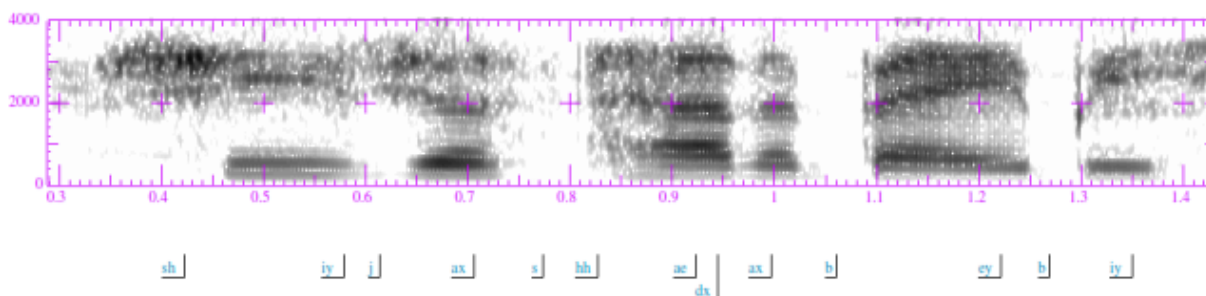


Figura 2.3 Spectrograma propoziției ‘Ea a avut doar un copil’

Deoarece semnalul de vorbire este unul nestaționar, analiza spectrală nu poate fi făcută pe întreg semnalul, ci pe cadre scurte (20-30ms), cvasi-staționare. Astfel semnalul în domeniul timp este împărțit în ferestre, iar pentru fiecare fereastră sunt extrași parametrii de tip spectral sau cepstral.

Cei mai folosiți coeficienți în recunoașterea vorbirii sunt cei MEL(*Mel Frequency Cepstral Coefficients*) cepstrali MFCC, deoarece sunt în număr mai mic față de cei spectrali (13 în loc de 64 sau 128 în celălalt caz), nu sunt corelați între ei.

Coeficienții MFCC sunt bazați pe anvelopa spectrală a logaritmului semnalului, ce folosește o scară neliniară de frecvență, deoarece simulează mai bine modul neliniar de percepție al omului. Semnalului $x(n)$ i se aplică o transformată Fourier discrete (DFT) cu formula:

$$X_a[k] = \sum_{n=0}^{N-1} x[n] e^{-j2\pi nk/N}, 0 \leq n < N$$

Se folosește apoi un banc de M filtre ($m=1,2,\dots,M$) triunghiulare, ce sunt plasate pe axa frecvențelor astfel încât frecvența centrală a filtrelor să urmeze scara mel. Banda lor crește o dată cu indexul m , conform figurii 2.4 [B01] :

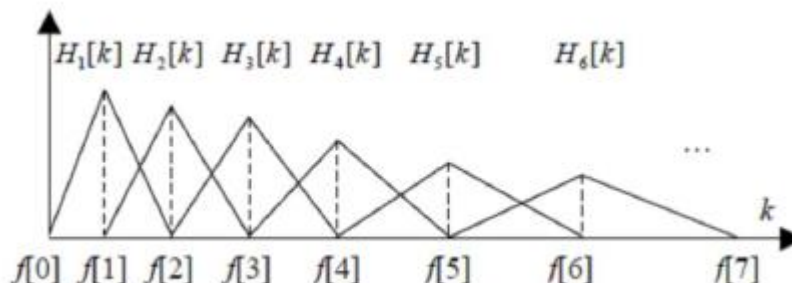


Figura 2.4 Bancul de filtre pe frecvențe în scara MEL

După cum se poate observa din figură, filtrarea cu scala Mel subliniază frecvențele joase. Relația între frecvențele MEL și cele liniare este, figura 2.5 :

$$\text{Frecvența_mel} = 2595 * \log\left(1 + \frac{\text{Frecvența_liniară}}{700}\right)$$

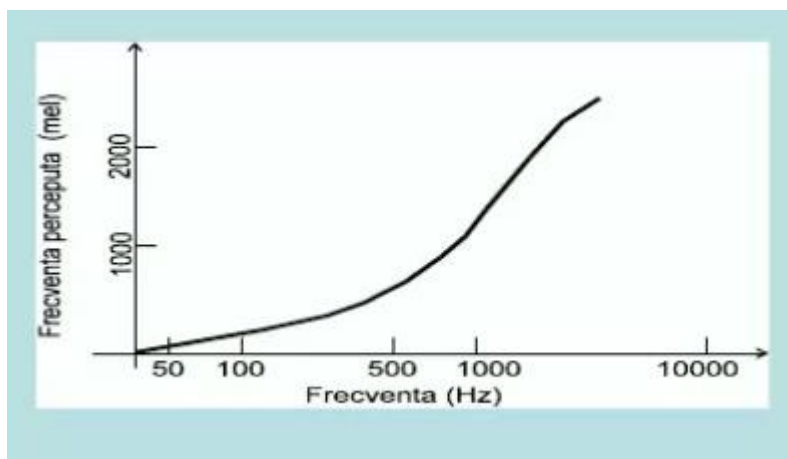


Figura 2.5 Scara Mel în percepția auditivă

Apoi, se calculează energia logaritmică la ieșirea fiecărui filtru:

$$S[m] = \ln\left(\sum_{n=0}^{N-1} |X_a[k]|^2 * H_m[k]\right) \quad , \quad 0 \leq m \leq M$$

Această reprezentare este nivelată și ortogonalizată prin aplicarea transformatei cosinus discrete (*DCT* – *Discrete Cosinus Transform*) cu ajutorul formulei :

$$c[n] = \sum_{m=0}^{M-1} S[m] * \cos\left(\pi n \left(m + \frac{1}{2}\right) M\right), \quad 0 \leq n \leq M$$

și astfel rezultă 7 coeficienți cepstrali pentru fiecare vector mel. M poate varia de la 24 la 40, dar pentru recunoașterea de vorbire se folosesc 13. Acești vectori, în mod normal, se calculează la fiecare 10ms cu ferestre Hamming de 25ms în interiorul semnalului. Coeficienții MFCC sunt creați după modelul din figura 2.6 :

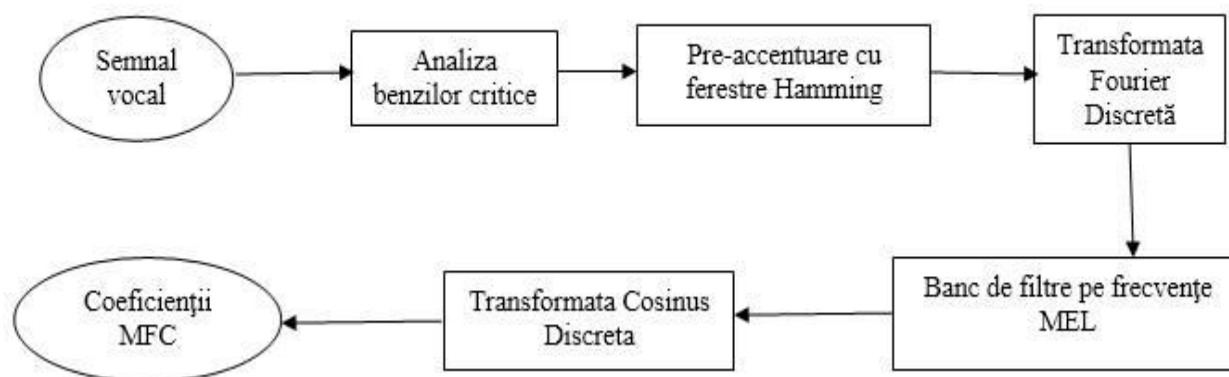


Figura 2.6. Schema bloc pentru analiza coeficienților MFC

2.1.2.2 Platforma HMM/GMM

Așa cum am văzut mai sus, un HMM este un automat cu un număr finit de stări, format dintr-un set de stări conectate prin arce, ce reprezintă tranzițiile. Prima și ultima stare a unui HMM este non-emisivă.

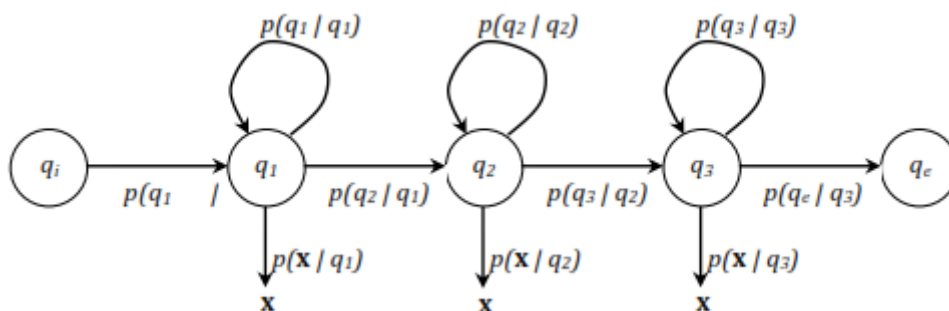


Figura 2.7. Reprezentarea unui HMM ca un automat probabilist cu număr finit de stări

Probabilitățile observațiilor de ieșire a unei stări q_i sunt modelate de o funcție densitate de probabilitate, sub forma unei sume de K funcții gaussiene.

Dacă X este o variabilă aleatoare continuă, atunci densitatea de probabilitate a lui X este o funcție $f(x)$ pentru care:

$$P(a \leq X \leq b) = \int_a^b f(x) dx,$$

deci probabilitatea ca X să aibă valori între a și b reprezintă aria densității între a și b . Orice funcție arbitrară a densității de probabilitate poate fi construită însumând N Gaussiene ponderate (mixture de Gaussiene):

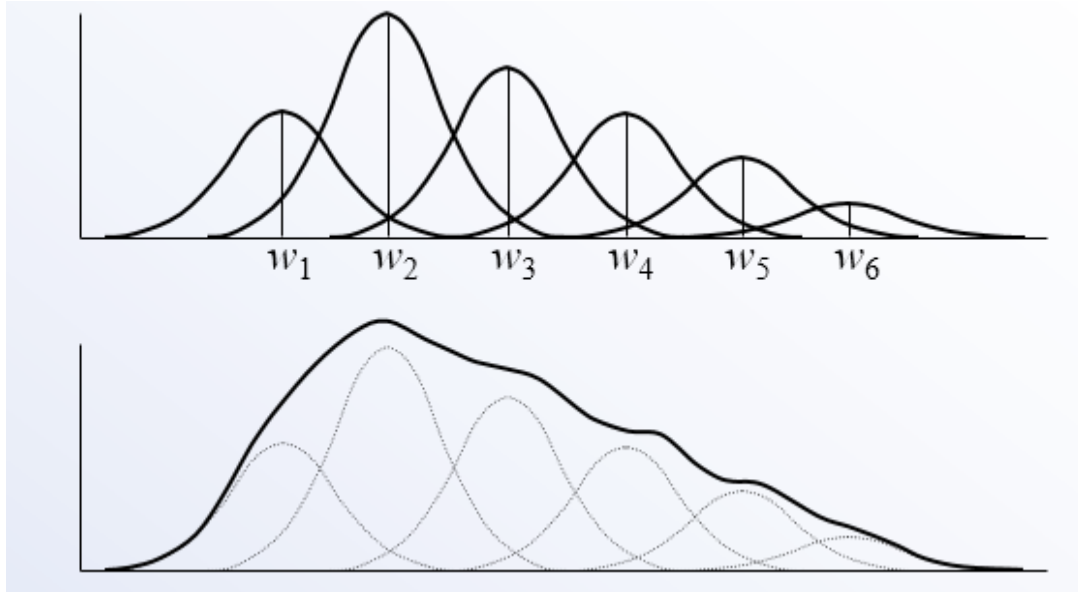


Figura 2.8. Însurarea a șase gaussiene ponderate

Astfel se formează un GMM (din engleză Gaussian Mixture Model).

În cazul în care ieșirea unei stări este o mixtură gaussiană (GMM), funcția densitate de probabilitate se calculează cu formula:

$$b_i(x) = p(x|\lambda_j) = \sum_{k=1}^K w_{ik} \mathcal{N}(x; \mu_{ik}, \Sigma_i),$$

în care x este un vector D -dimensional de date, cu valori continue (vectorul trăsăturilor), w_{ik} sunt mixturile ponderate, iar $\mathcal{N}(x; \mu_{ik}, \Sigma_i)$, $i = 1, \dots, K$, sunt componentele densităților Gaussiene, cu vectorul mediilor μ_i și matricea de varianță Σ_i :

$$\mathcal{N}(x, \mu_{ik}, \Sigma_i) = \frac{1}{(2\pi)^{\frac{D}{2}} |\Sigma_i|^{\frac{1}{2}}} \exp\left(-\frac{1}{2}(x - \mu_{ik})^T \Sigma_i^{-1}(x - \mu_{ik})\right)$$

Când realizăm modelarea vorbirii cu ajutorul HMM-urilor putem face două presupuneri [B01], [C02] :

Independența observațiilor: probabilitatea ca o stare să genereze un anumit vector acustic este independența de vectorii generați în stările anterioare.

Proces Markov: tranziția între stările HMM este un proces Markov de ordinul întâi în care probabilitatea tranziției următoare depinde numai de starea actuală.

Cele trei probleme ale platformei HMM/GMM sunt: evaluarea, decodare și estimarea, pentru care există 3 algoritmi:

- Antrenarea (estimarea diversilor parametri ai unui sistem HMM/GMM) folosește unul din algoritmi Forward-Backward sau Baum-Welch (o particularizare a lui expectation-maximization-EM);
- Decodarea (identificarea celei mai probabile secvențe de stări ce ar putea genera secvența de observații) se rezolvă cu algoritmul Viterbi;
- Evaluarea (aflarea probabilității ca o secvență de observații să fie emisă de un HMM) se realizează cu algoritmul Forward.

Cheia algoritmului Forward se rezumă la calculul probabilității care urmează $\alpha_t(q_t) = p(x_1, \dots, x_t, q_t = q_t | \lambda)$, adică probabilitatea de observare a secvenței $x_1, \dots, x_t$, fiind în starea q_t la momentul t . Cu ajutorul presupunerilor Markov putem folosi formula:

$$\alpha_t(q_t) = \sum_{q_{t-1}} \alpha_{t-1}(q_{t-1}) a_{q_{t-1}q_t} b_{q_t}(x_t)$$

Algoritmul Viterbi funcționează în maniera asemănătoare cu cel anterior, doar că sumarea nu se mai realizează la fiecare pas, ci se calculează valoarea maximă.

2.1.3 Modelarea fonetică

O unitate de vorbire potrivită trebuie să aibă cel puțin trei caracteristici, potrivit autorilor [H01] :

- Unitatea să fie precisă: să aibă realizarea acustică care apare în diferite contexte;
- Unitatea să fie antrenabilă: să dispunem de suficiente date astfel încât să estimăm parametrii asociați ei;
- Unitatea să fie generalizabilă, astfel încât să putem construi orice cuvânt dintr-un set de unități de bază.

Sistemele de recunoaștere a vorbirii nu folosesc cuvinte, deoarece acestea nu sunt nici antrenabile, nici generalizabile, ci se bazează pe foneme, ca unități de vorbire de bază, ce sunt independente de vocabularul cerinței de recunoaștere. Modelul fonetic nu este potrivit deoarece, se presupune, că un fonem este pronunțat identic, indiferent de context, iar oamenii nu pot pronunța un cuvânt ca un set de foneme independente.

Dacă am utiliza foneme dependente de context și am avea la dispoziție suficiente informații pentru antrenare și implicit pentru estimarea corectă a parametrilor modelelor corespunzătoare, am avea o precizie de recunoaștere mult mai bună.

Modelul de trifeneme este un model fonetic ce ține cont și de fonemele din dreapta și stânga sa. Dacă găsim două foneme identice în contexte diferite, ele sunt considerate trifeneme diferite. Folosirea de trifeneme ar conduce la o precizie mai puternică, deoarece conțin și efectele co-articulare, însă datorită numărului mare de modele diferite, am avea un număr ridicat de parametri estimați, ce trebuie de asemenea antrenați.

Scara mel pentru percepția auditivă bazată pe HMM este formată din HMM-urile unităților ‘subcuvintelor’ numite foneme. Numărul stărilor HMM distincte într-un sistem depinde de mărimea setului de unități al subcuvintelor.

Modelul fonetic conține un dicționar de pronunții, mapând toate cuvintele din vocabular cu o secvență de foneme. Construcția acestuia se poate face manual sau automat, însă ultimul mod nu oferă o precizie la fel de bună ca primul.

2.1.4. Principiile de bază ale modelării limbajului natural

2.1.4.1. Modele de limbă statistice

Cea mai mare diferență între un sistem de recunoaștere a cuvintelor izolate și unul de vorbire continuă este modelul de limbă. Scopul sau în cadrul decodării, așa cum este arătat în figura 2.1., este de a estima probabilitatea ca o secvență de cuvinte $W=w_1, w_2, \dots, w_n$ să fie o secvență validă pentru o sarcină de recunoaștere. Această probabilitate este complementare celei de la modelul acustic.

Probabilitatea se poate calcula cu formula :

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

și se poate observa că depinde de istoria cuvântului, adică de secvența de cuvinte anterioare lui, ce nu poate include un număr infinit de cuvinte, ci doar m cuvinte. Astfel s-a ajuns la modelul de limbă de tip n -gram [R01], ce este un model mult mai evoluat. Acesta prezice următorul cuvânt pe baza celor $n-1$ cuvinte anterioare.

Se folosesc modele de tip trigram ($n=3$), unde este nevoie de o istorie de două cuvinte pentru a-l putea prezice pe al treilea sau bigram ($n=2$).

În recunoașterea de vorbire, sunetele de intrare se pot confunda foarte ușor, deoarece multe cuvinte sună foarte asemănător.

2.1.4.2. Construcția modelelor de tip n -gram

În cazul unui model de limbă bigram, probabilitățile $p(w_i|w_j)$ pentru fiecare pereche de cuvinte (w_i, w_j) se calculează cu formula probabilității maxime (Maximum Likelihood-ML) :

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

în care se numără de câte ori cuvântul w_i este urmat de w_j față de alte cuvinte;

Pentru un trigram formula se modifică astfel:

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

Avem nevoie de o cantitate mare de date de antrenare pentru a obține un model n-gram precis. Există însă și probleme precum existența unor n-grame ce nu apar în textul de antrenare însă apar în textele de evaluare, indiferente de mărimea corpusului, sau existența unor n-grame ce apar de puține ori (mai puțin de zece ori). Probleme se rezolvă cu metodele de back-off, care folosesc multiple modele de limbă pentru a crea un model de limbă interpolat, care să poată utiliza toate părțile ce îl compun.

2.1.5. Gramatici cu stări finite FSG (Finite State Grammar)

O gramatică cu stări finite FSG este un graf în care fiecare nod reprezintă un cuvânt al vocabularului sarcinii de recunoaștere, iar arcele dintre noduri sunt tranziții între cuvinte. Astfel sunt specificate toate secvențele de cuvinte permise de gramatica sarcinii respective. Dacă atribuim acele costuri, specificăm și probabilitatea ca un cuvânt să fie precedat de altul. În cazul sarcinilor de mici dimensiuni (recunoaștere de cifre, apelarea unui număr de telefon), acest model poate fi folosit cu succes.

2.1.6. Evaluarea sistemelor

Evaluarea unui sistem de recunoaștere a vorbirii necesită existența unei baze de date de înregistrări, numită baza de date de evaluare și se întemeiază pe transcrierile clipurilor audio din bază. Acestea descriu cât de bun este un sistem de recunoaștere a vorbirii pentru a transcrie vorbirea dintr-un alt domeniu. Criteriile de bază pentru evaluarea unui model de limbă sunt rata de cuvinte în afara vocabularului (Out of vocabulary OOV) și perplexitatea (PPL), ce măsoară capacitatea de predicție a unui model de limbă. Cu cât e mai mică perplexitatea, cu atât e mai bună capacitatea de predicție a acelui model [H01].

Perplexitatea unui cuvânt din afara vocabularului este infinită și acest lucru face evaluarea performanței unui model de limbă dificilă. În acest caz, procentul cuvintelor OOV trebuie specificată când evaluăm un model de limbă.

Evaluarea completă a sistemului de recunoaștere a vorbirii este realizată prin compararea transcrierilor clipurilor audio de evaluare cu transcrierile generate în urma procesului de decodare, conform cu figura 2.9 :

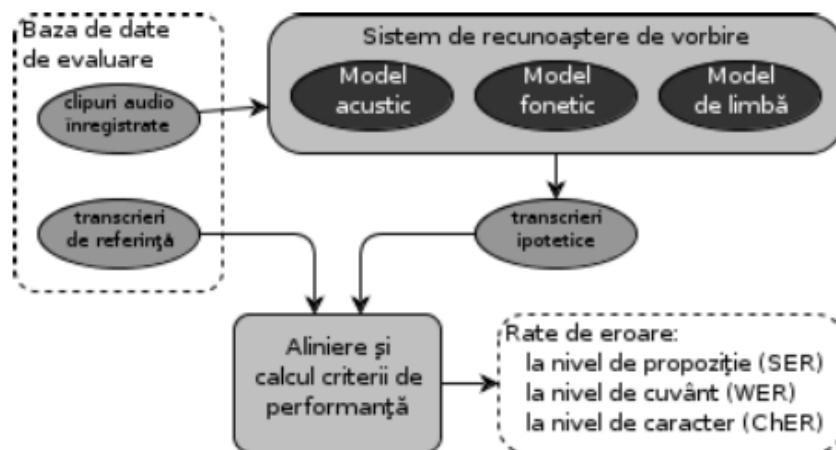


Figura 2.9 Evaluarea unui sistem de recunoaștere

Rata de eroare la nivel de propoziție (SER-Sentence Error Rate) este un criteriu de performanță utilizat în evaluare și este raportul între numărul de propoziții eronate (ce conțin cel puțin o greșeală) și numărul total de propoziții din transcrierea de referință:

$$\text{SER}[\%] = 100 * \frac{\text{\#PROPOZIȚII ERONATE}}{\text{\#PROPOZIȚII ÎN TRANSCRIEREA DE REFERINȚĂ}}$$

Aceasta rată se folosește doar în situațiile în care orice propoziție ce conține o greșeală este inutilă (sistem de recunoaștere a unui IBAN, CNP).

Metrica standard de evaluare a sistemelor de recunoaștere a vorbirii este rata de eroare a cuvintelor (WER-Word Error Rate). Această rată se referă la cât de mult cuvântul returnat de recunoașcător (numit și sir ipotetic) diferă de transcrierea de referință. Primul pas în aflarea ratei de eroare constă în determinarea distanței minime între cuvântul corect și cel recunoscut. Rezultatul acestui calcul va da numărul minim de cuvinte substituite, inserate sau șterse și astfel, formula ratei de eroare va fi:

$$\text{WER}[\%] = 100 * \frac{\text{\#INSERTII} + \text{\#SUBSTITUȚII} + \text{\#ȘTERGERI}}{\text{\#NUMĂRUL TOTAL DE CUVINTE DIN TRANSCRIEREA CORECTĂ}}$$

În cazul în care, procentul de caractere eronate din cuvântul ipotetic este redus, eroarea de substituție se poate redresa mult mai ușor prin utilizarea ratei de eroare la nivel de caracter ChER (Character Error Rate). Este foarte asemănătoare cu rata de eroare la nivel de cuvânt, diferența dintre ele limitându-se la nivelul unității folosite de baza folosite în comparație: caracterul și nu cuvântul:

$$\text{ChER}[\%] = 100 * \frac{\text{\#ERORI DE INSERTIE DE CARACTER} + \text{\#ERORI DE SUBSTITUȚIE DE CARACTER} + \text{\#ERORI DE ȘTERGERE DE CARACTER}}{\text{\#CARACTERE ÎN TRANSCRIEREA CORECTĂ}}$$

2.1.7 Unelte folosite în recunoașterea vorbirii [Sphinx]

Pentru lucrarea de față am folosit setul de instrumente, pus la dispoziție în mod gratuit pentru recunoașterea vorbirii de CMU Sphinx.

Pe lângă versiunile oferite de Sphinx, există numeroase sisteme utilizate în recunoașterea vorbirii precum : HTK [Y01] , ISIP[D01], AVCSR[L01]. Pentru a facilita toate inovațiile în domeniul recunoașterii vorbirii, a fost creat un Sphinx-4: o platformă la dispoziția utilizatorului ce încorporează metodologiile de ultimă generație și acoperă și ariile în curs de dezvoltare. Include componente separate pentru fiecare sarcină specifică, iar modulele sale pot fi înlocuite în timpul rulării.

Sphinx-4 este un system de recunoaștere a vorbirii scris în limbajul de programare Java. A fost creat în urma unei colaborări dintre grupurile Sphinx de la Universitatea Carnegie Mellon, Laboratoarele Sun Microsystems, Laboratoarele Mitsubishi Electric Research Labs (MERL) și Hewlett Packard (HP), cu unele contribuții de la Universitatea din California la Santa Cruz (UCSC) și Institutul de Tehnologie Massachusetts (Massachusetts Institute of Technology (MIT)).

Principalele trei module din arhitectura Sphinx-4 sunt:

- FrontEnd
- Decodor
- Linguist

Disponând în momentul de față de cei mai mulți parametrii configurabili, de acest aspect ocupându-se ConfigurationManager. Sphinx-4 folosește pentru a putea oferi statistici în legătură cu procesul de decodare precum rata de eroare a cuvintelor, viteza de rulare, memoria utilizată diverse ustensile (din engleză Tools) și de asemenea Utilitati (din engleză Utilities) ce suportă procesarea la nivel de aplicație a rezultatelor decodării. Se pot obține grile de rezultate, scoruri de confidențialitate sau înțelegerea limbajului natural.

Scopul modulului FrontEnd este acela de a transforma un semnal de intrare (audio) într-o secvență de ieșire de trasături. Conține mai multe lanțuri paralele de procesare astfel încât se pot calcula simultan tipuri diferite de parametri, precum MFCC și PLP din același semnal de intrare sau din semnale diferite. Cu ajutorul arhitecturii FrontEnd ce conține blocuri de procesare de date (din engleză DataProcessor) și de date (din engleza Data), Sphinx-4 oferă suport pentru următoarele tehnici: citirea de la diferite formate de intrare, citirea de la dispozitivele sistemului audio pentru funcționarea în mod live, pre-accentuare, folosirea de ferestre (ex. ferestre Hamming și Hanning), transformata Fourier discretă (FFT), transformata cosinus discretă (DCT).

Modului Lingvist generează la ieșire un grafic de căutare (din engleză SearchGraph) și poate fi configurat cu diverse implementări lingvistice și este caracterizat de trei componente:

Modelul de limbă format din gramatici bazate pe grafice sau modele stocastice N-gram. Suportă formate precum: gramatică JSGF [JSGF], LMG, FST, model n-gram simplu, model trigram.

Dictionarul furnizează pronunțiile pentru cuvintele din modelul de limbă și împarte cuvintele în secvețe formate din unități găsite în modelul acustic.

Modelul acustic o mapare între unitățile de vorbire și un HMM, ce ia în calcul și poziția cuvântului și contextul. Acesta permite varierea mai multor parametri ce caracterizează un HMM printre care: mixturile gaussiene, matricile de tranziții, înălțimea gaussianelor. Sphinx-4 este capabil să implementeze un singur model acustic ce încarcă și folosește modelele acustice generate de antrenorul Sphinx-3.

Cel de-al treilea modul, blocul de decodare folosește parametrii rezultați din primul modul împreună cu SearchGraph pentru a genera ipotezele rezultate. Cea mai importantă componentă a sa o reprezintă managerul de căutare (din engleză SearchManager) și poate folosi diverși algoritmi de decodare precum Viterbi, A*, Bushderby, decodare paralelă [W01].

2.2 Spoken Term Detection [STD]

2.2.1 Introducere

Procesarea de informații a devenit o activitate principală în lumea actuală, comunicațiile vorbite constituind sursa majoră a acestor informații. Procesul este o inițiativă pentru a facilita dezvoltarea tehnologiilor de găsire a informațiilor în arhivele de înregistrări. STD în cazul limbilor bine

documentate are la bază un sistem de recunoaștere automată a vorbirii dependent de o bază mare de date de antrenare, ce include atât înregistrări (pentru modelarea acustică), cât și text (pentru modelarea de limbă). Acuratețea sistemului ASR pentru aceste limbi, precum și tehnicile îmbunătățite de indexare și căutare asigură calitatea procesului STD.

2.2.2 Sarcina procesului și arhitectura sistemului

Procesul constă în căutarea unui cuvânt rostit în interiorul unei baze de înregistrări, de asemenea rostite. Am dispus de două seturi de înregistrări multilingve: unul cu exemplele de întrebări și celălalt set cu conținutul în care are loc căutarea. Scopul unei asemenea evaluări este de a detecta rapid un termen, adică o secvență de cuvinte rostite consecutiv, în cadrul unor înregistrări.

Vom folosi arhitectura STD propusă în Campania NIST 2006 descrisă în figura 2.10, ce separă partea de indexare de cea de căutare și nu efectuează sondarea direct pentru fiecare termen. Beneficiul acestei arhitecturi este acela de a permite măsurători uniforme asupra resurselor de operare: viteza de indexare, marimea indexului, viteza de căutare, precum și efectuarea unei căutări mai rapide, deoarece indexarea o simplifică.

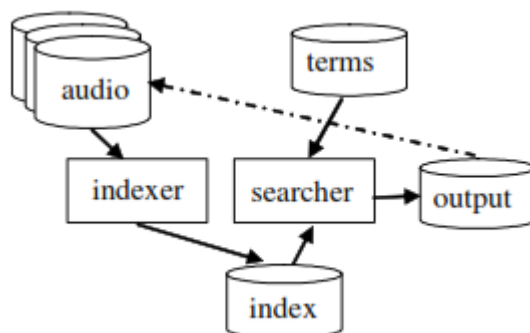


Figura 2.10 Arhitectura universală a sistemului STD [STD]

Un aspect important îl constituie timpul de căutare pentru un anumit cuvânt, care trebuie să fie mici (câteva secunde). Pentru a dezvolta acest sistem, ne bazăm pe sistemul de recunoaștere automată a vorbirii pentru limba română, dezvoltat anterior, pe care îl folosim pentru indexare, astfel încât toate înregistrările să fie transformate în șiruri de foneme. În timpul căutării, cuvântul rostit țintă este transcris de sistemul ASR și astfel vom căuta un rând de foneme asemănător cu cel pe care îl avem și cu care depășește un prag de similitudine.

2.2.3. Termenii utilizați. Fișierele de intrare și ieșire ale sistemului

Un termen se referă la un cuvânt sau un grup de cuvinte care au un singur sens în orice context, astfel încât oriunde l-am găsi, semnificația este aceeași. Deși un termen, de obicei are o singură formă de pronunție, orice variație sau model de plural ar trebuie recunoscută în cazul căutării noastre. Termenii vor fi prezentați folosind ortografia limbii native: ASCII pentru engleză, UTF-8 pentru arabă și GBK pentru limba mandarin chineză.

Termenii sunt constituiți din unul sau mai multe cuvinte. Conceptul de cuvânt ia forme diferite de la limbă la limbă.

Aparițiile de referință ale cuvintelor căutate pot fi găsite în fișierul RTTM furnizat, ce conține informații precum numele înregistrării ce conține termenul, canalul, timpul când începe exprimat în secunde, durata cuvântului.

Fișierul de ieșire va fi sub forma unui fișier text în format XML pentru fiecare combinație a categoriilor de limbi, ce va conține următoarele informații:

-limba

-tipul sursei;

-indexul timpului de generare și a măsurătorilor legate de dimensiuni;

-termenii căutați împreună cu următoarele informații pentru fiecare cuvânt: timpul de căutare și pentru fiecare presupusă găsim a cuvântului aflăm: locația (timpul unde începe și se termină în interiorul înregistrării), scorul de detecție și decizia binară luată.

Generarea la ieșire a deciziei luate împreună cu scorul de detecție pentru fiecare posibilă apariție are un beneficiu semnificativ asupra evaluării sistemului. Putem considera corectă ieșirea sistemului dacă termenii apar în transcriere ca o potrivire perfectă, corespunzând de asemenea și timpilor.

2.2.4. Metodologia evaluării

Performanțele unui system STD sunt caracterizate de două tipuri de erori: alarme false și detecții ratate și pot fi evaluate, de asemenea prin două metode:

a)grafic folosind curba erorii de detectare a compromisului (DET) ce prezintă o vedere intuitivă a performanței;

b)pentru un anumit punct de operare de pe curbă, se folosește valoarea ponderată a termenului (TWV), folosită pentru optimizarea performanței sistemului.

2.2.4.1 Curba erorii de detectare a compromisului

Performanțele de bază ale detecției sunt caracterizate în principal de curbele standard ale erorilor de detecție a compromisurilor DET (Detection Error Tradeoff) dintre probabilitatea de ratare (P_{Miss}), când un cuvânt nu este detectat deși el apare și probabilitatea de alarmă falsă (P_{FA}), în cazul în care detecția sistemului nu corespunde cu existența cuvântului. Ambele probabilități sunt funcții ce depind de pragul de detecție Θ , după cum urmează:

$$P_{Miss}(\text{cuvant}, \Theta) = 1 - \frac{N_{corect}(\text{cuvant}, \Theta)}{N_{true}(\text{cuvant})}$$
$$P_{FA}(\text{cuvant}, \Theta) = \frac{N_{fals}(\text{cuvant}, \Theta)}{NNT(\text{cuvant})}, \text{ unde}$$

- $N_{corect}(\text{cuvant}, \Theta)$ reprezintă numărul de detecții corecte ale cuvântului cu un scor mai mare sau egal ca pragul Θ ;
- $N_{fals}(\text{cuvant}, \Theta)$ este numărul detecțiilor incorecte ale cuvântului cu un scor mai mare sau egal decât pragul Θ ;

- $N_{\text{true}}(\text{cuvant})$ este numărul real al aparițiilor cuvântului căutat în context, iar $N_{\text{NT}}(\text{cuvant})$, numărul oportunităților pentru detecțiile incorecte ale cuvântului (= "Non-target"); $N_{\text{NT}}(\text{cuvant})$ se alege în mod arbitrar, proporțional cu numărul de secunde ale vorbirii în setul folosit pentru testare:
- $N_{\text{NT}}(\text{cuvant}) = n_{\text{tps}} * T_{\text{vorbire}} - N_{\text{true}}(\text{cuvant})$, unde n_{tps} este numărul de încercări de vorbire pe secundă (setat arbitrat la 1).
- T_{vorbire} reprezintă totalul discursului (exprimat în secunde).

Cele două probabilități vor fi calculate separat pentru fiecare termen și după se vor media peste toți termenii selectați cu aceeași pondere. Acest lucru este valabil doar pentru cuvintele cu un număr diferit de zero de apariții în setul de înregistrări, astfel încât P_{FA} să poată fi definită. Curbele DET se evaluează ca o funcție de limbă și tipul sursei pentru diverse înregistrări și query-uri.

2.2.4.2 Performanțele de detectare ale sistemului

A măsura "valoarea" unui sistem se referă la utilitatea acelui sistem pentru un utilizator, deoarece un sistem perfect va răspunde întotdeauna corect la o cerință, fiecare răspuns greșit sau derutant reducând această valoare alocată sistemului.

Performanța de detectare a sistemului va fi măsurată prin alocarea unui cost fiecărei recunoașteri corecte și un cost (o valoare negativă) fiecărei ieșiri incorecte. Se folosesc două definiții ale valorii totale a sistemului numite: valoarea ponderată a aparițiilor și valoarea ponderată a termenului (TWV). Valoarea ponderată a aparițiilor se calculează prin acumularea valorilor pentru fiecare detecție corectă din care se subtrage costul pentru fiecare detecție greșită, fiecare cuvânt contribuind în mod egal. Valoarea ponderată a termenilor (Term Weighted Value(θ)) se evaluează inițial prin calculul probabilităților de rată și alarma falsă, apoi cu ajutorul lor și a unei probabilități aprior alege, de calcul a valorilor specifice fiecărui cuvânt, se face o medie ce include toți termenii și se determină valoarea totală a sistemului, după formula:

$$\text{TWV}(\theta) = 1 - \text{medie}\{P_{\text{Miss}}(\text{cuvânt}, \theta) + \beta * P_{\text{FA}}(\text{cuvânt}, \theta)\}, \text{unde } \beta = \frac{C}{V}(\text{Pr}_{\text{cuvânt}}^{-1} - 1);$$

Folosim raportul cost/valoare, C/V egal cu 0.1 și probabilitatea anterioră a cuvântului $\text{Pr}_{\text{cuvânt}}^{-1}$ este 10^{-4} . Valoarea maximă posibilă va fi de 1.0 și corespunde unui sistem "perfect": fără ratări și fără alarme false. Pot fi posibile și valorile nule (când sistemul nu obține la ieșire nimic) sau chiar valori negative. Față de valoarea ponderată a aparițiilor, ce a termenilor prezintă un avantaj fiind mai puțin susceptibilă la frecvența de apariție a cuvintelor, excluzându-le pe cele care nu au nici o apariție. Ieșirea sistemul trebuie să includă de asemenea o indicație binară a detecției, aleasă astfel încât să maximizeze valoarea de ieșire a sistemului.

Întrucât curbele DET reprezintă performanța pentru toate valorile posibile ale pragului θ , doar două puncte prezintă interes deoarece arată dacă pragul de decizie actual al sistemului este optim. Primul este valoarea ponderată actuală a termenilor (ATWV), iar al doilea punct este reprezentat de valoarea ponderată maximă a termenilor (MTWV), ce corepunde cu locul de pe curbă unde pragul θ coincide cu valoarea TWV maximă.

Metrica principală folosită pentru comparație a fost Actual Term Weighted Value (ATWV) și reprezintă abilitatea sistemului de a preciza punctul optim de operare pe baza scorului. Diferența dintre ATWV și MTWV indică beneficiile pe care le vom avea dacă alegem un prag de decizie mai bun.

2.2.5 Procesul de căutare

În cazul în care am obține pentru sistemul de recunoaștere ASR o precizie de 100%, procesul de detectare a cuvântului rostit s-ar reduce la căutarea unui șir de caractere într-un conținut textural și ar constitui o problemă simplă. Deoarece nu suntem aproape de acest caz ideal, vom căuta în baza noastră de înregistrări un șir similar cuvântului căutat. Căutarea exactă a unui cuvânt are rezultate STD foarte slabe: o probabilitate de alarmă falsă (FAP) de 0.1% și o probabilitate de rată (MP) de 99%. Cum am precizat anterior, erorile în recunoașterea limbajului vorbit la nivel de fonem pot fi substituții, inserții și ștergeri. Aceste erori sunt numărate automat prin aplicarea algoritmului DTW, unul din cei mai vechi și mai importanți algoritmi folosiți în recunoașterea vocală. Numărul mare de cuvinte ce nu fac parte din vocabularul (OOV) ce par a face parte din termenii de interogare va crește numărul de erori.

2.2.5.1 Etapele parcurse în sistemul de căutare

O metodă simplă de recunoaștere a unui cuvânt izolat este să fie comparat cu un număr predefinit de șabloane, dintre care se alege acela care se potrivește cel mai bine.

Apar însă, în această situație, două probleme:

- ce formă trebuie să ia șabloanele;
- cum vom compara noi cuvântul cu semnalul vocal.

a)Prima etapă în procesul de recunoaștere este analiza vorbirii, în timpul căreia semnalul este procesat, astfel încât să extragem cele mai importante caracteristici, numite trăsături sau parametri. Precum am detaliat în partea de recunoaștere automată a vorbirii, am folosit coeficienți cepstrali MFCC. Prin folosirea doar a celor mai importante caracteristici ale vorbirii, cantitatea de date utilizată pentru comparații este mult redusă și prin urmare vom avea mai puține comparații și nevoie de un timp mai scurt.

b)A doua etapă a constat în clusterizarea acestor coeficienți, adică procesul de grupare a unor obiecte în grupuri numite clusteri în așa fel încât acestea să fie similare, găsind practic reprezentanții unor grupuri omogene.

Aceast algoritm se realizează parcurgând următoarele etape:

- Gruparea în K-means: partiționăm vectorul de intrare în k seturi prin selecție aleatoare sau pe baza distanței maxime), așa cum este arătat și în figura 2.11 :

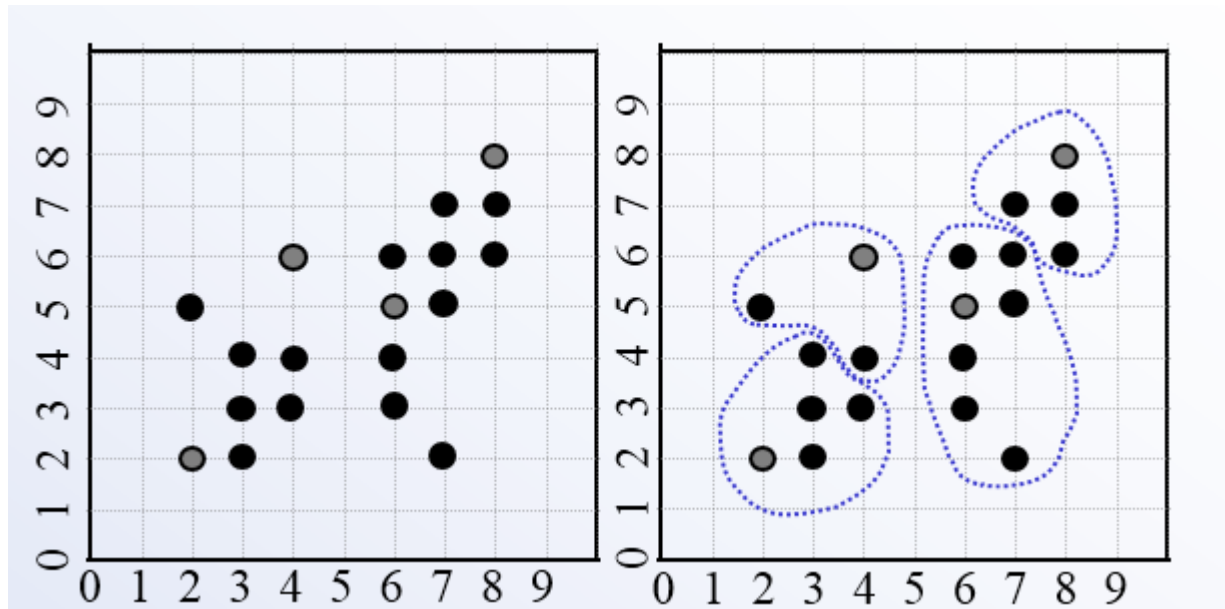
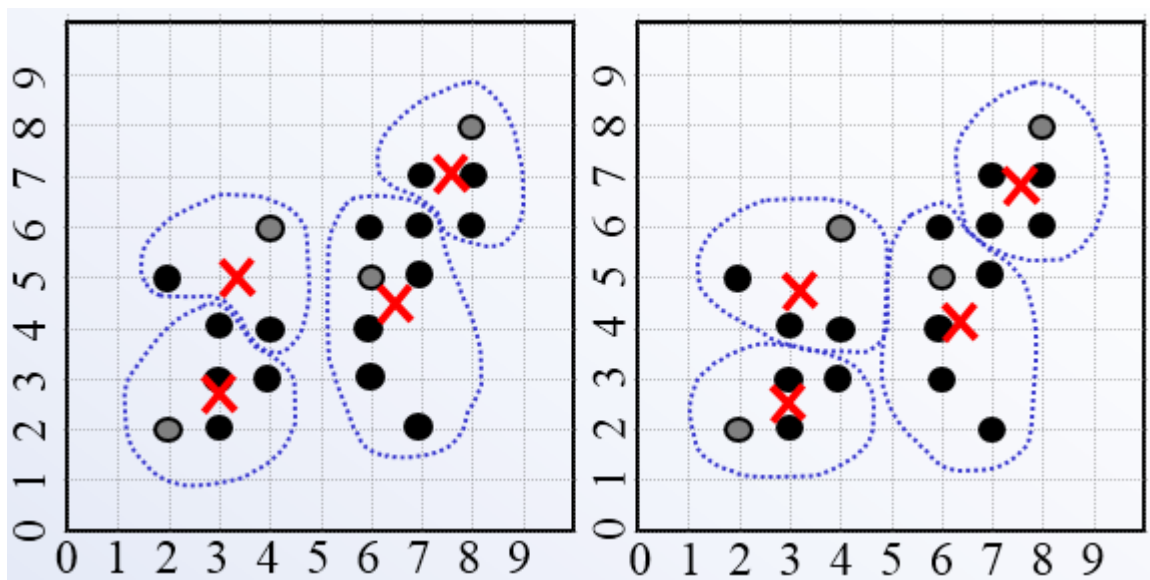


Figura 2.11 Gruparea punctelor în patru clustere

- Căutarea: pentru fiecare vector de antrenare găsim cel mai apropiat cuvânt de cod și îi alocăm grupului (clusterului) de care aparține acel cuvânt.
- Actualizarea centroidului: pentru fiecare grup, recalculăm centroidul, luând în calcul și noua observație adăugată;
- Se repetă ultimii doi pași până când distanța medie scade sub un anumit prag sau nu se mai modifică, figura 2.12.



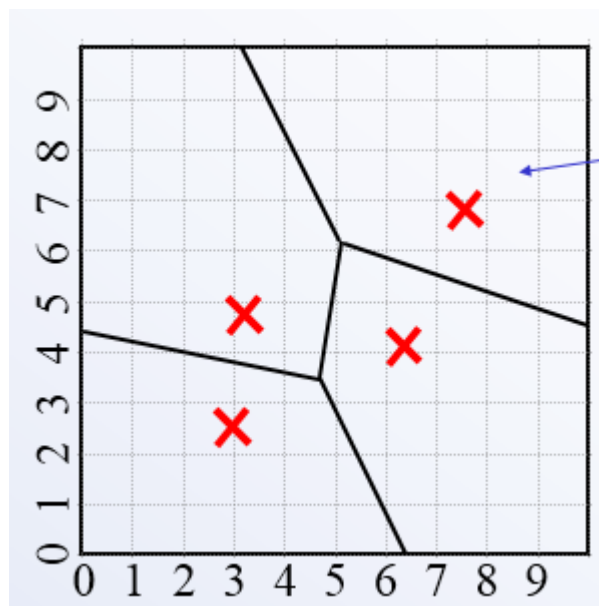


Figura 2.12 Calculul centrozilor fiecărui cluster până în momentul convergenței algoritmului

Algoritmul converge destul de repede în practică și returnează cele mai bune cluster găsite prin execuția a mai multor iterații. Soluția depinde de numărul setului inițial de cluster, ce trebuie definit înainte de a-l parcurge.

c) Împărțim fiecare înregistrare în ferestre proporționale cu lungimea termenului căutat, de 1.5 ori mai mari.

d) Am folosit algoritmul DTW pentru a găsi alinierea neliniară optimă în timp pentru secvența de simboluri obținute din procesarea textului și secvența de segmente de voce înregistrate. Principiul de funcționare este următorul: se creează o matrice având ca dimensiuni lungimile celor două secvențe care trebuie comparate. În fiecare celulă a matricii se pune o valoare numerică dată de costul dintre elementele corespunzătoare ale celor două secvențe, calculat cu ajutorul distanței euclidiene.

Diferențele dintre formele de undă a unor cuvinte diferite pot fi văzute în figura 2.13.

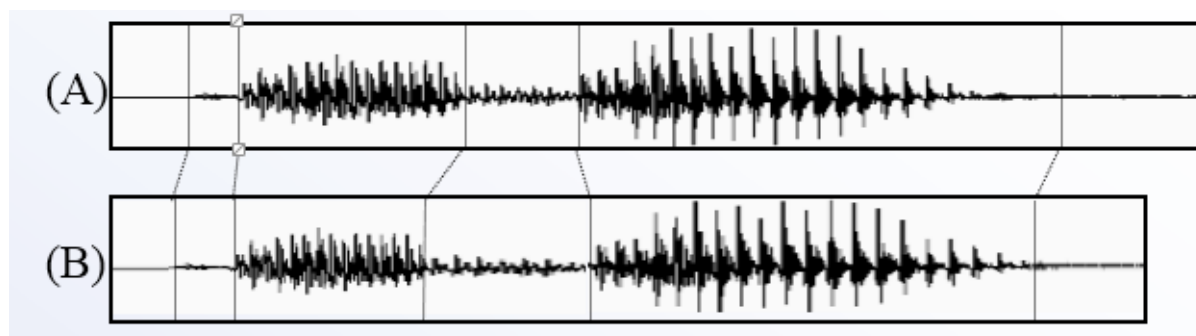


Figura 2.13 Formele de undă a două semnale diferite

Dacă A și B ar fi identice, drumul de cost minim ar fi o linie dreaptă ce trece prin 0 (diagonala principală). Deoarece cele două semnale nu sunt la fel, calea prin matrice arată cea mai bună potrivire între cadrele dintre A și B.

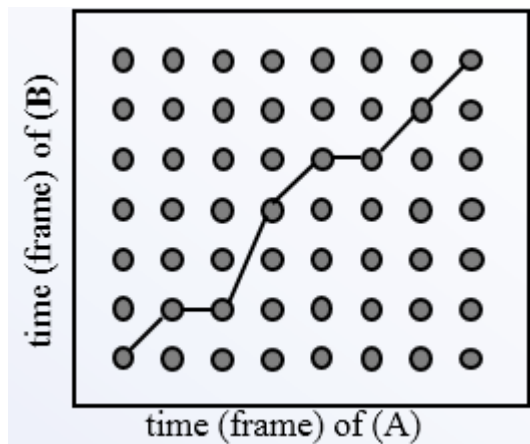


Figura 2.14 Rezultatul rulării algoritmului DTW

Vectorul de aliniere este reprezentat de această cale între fonemele de pe abscisă ce aparțin cuvântului original din înregistrare (cel de referință) și cele ale cuvântului căutat de pe orizontală. Costurile au fost alese astfel încât să se poată acoperi toate tipurile de erori (substituiri, ștergeri și inserții). Tranzițiile orizontale corespund ștergerilor. Cele verticale, care corespund inserțiilor, au un cost mai mic decât cele orizontale, adică două foneme au fost recunoscute în locul uneia. De multe ori, fonemul recunoscut este cel corect și încă unul, sau cel corect de două ori, ca în figura 2.15.

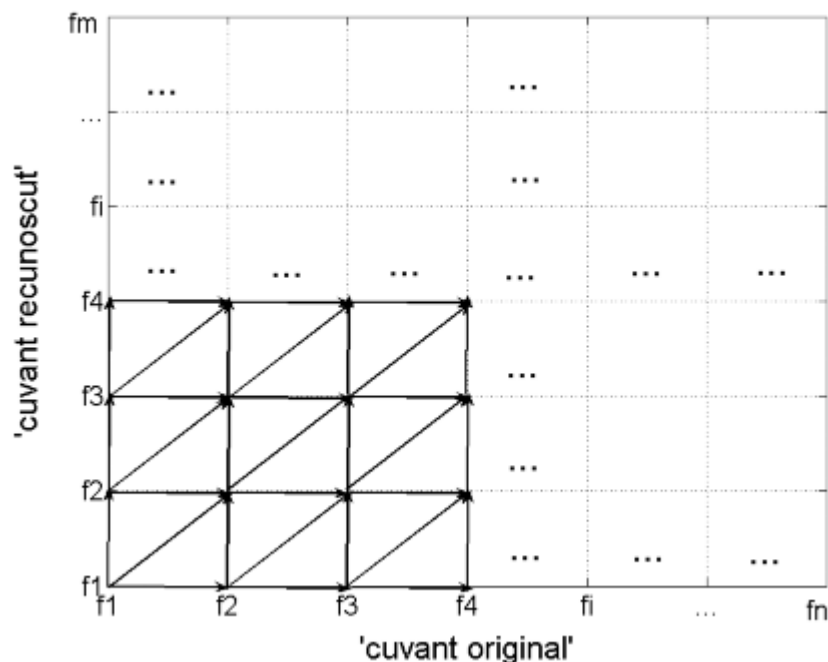


Figura 2.15 Alinierea cuvintelor cu algoritmul DTW [B01]

2.3 Detecția Genului

2.3.1 Introducere

Tehnologiile de recunoaștere a vorbitorului au fost studiate activ de câteva decenii, iar rezultatele au ajuns la crearea unor sisteme de ultimă generație, astfel încât telefoanele bazate pe servicii integrate cu recunoaștere a vorbirii, a vorbitorului sau a limbii vor suplimenta sau chiar înlocui serviciile cu operatori umani.

Scopul acestei aplicații este acela de a identifica genul vorbitorului pe baza vocii sale. Diferențele de gen în vorbirea umană se datorează diversității fiziologice precum lungimea tractului vocal, grosimea vocii sau stiluri de vorbire diferite, aspect evidențiat în figura 2.16.

Principiile recunoașterii genului unui vorbitor sunt aceleași ca ale sistemului recunoașterii automate ASR, constituit din două părți: una de admitere în care este creat un model al vorbitorului cu ajutorul unui model de fundal și o parte de recunoaștere în cadrul căreia luăm o decizie binară.



Figura 2.16 Cuvântul ,minge' rostit de doi vorbitori diferiți: a) femeie și b) bărbat

Figura arată de ce simpla abordare de recunoaștere a semnalelor bazându-ne doar pe forma similară a sunetelor nu este de ajuns pentru sarcinile de recunoaștere a vorbirii. Jumătatea superioară a figurii arată forma de undă a semnalului brut, iar cea inferioară afișează versiunile prelucrate ale semnalelor evidențiind formatele specifice. Se poate observa că vorbirea este dependentă de persoană, lucru ce nu mai reprezintă o surpriză, deoarece vocile diferite sună diferit.

Principala caracteristică ce deosebește cele două genuri o reprezintă frecvența fundamentală F_0 , ce are valoare tipică în cazul vorbirii masculine de 110Hz și 200Hz în cazul celei feminine. Înălțimea vocii în cazul copiilor este atât de variată, încât aceștia pot fi catalogați ca un al treilea gen. Valorile uzuale ale acestei frecvențe pentru persoanele cu vârsta cuprinsă între 20-70 ani sunt cuprinse între 80-170 Hz pentru bărbați, 150-260 Hz pentru femei și 300-500 pentru copii.

Detectarea automată a genului unui vorbitor poate prezenta multiple avantaje:

- Identificarea genului și eliminarea componentelor caracteristice ce va conduce la o rată de compresie mai mare a semnalului, o cantitate mai mare de informații ce pot fi transmise și o economie asupra benzii ;
- Sortarea apelurilor telefonice după gen în cazul sondajelor;
- Facilitarea sistemelor de recunoaștere a vorbitorului prin împărțirea vorbirii în jumătăți, ce va conduce la reducerea calculelor și îmbunătățirea vitezei sistemului.

2.3.2 Uneltele folosite

Laboratoarele de informatică ale Universității din Maine (din limba franceză The Laboratoire d'Informatique de l'Université du Maine (LIUM)) au dezvoltat diverse pachete software scrise în limbajul de programare Java ce pot fi folosite în mai multe domenii precum: detecția vorbirii/ liniștii, detecția vorbitorului, calcularea coeficienților MFCC, segmentarea și clusterizarea înregistrărilor.

LIUM oferă unelte atât pentru dezvoltatori ce îmbunătățesc performanțele WER, cât și pentru utilizatori în cazul detecției spontane a vorbirii, cercetări în domeniul recunoașterii, traducere automată. LIUM_SpkDiarization conține un set de unelte ce pot fi folosite pentru procesul cunoscut sub numele din engleză diarization, în care se pleacă din semnalul audio al utilizatorului pe care îl clusterizăm pe baza diverselor metrici. Setul a fost dezvoltat pentru campania de evaluare French ESTER2, cele mai bune rezultate obținându-se pentru spectacolele de televiziune și radio.

Segmentarea se referă la împărțirea fișierelor audio în diverse părți omogene și flexibile, iar procesul de diarization semnifică identificarea unică a vorbitorilor într-un fișier. Procesul poate îmbunătăți transcrierea automata, prezentând interes și în indexarea documentelor multimedia.

2.3.3 Descrierea sistemului [Diarizare]

LIUM utilizează în aplicațiile asociate vorbitorilor metode de grupare a înregistrărilor (în engleza diarization), proces ce constă, în terminologia NIST, în asocierea unor regiuni temporale din înregistrările audio cu etichetele corespunzătoare. Etichetele sunt formate din numele vorbitorilor, genul lor, tipul canalului utilizat (de banda largă sau banda îngustă), mediu din fundal (liniște, muzică, zgomot) sau orice altă caracteristică a semnalului. Scopul este acela de a detecta segmentele audio omogene, fiecare conținând vocea unui singur vorbitor [S01].

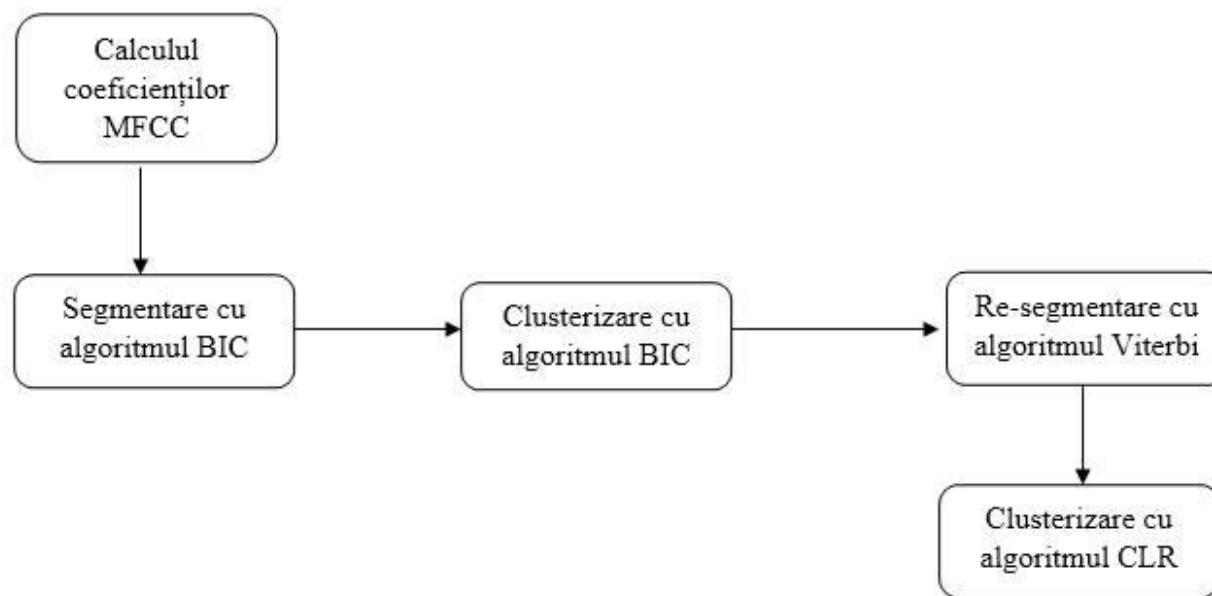


Figura 2.17 Arhitectura sistemului de diarization

Etapele care stau la baza sistemului sunt cele din figura 2.17:

1) Cadrele acustice sunt formate, în mod identic cu capitolul 2.1., și rezultă 13 coeficienți MFCC ce conțin și energia C_0 , după care îndepărtăm coeficientul C_0 și calculăm derivatele de ordin întâi ale acestora.

2) Segmentarea este nucleul procesului de diarization și are ca obiectiv divizarea fișierului audio în segmente omogene pentru vorbitori. Are loc în două faze: prima detectează automat punctele de schimbare ce corespund marginilor segmentelor pe baza unui raport de probabilitate generalizat GLR (din engleza Generalized likelihood ratio), calculat cu ajutorul gaussianelor cu matrici de convergență complete. Aceste gaussiene sunt estimate pentru ferestre glisante de cinci secunde de-a lungul întregului semnal. Se estimează curba distanței GLR pentru aceste ferestre, cele mai mari puncte corespunzând schimbării între vorbitori. Următoarea fază folosește distanțe BIC cu scopul de a asocia segmentele consecutive găsite ca aparținând aceluiași vorbitor. Se pot folosi și alte metode de măsurare precum KL2, divergența gaussiană și ΔBIC . BIC a fost introdus pentru prima dată de Chen și Gopalakrishnan [C01] și reprezintă un model de selecție, asimptotic, optimal, bazat pe criteriul Bayesian folosit pentru a decide care dintre cele N_c modele parametrice reprezintă cât mai bine M eșantioane ale datelor $x_i \in R^d$, $i=1,2,\dots,M$. Eșantioanele x_i reprezintă vectori, d -dimensionali ce conțin trasături extrase la punctul anterior 1). În cazul segmentării vorbitorului, algoritmul se reduce doar la două modele angajate ($N_c=2$), adică se analizează doar două ferestre vecine X și Y în jurul timpului t_j , cu scopul de a decide dacă t_j este sau nu, un punct de schimbare al vorbitorului.

3) Toate segmentele sunt grupate în mulțimi (din engleza cluster), ce le sortează inițial după nume și apoi după momentul de început al înregistrării. Erori precum a avea două mulțimi distincte ce corespund aceluiași vorbitor sau asocierea segmentelor a doi vorbitori într-o singură mulțime au o penalitate destul de mare în metrica de evaluare NIST. Există mai mulți algoritmi de grupare, secvențiali sau ierarhici, cei din urmă divizându-se la rândul lor în două categorii (aglomeratici sau divizivi).

Algoritmul de clusterizare diviziv aglomerativ pornește de la un număr N de clustere, fiecare cluster conținând un singur segment, cu propriul lui model antrenat de o gaussiană cu matricea de covarianța completă. La fiecare pas sunt calculate distanțele dintre perechile de clustere, pe baza cărora clusterelor care se aseamnă cel mai mult sunt unite într-o mulțime nouă, căreia îi antrenăm propriul model. Algoritmul se încheie când întâlnim un criteriu de stopare. Există diverse metrice și criterii de stopare propuse în [Siu1991, Solomonoff1998, Chen1998a, Reynolds1998]. Unul dintre acestea este metrica Delta BIC (din engleza Bayesian Information Criterion) ce poate atât selecta clusterelor în grupuri cât și încheie procesul, atunci când $\Delta BIC_{i,j} > 0$, i și j reprezentând două clustere apropiate. Se pot folosi două tipuri de metrice Delta BIC, ce diferă prin factorul de penalizare:

Factorul de penalizare local P_l ce depinde de n_i și n_j , lungimile clusterului i , respectiv al clusterului j prin formula:

$$P_l = \frac{1}{2} \left(d + \frac{d(d+1)}{2} \right) + \log(n_i + n_j)$$

Factorul de penalizare global P_g ce folosește lungimea întregului semnal, n :

$$P_g = \frac{1}{2} \left(d + \frac{d(d+1)}{2} \right) + \log(n)$$

Experimentele au dovedit, că factorul local dă rezultate mai bune.

4) Cea de-a patru etapă constă într-o decodare cu algoritmul Viterbi pentru a genera o segmentare nouă, ce va îndepărta porțiunile de muzică sau zgomot. Clusterelor create sunt modelate de un model Markov ascuns HMM cu o singură stare, reprezentate de un model de mixturi gaussiene cu opt componente. Segmentele de lungime mare sunt tăiate în punctele cu cea mai mică energie pentru a obține doar segmente mai scurte de 20 secunde.

5) Ultimul pas constă în clusterizare folosind algoritmul CLR (Cross Likelihood ratio), cel de-al doilea algoritm ierarhic, aglomerativ. Trăsăturile folosite nu au fost normalizate cu scopul de a păstra informațiile mediului din fundal, ce ajută în diferențierea vorbitorilor. În momentul de față, fiecare cluster conține vocea unui singur utilizator, dar îi putem asocia acestuia mai multe. Pentru a putea efectua și ultima operație din proces, trebuie să eliminăm influența mediului din modelele clusterelor, prin normalizarea coeficienților. Astfel vom obține o relație de unu-la-unu între vorbitori și cluster.

La fiecare iterație, două cluster care maximizează rata scorului de încrucișare (CLR) [C01] sunt unite într-unul singur. Algoritmul se încheie când unitatea de măsură depășește un prag setat anterior.

2.3.4 Modelarea vorbitorului

Un sistem deterministic pas-cu-pas de diarization a fost dezvoltat de Meignier et al. și este folosit în detecția de gen și bandă. Acesta estimează automat numărul de vorbitori și se realizează având la bază modelele gaussiene GMM (cu 128 de componente diagonale) pentru fiecare din cele patru combinații de gen (feminin/masculin) și bandă îngustă/bandă largă. Fiecare cluster este etichetat în concordanță cu caracteristicile modelelor gaussiene care maximizează probabilitatea coeficienților conținuți. Fiecare model a fost extras dintr-o ora de înregistrare a fișierelor de antrenare date de ESTER. În urmă procesul de diarization au fost obținute segmente mai scurte de 20 de secunde, ce conțin vocea unui singur vorbitor, cu banda și genul cunoscut pentru fiecare segment în parte. Din fiecare înregistrare este eliminată liniștea și momentele de tăcere.

După ce extragem vectori cu coeficienți din înregistrările vorbitorului putem antrena un model al acestuia și să îl stocăm în baza de date a sistemului. Întrucât folosim un model independent de vorbitor ce include o corespundență foarte mică sau chiar deloc între cadrele înregistrărilor de referință și cele de testare, alinierea la nivel de cadru nu este posibilă. Prin urmare, este utilă partiționarea semnalului în foneme sau clase mari de foneme, adică structurarea fonetică a modelului fiecărui vorbitor.

În aplicațiile de vorbire un rol important îl are adaptarea modelelor acustice la condițiile noi de operare, deoarece există o varietate mare de vorbitori, de stiluri, medii de înregistrare [Gender].

Prin antrenarea unui GMM ne referim la estimarea parametrilor $\lambda = \{P_k, \mu_k, \Sigma_k\}_{k=1}^K$, pornind de la eșantioane de antrenare $= \{x_1, x_2, \dots, x_T\}$, și folosind estimarea probabilității maxime. Cu cât este mai mare valoarea acestei probabilități, cu atât putem afirma mai corect că vectorii necunoscuți provin din modelul λ . Unul din algoritmi de bază este EM (din engleză Expectation-maximization), ce poate fi

folosit pentru a maximiza această probabilitate având anumite date inițiale. Acesta pornește inițial de la o gaussiană învățată pentru toate datele de antrenare, după care, iterativ, gaussienele sunt împărțite și antrenate până ajungem la numărul de componente dorit. În general, sunt suficiente cinci iterații pentru ca parametri să convergă, dar ne oprim în momentul în care câștigul de probabilitate dintre două iterații este mai mic decât un anumit prag. În sistemele de recunoaștere independente de vorbitor, bazate pe GMM-uri, se folosește un model universal de fundal (universal background model-UBM) ce este antrenat pornind de la o bază de înregistrări de zeci sau sute de ore, aparținând mai multor persoane. Modelul UBM este reprezentat de distribuția vectorilor de coeficienți, independent de vorbitor, iar în momentul în care un membru nou este introdus în sistem, parametri acestui model sunt adaptați la trăsăturile sale. Modelul adaptat rezultat va fi folosit ca model al vorbitorului. Având datele necesare pentru a antrena un model universal, există mai multe abordări posibile:

- Prima metoda constă în unirea tuturor înregistrărilor și antrenarea lor cu un algoritm EM (Expectation-maximization), cu observația ca datele să fie repartizate echilibrat între subgrupe. De exemplu, dacă vrem un model independent de date este necesar ca înregistrările de voci feminine și masculine să fie balansate, pentru a nu înclina spre direcția dominantă.
- Cea de-a doua metodă se referă la antrenarea modelelor UBM pentru fiecare subgrupă în parte și unirea modelului final. În acest caz, putem folosi și înregistrări neechilibrate între grupuri.

Putem observa cele două abordări în figura 2.18 :

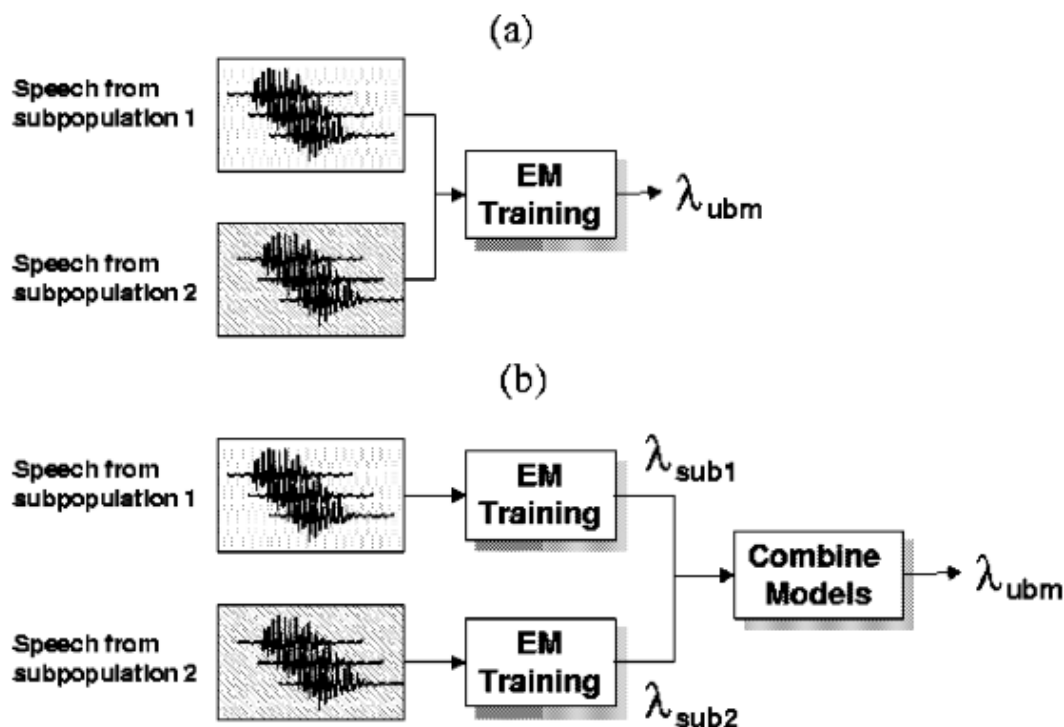


Figura 2.18 Abordările pentru crearea unui model UBM

Pentru aplicația noastră, este utilă antrenarea a două modele de fundal diferite: unul pentru vorbitorii feminini și altul pentru cei masculini, situație în care vom antrena modelul noului venit în funcție de genul său.

Putem adapta toți parametrii sau doar o parte din ei pornind de la modelul de fundal.

Având eșantioanele $\mathcal{X}=\{x_1, x_2, \dots, x_T\}$ și parametrii UBM (din engleza universal background model) $\lambda_{\text{UBM}} = \{P_k, \mu_k, \Sigma_k\}_{k=1}^K$, obținem vectorii mediilor adaptate (μ_k') cu metoda MAP (din engleza maximum a posteriori) ca o sumă ponderată a datelor de antrenare ale vorbitorilor și media UBM:

$$\mu_k' = \alpha_k x_k + (1 - \alpha_k) \mu_k, \text{ unde } (2.3.3)$$

$$\alpha_k = \frac{n_k}{n_k + r} \quad (2.3.4)$$

$$n_k = \sum_{t=1}^T P(k|x_t) \quad (2.3.5)$$

$$x_k = \frac{1}{n_k} \sum_{t=1}^T P(k|x_t) x_t \quad (2.3.6)$$

Pentru fiecare mixtură și fiecare parametru s-a folosit un coeficient de adaptare, dependent de înregistrări α_k , cu μ_k vectorul mediilor, r parametrul de legătură fixat și n_k un număr probabilistic.

În procesul de recunoaștere, modelul adaptat cu ajutorul algoritmului maxim a posteriori (MAP), ce implică derivarea unui model GMM specific vorbitorului dintr-un model UBM, este asociat cu modelul UBM și denumit în mod obișnuit *Gaussian mixture model-universal background model* sau modelul GMM-UBM.

Scorul de potrivire depinde atât de modelul țintă (λ_{tinta}) cât și de modelul de fundal (λ_{UBM}) prin intermediul logaritmului probabilității:

$$\text{Cost}(\mathcal{X}, \lambda_{\text{tinta}}, \lambda_{\text{UBM}}) = \frac{1}{T} \sum_{t=1}^T \{\log p(x_t | \lambda_{\text{tinta}}) - \log p(x_t | \lambda_{\text{UBM}})\}$$

și masoară diferența dintre cele două modele privind generarea observațiilor $\mathcal{X}=\{x_1, x_2, \dots, x_T\}$;

Atunci când indentificarea vorbitorului se face într-un cadru strâns, se poate folosi doar GMM-uri fără a fi nevoie neapărat de modelul UBM, însă tehnica nu este la fel de precisă în modelarea trăsăturilor.

2.3.5 Decodarea

Problema decodării coincide cu problema recunoașterii. Deși la recunoaștere ne interesează doar modelul cel mai probabil, există situații (cum este cel al vorbirii continue) în care avem nevoie să cunoaștem întreaga secvență de stări. Prin vorbire continuă ne referim la cazul în care trebuie să găsim pentru o secvență de observații cea mai probabilă secvență de modele, necunoscând granițele între cuvinte. În acest scop, se utilizează algoritmul Viterbi, folosit de asemenea în capitolul 2.1.2.1 în sistemele de recunoaștere automată a vorbirii, unde înlocuim funcția de însumare cu una de maximizare. Acest algoritm a fost aplicat pentru prima oară în probleme de vorbire de Vintsyuk (1968). Fiecare stare a unui HMM [P01] reprezintă un cluster, iar funcțiile densității de probabilitate ale fiecărui astfel de cluster sunt modelate de GMM-uri. Procesul se rezumă la evaluarea modelelor HMM rezultate în urma transformării unui set de mixturi gaussiene, adică probabilitatea maximă de generare a secvenței de observații de la momentul 1 la momentul t :

$$\alpha_j(t) = \max_i \{ \alpha_i(t-1) a_{ij} \} b_j(o_t) \quad (2.3.8), \text{ unde,}$$

$$\alpha_1(1) = 1 \quad (2.3.9)$$

$$\alpha_j(1) = a_{ij} b_j(o_1) \quad (2.3.10)$$

a_{ij} este probabilitatea de a trece din stare i în stare j , iar $b_j(o_t)$ reprezintă probabilitatea să se emită simbolul o_t la starea j .

Relația care stă la baza algoritmului Viterbi este:

$$\mathcal{A}_j(t) = \max_i \{ \mathcal{A}_i(t-1) + \log(a_{ij}) \} + \log(b_j(o_t)) \quad (2.3.11)$$

Acesta găsește drumul optim într-o rețea de noduri construită între ferestrele de timp și stările modelelor Markov, adică costul celei mai probabile căi (figura 2.19). Termenul $\log(b_j(o_t))$ din formula (2.3.11) reprezintă scorul nodului, iar termenul $\log(a_{ij})$ scorul tranziției.

Algoritmul este realizat astfel:

Pasul 1: Inițializarea

$$\begin{aligned} \mathcal{A}_1(1) &= \pi_i b_i(o_1) \\ B_i(1) &= 0 \quad 1 \leq i \leq N \end{aligned}$$

Pasul 2: Inducția

$$\begin{aligned} \mathcal{A}_j(t) &= \max_i [\mathcal{A}_i(t-1) a_{ij}] b_j(o_t) \quad 2 \leq t \leq T; 1 \leq j \leq N \\ B_j(t) &= \arg \max_i [\mathcal{A}_i(t-1) a_{ij}] \quad 2 \leq t \leq T; 1 \leq j \leq N \end{aligned}$$

Pasul 3: Încheierea

$$\begin{aligned} \text{scorul cel mai bun} &= \max(\mathcal{A}_j(t)) \\ s_T^* &= \arg \max_i B_i(t) \end{aligned}$$

Pasul 4: Backtracking

$$s_t^* = B_{s_{t+1}^*}(t+1) \quad t=T-1, T-2, \dots, 1$$

$S^* = (s_1^*, s_2^*, \dots, s_T^*)$ este cea mai bună secvență de stări

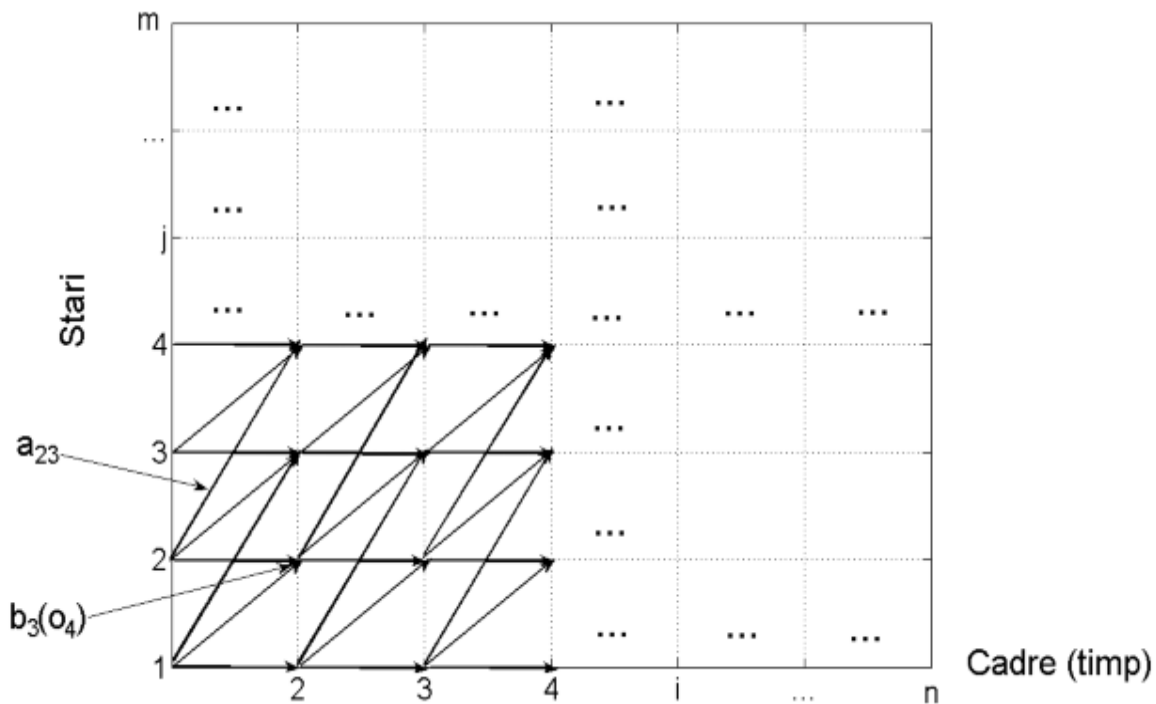


Figura 2.19 Algoritmul Viterbi [B01]

Întrucât foarte multe probabilități au valori mici, se aplică un logaritm asupra acestora ($\log P(w)$), ceea ce conduce la adunarea costurilor și nu înmulțirea probabilităților. În final, drumul va consta în alegerea celui mai mare cost și nu a celei mai mari probabilități.

Algoritmul Viterbi realizează implicit alinierea semnalului vocal cu modelele recunoscute. Rezultatul decodării va consta în găsirea secvenței de stări HMM ce a generat semnalul de voce și returnarea granițelor care indică segmentele ocupate de fiecare din cele patru modele în semnal vocal. Complexitatea acestuia depinde de numărul de stări pe care le au HMM-urile, de aceea se preferă soluția cu numărul de stări cel mai mic. Calea cea mai bună de-a lungul stărilor poate fi definită și printr-un proces de diarization ce folosește numele GMM-urilor pentru a înlocui numele grupurilor, iar stările aliniate reprezintă segmentele din fiecare grup.

2.3.6 Evaluarea performanțelor

Performanțele sistemului de detecție a genului sunt măsurate printr-o asociere de una-la-unu între ID-urile vorbitorilor de referință și cei ipotetici. Putem analiza și performanțele proceselor intermediare ce compun sistemul final. Astfel, metrica de bază în cazul acțiunii de diarization o reprezintă rata de eroare DER (din engleză diarization error rate), suma erorilor de ratare, alarma falsă sau eroare a vorbitorilor. Pentru performanțele metodelor de clusterizare se analizează puritatea grupului, adică raportul între numărul de cadre ale vorbitorului dominant și numărul total de cadre din cluster și acoperirea grupului care se ocupă de dispersia vocii unui anumit vorbitor de-a lungul clusterelor.

CAPITOLUL 3

EXPERIMENTE

3.1. Recunoașterea secvențelor audio ce conțin cifre

3.1.1. Construcția unui sistem de recunoaștere a vorbirii continue [ASR]

Sistemul de recunoaștere a clipurilor audio ce conțin cifre este un sistem de recunoaștere a vorbirii continue cu vocabular redus ce conține doar zece cifre. Din figura 2.1 a arhitecturii generale a unui sistem de recunoaștere automată a vorbirii (RAV) reiese faptul că avem nevoie de două procese [Speech]:

- a) extragerea de parametrii vocali din semnalul vocal (folosind mesajul vorbit) utilizați apoi în procesul de recunoaștere;
- b) utilizarea unor modele (acustic, fonetic și lingvistic) dezvoltate anterior.

Pentru cazul nostru, sistem de recunoaștere cu vocabular redus, se folosesc modelele de limbă de tip gramatical cu stări finite, deci nu avem nevoie de resurse textuale pentru a le forma.

La baza construcției unui asemenea sistem stau următorii pași:

- **Înregistrarea unei baze de date de clipuri audio**

Pentru a forma modelul acustic avem nevoie de un set de clipuri audio înregistrate. Astfel am înregistrat pentru fiecare vorbitor diferit din cei 69 folosiți, 100 de clipuri audio a câte 12 cifre fiecare.

- **Crearea dicționarului fonetic și a listei de foneme**

Dicționarul fonetic specifică modul în care sunt pronunțate cuvintele unei limbi. El va conține, pentru sarcina noastră, cele zece cifre ale sistemului zecimal: zero, unu, doi, trei, patru, cinci, șase, șapte, opt, nouă. Sistemul de recunoaștere CMU Sphinx folosește dicționarele fonetice în format text, cu câte un cuvânt pe fiecare linie, însă transcrierea fonetică a unui cuvânt apare pe aceeași linie. Fonemele le scriem separate prin spațiu. Lista de foneme se crează prin scrierea tuturor fonemelor ce vor fi modelate de modelul acustic pe câte o linie, în ordine alfabetică.

- **Antrenarea modelului acustic**

Necesită existența următoarelor resurse:

- baza de înregistrări audio ce conține vorbire;
- dicționarul fonetic cu toate cuvintele;
- un dicționar cu elemente ce nu sunt foneme (muzică, tuse, râset etc).

Se crează o listă cu clipurile audio ce vor fi folosite în procesul de antrenare și un fișier cu transcrierea textuală a acestor înregistrări. Vom folosi primele 50 și ultimele 30 din clipurile audio ale fiecărui vorbitor. Pentru orice proiect CMU Sphinx există un fișier de configurare unde modificăm parametrii modelului acustic și locația resurselor necesare. Avem nevoie apoi de coeficienții MFCC corespunzători clipurilor de antrenare. Astfel creăm 80 de fișiere cu coeficienți cepstrali pentru fiecare vorbitor și antrenăm modelul.

- **Crearea modelului de limbă (gramaticii)**

Vom construi pentru o sarcină de transcriere a secvențelor audio ce conțin cifre o gramatică cu stări finite. Presupunem că vorbitorul poate rosti numai cuvintele zero, unu, doi, trei, patru, cinci, șase, șapte, opt, nouă. Definirea grafului va începe cu definirea nodurilor de intrare și ieșire din gramatică, notate cu N (null), deoarece trecerea prin ele nu generează afișarea niciunui cuvânt. Creăm apoi nodurile ce corespund cuvintelor și le legăm de nodurile de intrare și ieșire cu arce direcționate. A fost creată o gramatică dacă se pronunță o cifră o singură dată. Adăugăm o tranziție de la nodul de ieșire la cel de intrare și încă două stări de null, pentru ca celelalte două acum au arce care intră și ies din ele.

Figura 3.1 prezintă grafic gramatica cu stări finite pentru sarcina propusă. Modelul e format din 14 noduri, dintre care 10 reprezintă cuvintele (cifrele de la zero la nouă). Tranzițiile arată modul în care se parcurge gramatica și ce este permis ca secvențe de cuvinte:

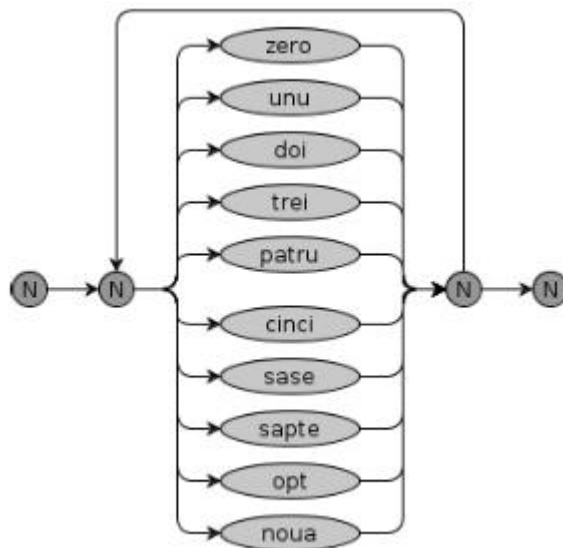


Figura 3.1 Gramatica cu stări finite a sarcinii de recunoaștere de cifre

Cifrele și succesiunile de cifre din înregistrările noastre apar cu probabilități aproape egale și din această cauză se folosește gramatica cu stări finite (FSG-Finite State Grammar) ce se poate implementa cu ajutorul formatului Java Speech Grammar (JSGF) [JSGF].

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

Tranzițiile, nodurile și probabilitățile de tranziție ale gramaticii cu stări finite sunt specificate în formatul intern folosit de CMU Sphinx, FSG astfel:

```
FSG_BEGIN <rodigits.numbers>
NUM_STATES 18
START_STATE 0
FINAL_STATE 1
TRANSITION 0 2 1.000000
TRANSITION 2 4 0.500041
TRANSITION 2 5 0.500041
TRANSITION 3 1 1.000000
TRANSITION 4 3 1.000000
TRANSITION 5 7 0.100000 nouă
TRANSITION 5 8 0.100000 opt
TRANSITION 5 9 0.100000 șapte
TRANSITION 5 10 0.100000 șase
TRANSITION 5 11 0.100000 cinci
TRANSITION 5 13 0.100000 patru
TRANSITION 5 14 0.100000 trei
```

```
TRANSITION 5 15 0.100000 doi
TRANSITION 5 16 0.100000 unu
TRANSITION 5 17 0.100000 zero
TRANSITION 6 2 1.000000
TRANSITION 6 3 1.000000
TRANSITION 7 6 1.000000
TRANSITION 8 6 1.000000
TRANSITION 9 6 1.000000
TRANSITION 10 6 1.000000
TRANSITION 11 12 1.000000
TRANSITION 12 6 1.000000
TRANSITION 13 6 1.000000
TRANSITION 14 6 1.000000
TRANSITION 15 6 1.000000
TRANSITION 16 6 1.000000
TRANSITION 17 6 1.000000
FSG_END
```

- **Evaluarea sistemului de recunoaștere și interpretarea rezultatelor**

Având la dispoziție cele 3 componente fundamentale ale unui sistem de recunoaștere (modelul acustic, modelul de limbă și modelul fonetic) putem decoda clipurile audio de evaluare și comparare a transcrierii textuale rezultate cu transcrierea textuală de referință în vederea evaluării sistemului de recunoaștere de secvențe audio ce conțin cifre.

După rularea procesului de decodare vom obține rezultatele pentru fiecare configurație în parte. Raportul de evaluare afișează pe fiecare linie rezumatul rezultatelor pentru diverși vorbitori din baza de date de evaluare, iar pe coloane avem informații precum: id-ul vorbitorului, numărul propozițiilor și a cuvintelor de evaluare, procentul de cuvinte substituie/ șterse/ inserate, procentul de cuvinte recunoscute corect, rata de eroare la nivel de propoziție, precum și la nivel de cuvânt.

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. Următoarele cinci secțiuni prezintă statistici despre perechile de cuvinte confundate cel mai des (Confusion Pairs), cuvintele cele mai frecvent inserate (Insertions), șterse (Deletions), substituie (Substitutions) și cuvintele recunoscute fals cel mai des (Falsely Recognized).

În final, sunt prezentate în raport, frazele ipotetice aliniate de la frazele de referință corespunzătoare, precizând tipurile de erori pentru fiecare frază.

- **Optimizarea modelului acustic**

Modelul acustic creat de noi cu CMU Sphinx este construit cu trifoneme (fonem cărui i se precizează vecinii din stânga și dreapta), fiecare trifonem fiind implementat ca HMM cu trei stări emise. Fiecare astfel de model-stare poartă denumirea de senone și poate fi comună mai multor trifoneme în funcție de

asemănarea dintre ele. Modelul în stare inițială are numărul de senone egal cu 200, iar numărul de densități de probabilitate pentru fiecare GMM egal cu 8.

Pentru a obține acest model acustic cu 200 de senone și 8 densități Gaussiene per senone s-a trecut prin mai multe etape ce includ:

- modele fonetice independente de context;
- modele fonetice dependente de context cu o singură densitate Gaussiană per senone, apoi cu 2 și respectiv 4 densități;

Am configurat procesul de decodare pentru a folosi pe rând fiecare din cele 4 modele și am centralizat rezultatele astfel încât să avem o evaluare finală și comparativă a diverselor configurații.

Pentru decodarea clipurilor audio de evaluare folosind metode fonetice independente de context, modificăm fișierul de configurare al proiectului, specificând alte modele acustice. Am folosit modelele cu 200 de senone și am variat numărul de densități gaussiene per senone.

Pentru a verifica dacă numărul de 200 de senone este optim sau nu pentru sarcina noastră de recunoaștere a secvențelor audio ce conțin cifre, vom antrena un model acustic cu un număr de 100 de senone folosite la modelarea fonemelor ce formează cuvintele. Variem și acum numărul de densități Gaussiene și evaluăm 4 modele acustice.

Rezultate privind evaluarea sistemului dependent de vorbitor

Tabel 3.1 Ratele WER obținute pentru un sistem dependent de vorbitor

WER[%]		Densități Gaussinene per senone			
		1	2	4	8
SENONE	100	2.5	1.7	1.7	1.3
	200	6.3	7.5	15	22.8

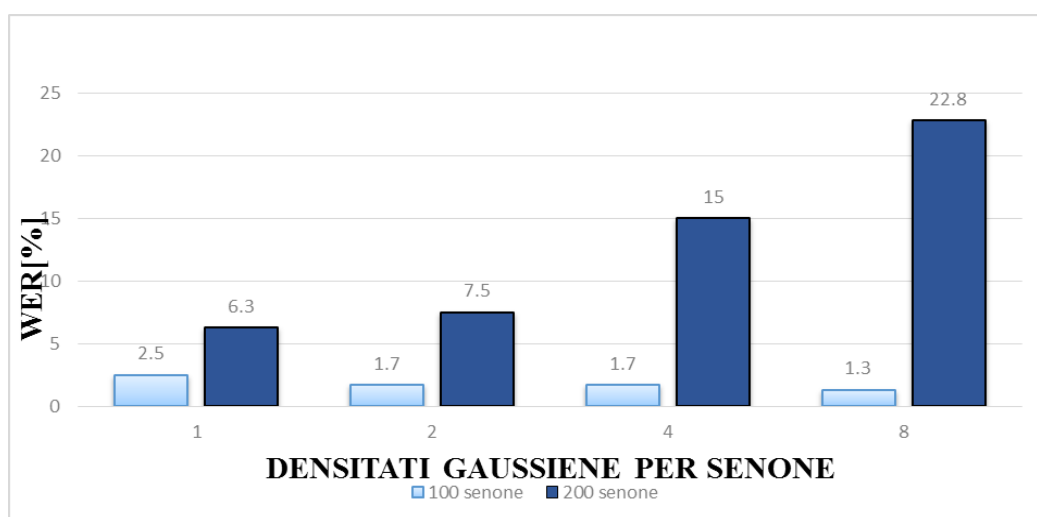


Figura 3.2 Reprezentarea grafică a ratelor WER obținute pentru un sistem dependent de vorbitor

Tabel 3.2 Ratele SER obținute pentru un sistem dependent de vorbitor

SER[%]		Densități Gaussinene per senone			
		1	2	4	8
SENONE	100	20	10	10	10
	200	30	40	60	75

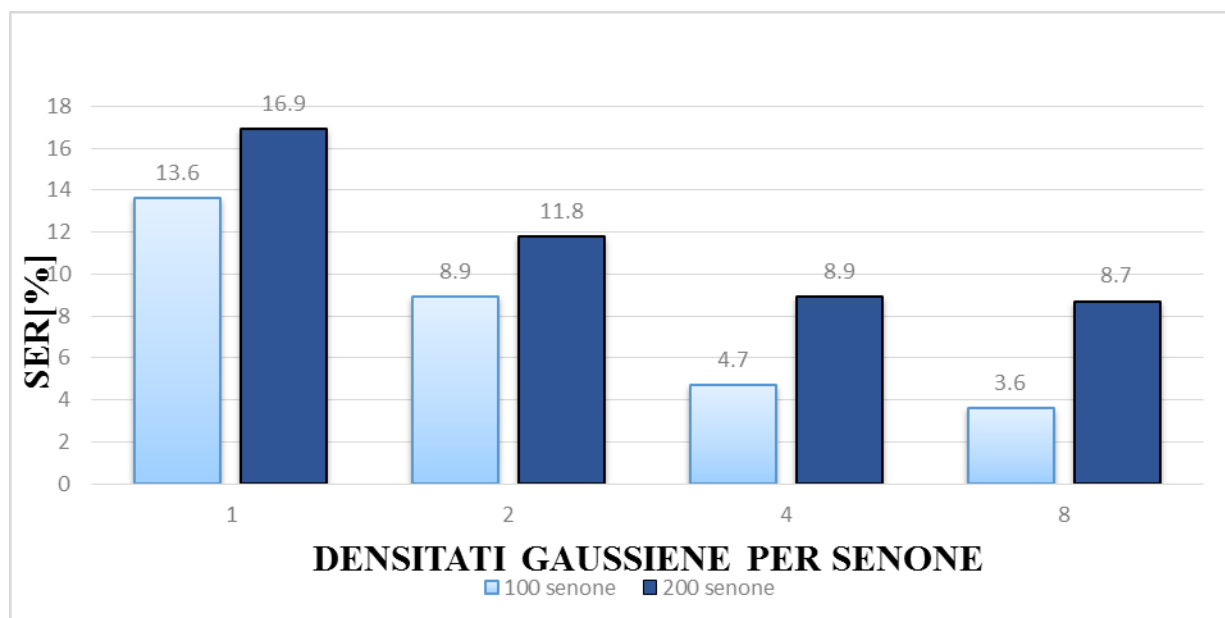


Figura 3.3 Reprezentarea grafică a ratelor SER obținute pentru un sistem dependent de vorbitor

Pentru un sistem dependent de vorbitor, dar independent de context am obținut o rata de eroare pe cuvânt de 3.3% și o rata de eroare pentru propoziții SER de 30%.

3.1.2 Construcția unui sistem de recunoaștere independent de vorbitor

Necesitatea unui astfel de sistem reiese din erorile foarte mari obținute în urma evaluării sistemului creat anterior pe fișierele de evaluare ale altor colegi. Rezultatele sunt de înțeles, deoarece modelul acustic din proiect a fost antrenat doar cu vocea noastră și cea mai bună recunoaștere o va face tot pentru noi. Am decis să creăm un alt proiect, independent de vorbitor în acest caz, folosind nu doar clipurile audio înregistrate de noi, ci toate clipurile din baza de date.

Pentru crearea acestuia ne-am folosit de :

- a) pentru antrenare: 80 de fișiere audio de la 69 de vorbitori (5520 fișiere audio);
- b) pentru testare: 20 de fișiere audio de la aceiași 69 de vorbitori (1380 fișiere audio);

Urmăm aceiași pași parcurși în cazul anterior, dependent de vorbitor, cu un număr de 200 și respectiv 100 de senone. Pentru fiecare situație folosim metode fonetice independente de context și dependente cu 1/2/4/8 densități Gaussiene per senone.

Rezultate privind evaluarea sistemului independent de vorbitor (vorbitori cunoscuți)

Tabel 3.3 Ratele WER obținute pentru un sistem independent de vorbitor

WER[%]		Densități Gaussinene per senone			
		1	2	4	8
SENONE	100	13.6	8.9	4.7	3.6
	200	16.9	11.8	8.9	8.7

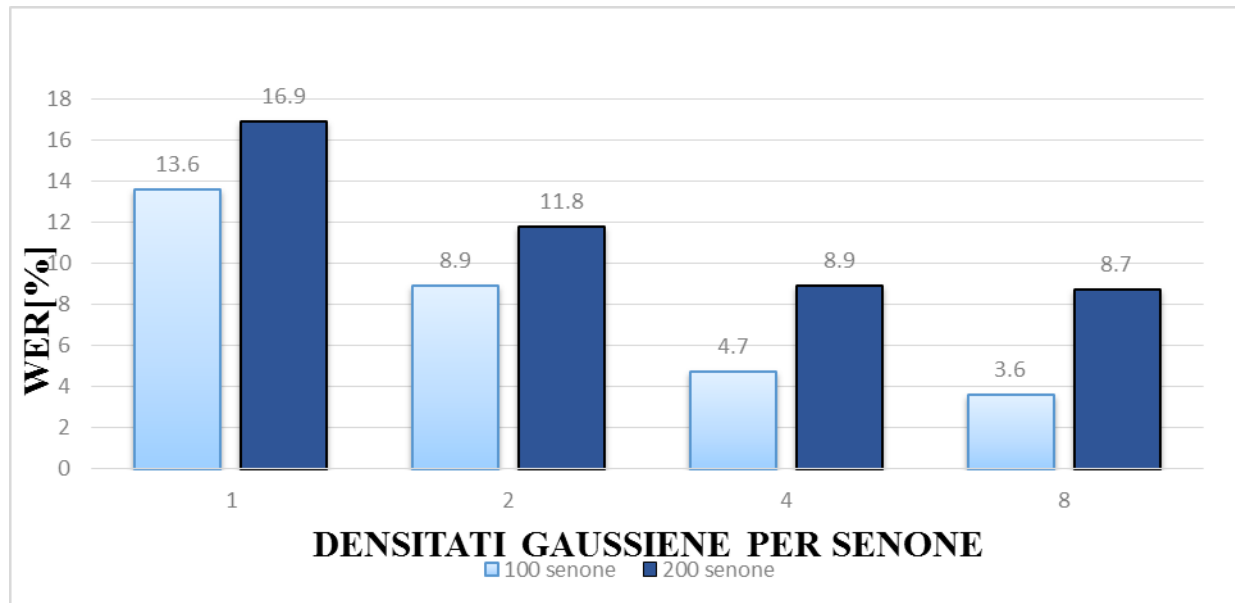


Figura 3.4 Reprezentarea grafică a ratelor WER obținute pentru un sistem independent de vorbitor

Tabel 3.4 Ratele SER obținute pentru un sistem independent de vorbitor

SER[%]		Densități Gaussinene per senone			
		1	2	4	8
SENONE	100	57.8	42.5	27	21.7
	200	67.2	56	48	46.4

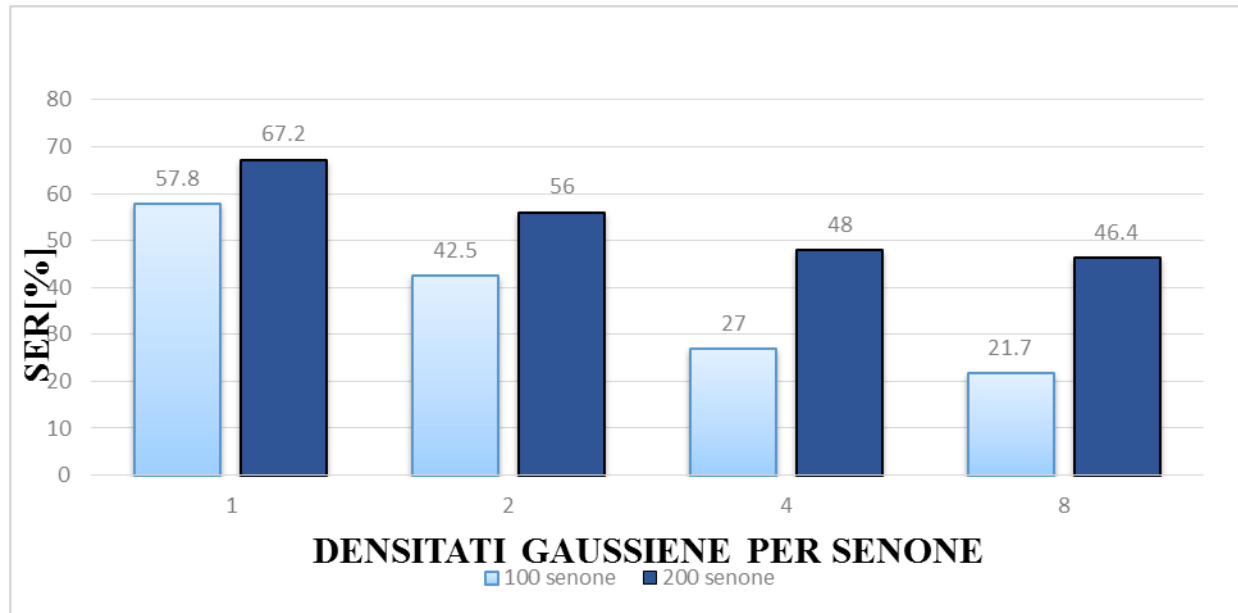


Figura 3.5 Reprezentarea grafică a ratelor SER obținute pentru un sistem independent de vorbitor

Comparând rezultatele obținute, observăm un trend de scădere al erorii pentru 100 de senone și decidem să extindem numărul densităților Gaussiene, trecând prin 16, 32, 64 și în final 128.

Tabel 3.5 Ratele WER obținute pentru un sistem independent de vorbitor, 100 de senone si un număr mai mare de densități gaussiene

WER[%]		Densități Gaussinene per senone			
		16	32	64	128
SENONE	100	3.5	2.9	2.4	2.4

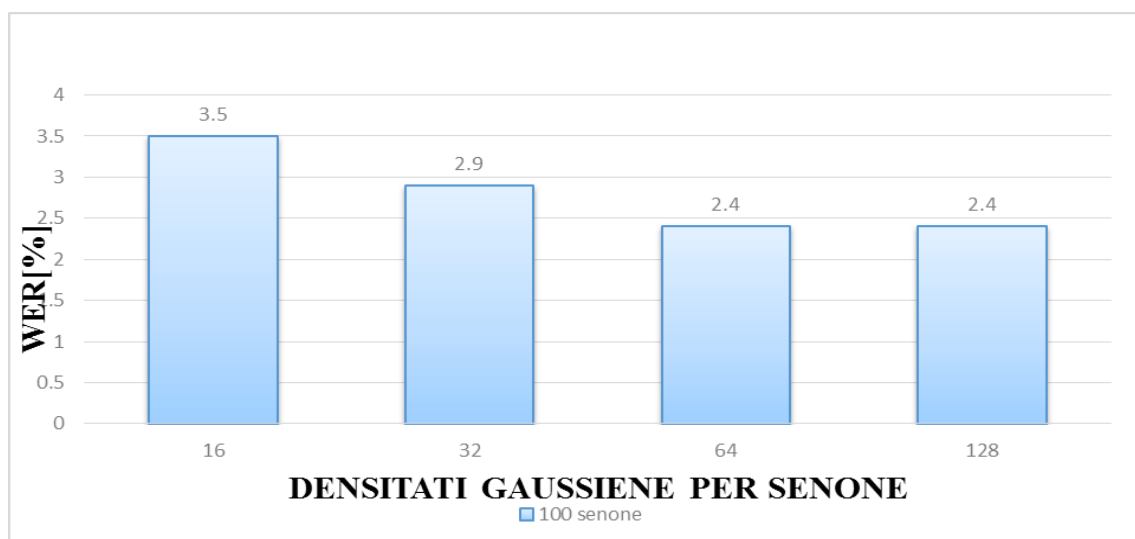


Figura 3.6 Reprezentarea grafică a ratelor WER obținute pentru un sistem independent de vorbitor, pentru 100 de senone

Tabel 3.6 Ratele SER obținute pentru un sistem independent de vorbitor, 100 de senone și un număr mai mare de densități gaussiene

SER[%]		Densități Gaussiene per senone			
		16	32	64	128
SENONE	100	21.8	17.9	17.2	16.7

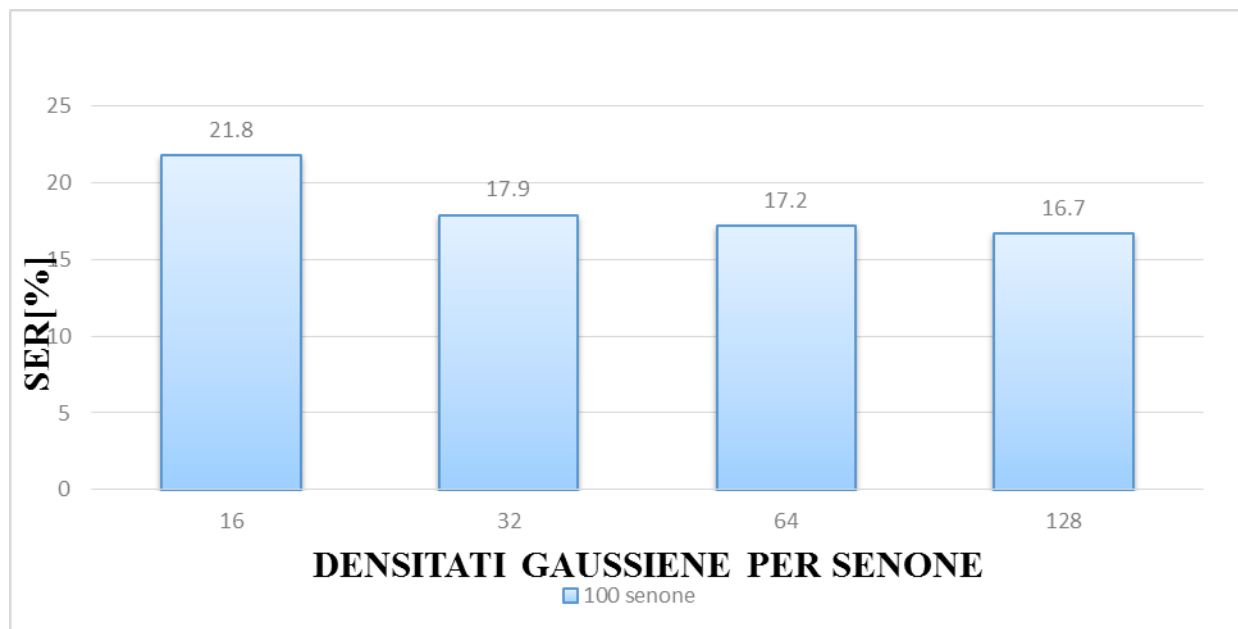


Figura 3.7 Reprezentarea grafică a ratelor SER obținute pentru un sistem independent de vorbitor, pentru 100 de senone

3.1.3 Concluzii

În urma evaluărilor, am selectat modelul optim ca fiind cel cu 128 densități Gaussiene și 100 de senone. Numărul mare al densităților Gaussiene oferă o aproximare mai bună a modelului.

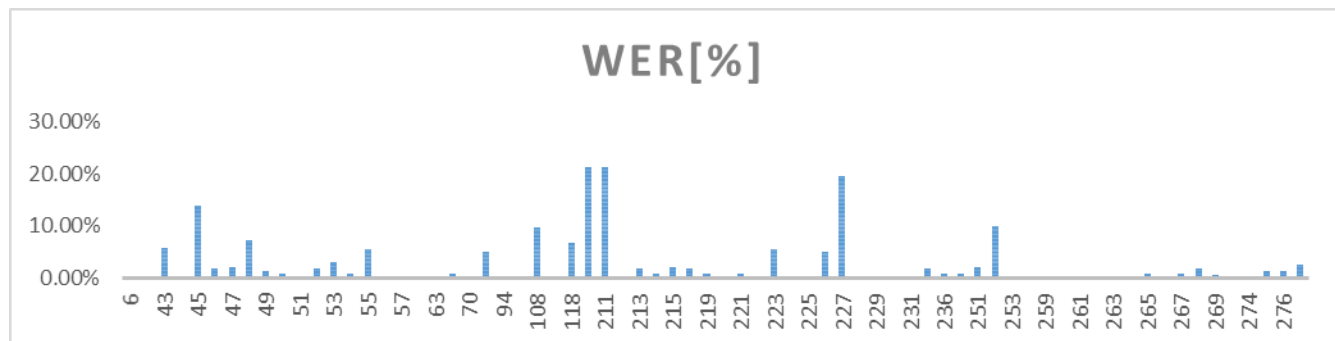


Figura 3.8 Reprezentarea grafică a ratelor WER pentru diverși vorbitori, obținute pentru un sistem independent de vorbitor

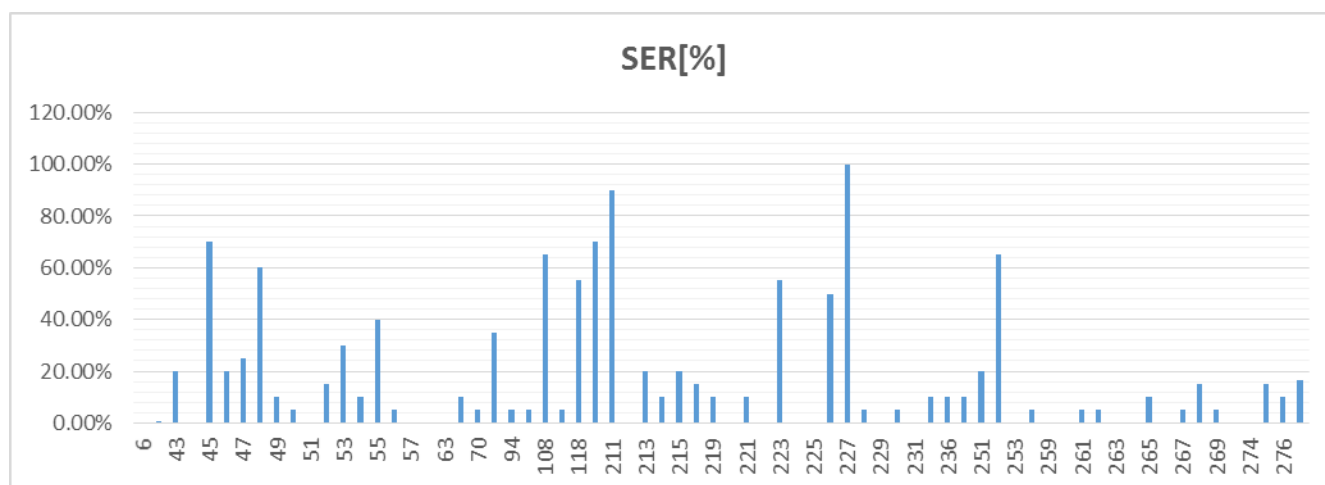


Figura 3.9 Reprezentarea grafică a ratelor SER pentru diverși vorbitori, obținute pentru un sistem independent de vorbitor

Pentru obținerea unui rezultat bun într-un sistem bazat pe modele Markov ascunse cu densități continue, este necesar să folosim îmbinarea dintre distribuțiile Gaussiene și modelele(fonetice) dependente de context.

3.2 MediaEval 2013: Sarcina de căutare a cuvintelor vorbite (*Spoken Web Search*)

Atribuțiile pentru MediaEval 2013 au constat în căutarea unui cuvânt rostit în interiorul unei baze de înregistrări, de asemenea rostite.

Procesul s-a bazat pe două seturi ce conțineau înregistrări multilingve: un set cu exemple de termeni căutați (unul sau mai multe exemple pentru un termen) și un set de documente audio unde vor avea loc căutarile și a implicat ca fiecare apariție a unui cuvânt cerut să fie identificată. Întrucât atât termenii cât și înregistrările audio pot conține diferite limbi, sistemele de căutare trebuie să fie independente de limbă.

Fișierele audio din baza de date a proiectului au fost înregistrate în diverse limbi (engleză, română, albaneză, cehă, bască, precum și 4 limbi africane) cu accente și condiții acustice diferite (unele realizate în cameră cu ajutorul microfoanelor, în timp ce altele pe stradă cu telefonul mobil). Dimensiunea bazei de înregistrări este în jur de 20 ore, iar numărul de cuvinte căutate este de aproximativ 400, ele fiind proporționale cu numărul de limbi utilizate.

3.2.1 Abordarea sarcinii

Abordarea acestei sarcini are la bază sistemul de recunoaștere automată a vorbirii pentru limba română, dezvoltat anterior.

Conform arhitecturii sistemului propusă de NIST [STD], avem două componente: o etapă de indexare, ce rulează transparent față de utilizator și una de căutare. În faza de indexare folosim sistemul ASR, ce preia o înregistrare și returnează un șir de foneme, reprezentând transcrierea fonetică a acesteia. Pasul

următor, acela de a găsi termenii dați se realizează prin căutarea transcrierii acestora în transcrierea fonetică a bazei de înregistrări. Scopul este acela de a găsi locurile în care asemănarea dintre șirurile de foneme (ale cuvintelor și bazei avute) depășește un anumit prag.

3.2.2 Sistemul de recunoaștere automată a vorbirii (ASR) pentru limba română

Modelul acustic pentru acest sistem a fost creat având la bază 64 ore de înregistrări aparținând unor diversi vorbitori. Cea mai bună performanță a fost o rată de eroare a cuvintelor (WER) de 18% pentru un discurs clar și continuu, în cazul în care a fost folosit un model de limbă antrenat cu 170 milioane de cuvinte, bazat pe N-gramme (rezultatul a fost preluat din experimentele anterioare realizate de echipa de cercetare).

Deoarece înregistrările sunt aleatorizate, nu a fost furnizată nici o informație privind limba vorbită în fisier. Pentru a reduce neconcordanța dintre aceasta bază de înregistrări și baza de antrenare a sistemului ASR, am filtrat înregistrările în limba română la 8kHz și am reantrenat modelul acustic.

Baza de înregistrări avută nu conține niciun tip de transcriere. Din acest motiv, este mai fezabil să folosim recunoașterea de cuvinte, însă folosim recunoașterea de foneme. Limba română are doar 26 foneme, deși numărul mare de limbi conținute în înregistrări triplează acest număr. Deoarece multe sisteme de recunoaștere automată a vorbirii folosesc clusterizarea fonemelor prin gruparea acestora în clase, în funcție de similitudinea acustică, putem avea mai puține clase decât foneme.

Pentru realizarea acestui lucru folosim modelele Markov ascunse (HMM) cu 3 stări per foneme, cu 4000 de senone, 8 densități Gaussiene, cu parametrii de voce MFCC. Modelele au fost antrenate să dividă spațiul trăsăturii vorbirii în clase de foneme, iar fonemele care nu aparțin limbii române vor fi clasificate în una din aceste clase. Dacă toate aceste instanțe, considerate ca neaparținând fonemelor din limba română ar fi recunoscute ca foneme românești, atât baza de înregistrări cât și cuvântul căutat ar avea același simbol. Rezultatele experimentale arată însă contrariul, adică faptul că multe foneme ocupă o regiune a vorbirii ce se suprapune doar parțial cu subspațiile fonemelor limbii române. Pentru baza de date a sistemului ASR din limba română avem o rată de eroare a fonemelor (PhER) de 36.8% (testată pe date în limba română).

3.2.3 Metodologia căutării

Algoritmul DTWSS (DTW String Search) folosește DTW pentru a alinia cuvântul căutat în conținut. Deoarece lungimea înregistrării o depășește pe cea a query-ului, căutarea trebuie să se realizeze doar în partea proporțională cu lungimea query-ului căutat numită “fereastra glisantă”, în care noi am împărțit înregistrarea. Aceste ferestre nu pot fi oricât de mari, deoarece fonemele unui termen căutat pot fi eventual găsite în interiorul ferestrei, deși este o potrivire falsă. Cuvântul este considerat găsit atunci când scorul algoritmului DTW depășește un anumit prag.

Pentru fiecare fereastră, alinierea la un scor s , complementar costului (egal cu valoarea PhER) :

$$S = (1 - \text{PhER}) \quad (1)$$

unde s este scorul de similitudine, atunci când detecția se bazează pe un prag determinat empiric. Procesul de segmentare a înregistrărilor în ferestre glisante și marcajele de timp fac procesul de

localizare a unui cuvânt posibil. Această metodă trebuie să ia în calcul și query-urile foarte scurte și implicit răspândirea potrivirilor DTW în context. Formula poate îmbunătăți căutarea, dacă este folosită sub formă:

$$s = (1 - PhER)(1 + \alpha \frac{L_Q - L_{Qm}}{L_{QM} - L_{Qm}})(1 + \beta \frac{L_W - L_S}{L_Q}) \quad (2)$$

unde, L_Q este lungimea cuvântului căutat, $L_{QM}=17$ și $L_{Qm}=4$ sunt lungimile maximă respectiv minimă a unui query găsit în setul de date, L_W este lungimea ferestrei glisante, α și β sunt parametrii configurabili, L_S este lungimea termenului găsit în înregistrare.

Penalizările în formula (1) sunt date de presupunerea că două cuvinte de lungimi diferite se pot potrivi într-un anumit conținut prin aceeași valoare PhER, iar query-ul cu lungimea cea mai mare este mai probabil să fie cel corect. Ponderea de raspândire a potrivirilor în cazul algoritmului DTW se explică prin faptul că, cu cât sunt mai compact aliniate fonemele, cu atât este mai mare probabilitatea ca ele să aparțină aceluiași termen. Introducerea parametrilor ajustabili (α și β) reprezintă o abordare nouă, valoarea lor optimă fiind obținută empiric.

3.2.4.Rezultate experimentale

Rezultatele obținute cu metoda primară, algoritmul DTW, pe setul de înregistrări de evaluare sunt arătate în Tabelul 1. Evaluarea se bazează pe valoarea ATWV care ia în calcul și locația cuvântului căutat. Rezultatele au fost obținute pentru o lungime a ferestrei glisante mai mare de 1.5 ori decât cea a cuvântului căutat. Cu cât este mai scurt query-ul, cu atât cresc șansele ca, cuvinte diferite, dar similare să obțină scoruri mai mari.

Tabel 3.7 Rezultatele algoritmului DTWSS

ATWV	$\alpha=0$	$\alpha=0.2$	$\alpha=0.4$	$\alpha=0.6$	$\alpha=0.8$	$\alpha=1$	$\alpha=1.2$
$\beta=0$	0	0.009	0.026	0.06	0.094	0.106	0.105
$\beta=0.2$	0	0.009	0.026	0.061	0.094	0.106	0.105
$\beta=0.4$	0	0.009	0.026	0.064	0.094	0.107	0.106
$\beta=0.6$	0	0.009	0.027	0.065	0.096	0.107	0.105
$\beta=0.8$	0	0.009	0.028	0.067	0.096	0.108	0.105
$\beta=1$	0	0.009	0.031	0.07	0.097	0.108	0.105
$\beta=1.2$	0	0.009	0.031	0.071	0.098	0.107	0.105

Experimentele au fost făcute pentru a determina valoarea optimă a parametrilor α și β , pentru metoda de căutare DTWSS, pornind de la formula 1($\alpha=0$; $\beta=0$) am determinat baza. Se observă că este mai bine să îi dăm parametrului α o valoare mai mare. În mod similar, cu cât este mai mare raspândirea potrivirii dată de DTW, cu atât este mai mică probabilitatea ca acela să fie cuvântul căutat. Există totuși o valoare optimă pentru α și β , explicate pe baza că, crescând α peste o anumită valoare, cuvintelor scurte li se alocă scoruri mari doar pentru că sunt de lungime mică, deși ele ar putea avea o valoare

PhER mare. Similar, o valoare mare a lui β înseamnă că potrivirile DTW vor primi un scor mai mare deși scorul lor real DTW (1-PhER) este mai mic.

Am ales următoarele combinații (α, β) : (1; 0.8), (1, 1).

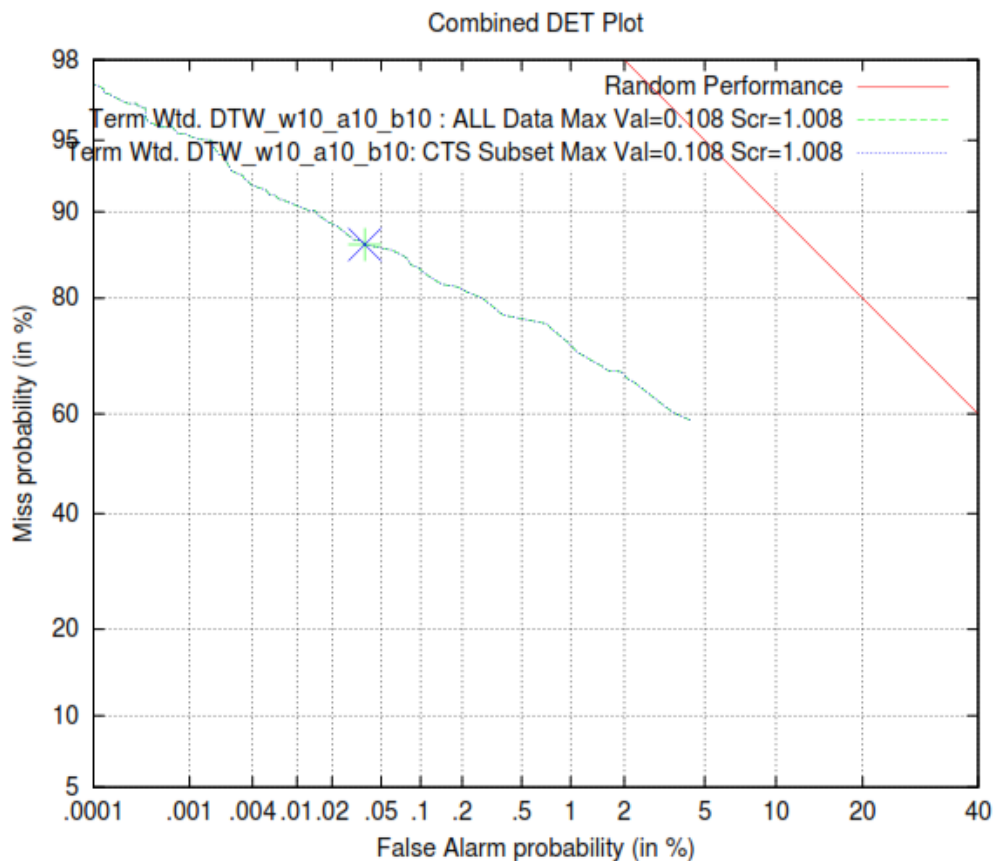


Figura 3.10 Rezultatele rulării principale cu baza de evaluare

Figura 3.9 arată curbele erorii de detectare a compromisului (DET) folosind metoda DTWSS, obținută pentru baza de înregistrări de evaluare.

3.2.5 Concluzii:

Sistemul de recunoaștere a termenilor (STD) s-a conturat ca un process în doi pași. Am folosit un ASR în limba română pentru recunoașterea de foneme, în cazul indexării bazei de înregistrări, în timp ce algoritmul DTW se utilizează în căutarea unui cuvânt dat. Rezultatele se îmbunătățesc dacă scorul deciziei este reglat în concordanță cu lungimea cuvântului și raspândirea de potriviri. Metoda aleasă este scalabilă pentru baze mari de date, deoarece am redus căutarea la o comparație de siruri de foneme.

CAPITOLUL 4

APLICAȚII DEMONSTRATIVE

4.1 Recunoaștere de cifre

4.1.1 Descriere

Aplicația demo de recunoaștere de cifre în limba română este o aplicație în Java ce preia cifrele rostite de utilizator, cu ajutorul unui microfon și le afișează textual în interfața grafică, cum este prezentat în figura 4.1.

Modelul acustic și modelul lingvistic folosite pentru aceasta aplicație sunt cele create anterior, care au dat cele mai bune rezultate de recunoaștere, în urma testării a diverse configurații.

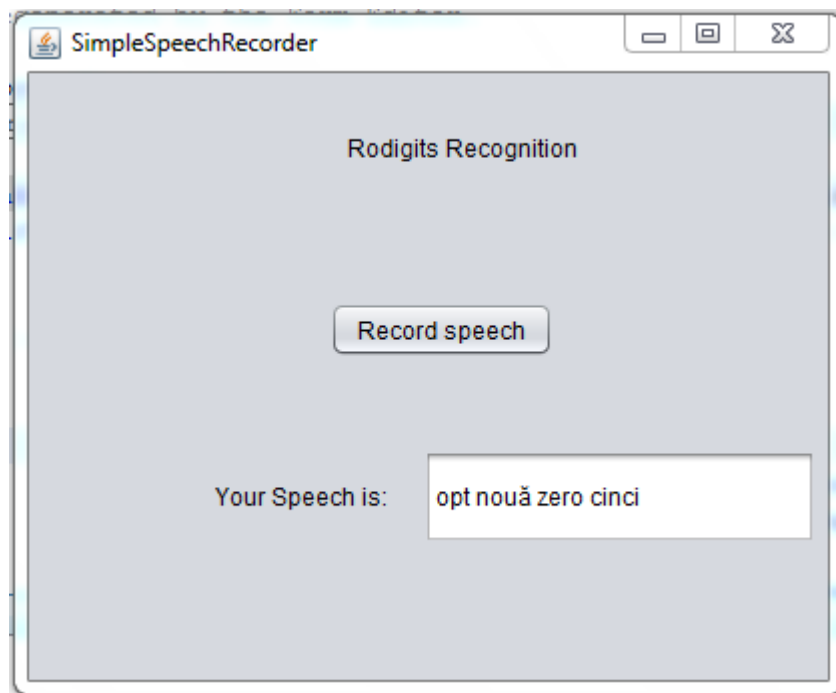


Figura 4.1 Interfața aplicației demonstrative

4.1.2 Interfața utilizatorului

Aplicația are o interfață grafică simplă construită cu ajutorul cadrului de lucru swing javax. Prin apăsarea butonului „Record speech” se începe înregistrarea utilizatorului, textul butonului se schimbă în „Stop recording”, iar în momentul în care apăsăm pe „Stop recording”, oprim înregistrarea și începem procesul de recunoaștere. Programul detectează automat sfârșitul mesajului rostit și face recunoașterea de cifre, dar este recomandată apăsarea butonului de stop pentru a diminua zgomotele din jur, ce pot influența negativ rezultatele recunoașterii. După ce recunoașterea este terminată, câmpul de text „Your Speech is:” va afișa șirul de cifre depistat. Se pot face mai multe recunoașteri consecutive prin reapăsarea butonului de înregistrare (Record speech), moment în care caseta de afișare se golește și procesul se reia.

4.1.3 Implementare

Pentru dezvoltarea proiectului am folosit unelte puse la dispoziție în mod gratuit de CMU Sphinx, ce oferă posibilitatea de a crea sisteme de recunoaștere a vorbirii continue, cu un vocabular larg, independente de vorbitor și mediul de dezvoltare Netbeans IDE. Sphinx-4 conține un set de clase ce folosesc Java Speech API (JSAPI) ca un motor în recunoașterea de vorbire. NetBeans reprezintă un suport integrat pentru dezvoltarea proiectelor ce includ o interfață grafică cu utilizatorul și oferă multe facilități.

Pentru început, am setat calea spre atributele necesare sarcinii noastre, specificând modelul acustic, dicționarul creat, fișierul cu gramatica și sursa de unde va proveni mesajul vorbit. Primele trei atribute au fost setate folosind obiectul de configurare, care a fost apoi pasat recunoscătorului de cifre.

Recunoașterea s-a bazat pe înregistrarea live a vocii utilizatorului cu ajutorul unui microfon și pentru a realiza acest lucru am folosit o clasă a bibliotecii Sphinx-4 numită LiveSpeechRecognizer împreună cu metodele sale. După ce microfonul este pornit, programul încearcă să recunoască ceea ce utilizatorul spune (cu metoda `recognizer.startRecognition(true)`). Cifrele recunoscute prin apelul metodei `recognizer.getResult()` sunt afișate în caseta de text corespunzătoare.

Pentru a putea realiza recunoașterea în mod transparent utilizatorului și a permite ca interfața să rămână activă, astfel încât să putem realiza și alte acțiuni în acest timp, se folosesc mai multe fire de execuție. `SwingWorker` este o clasă abstractă, ideală pentru a detașa unele sarcini de aplicație și doar să putem monitoriza statusul lor sau să le putem închide când dorim. Fiecare sarcină a unui fir de execuție este reprezentată de o instanță a lui `javax.swing.SwingWorker`.

Se pot furniza trei tipuri diferite de fire de execuție:

- Firul curent în cadrul căruia este apelată metoda `execute()`, ce programează un `SwingWorker` pentru a executa un fir și returnează răspunsul imediat;
- Firul lucrător (din engleza `Worker`) ce apelează `doInBackground()` și se ocupă de toate activitățile din fundal. Cu această metodă putem actualiza progresele aplicației.
- Event Dispatch Thread (EDT): toate activitățile Swing legate au loc aici și se folosesc metodele `process()` și `done()`.

Am utilizat metoda `doInBackground()`, în interiorul căruia pornim procesul de recunoaștere a cifrelor și recepționăm rezultatele. Pentru pornirea execuției codului apelăm metoda `execute()`. În momentul în care utilizatorul apasă pe butonul de înregistrare, `SwingWorker`-ul începe recunoașterea în fundal, apoi eliberează firul de execuție pentru a împiedica blocarea aplicației. Când este finalizată operația din fundal, automat este invocată metoda `done()` pentru a putem recupera rezultatul. Apelul metodei `done()` se poate face doar în firul de fundal numit event dispatching thread, singurul fir valid ce poate modifica starea componentelor de pe interfața grafică.

Atunci când calculul este terminat și rezultatul este disponibil, îl preluăm cu ajutorul metodei `getResult()` și astfel, putem să îl afișăm pe interfață.

4.1.4 Concluzii

Scopul principal al acestei aplicații demo a constat în a vedea cum sunt recunoscute secvențele de cifre rostite, indiferent de utilizator sau de modul său de pronunție și beneficiile pe care o astfel de aplicație le poate aduce tehnologiei. Există, bineînțeles, numeroase dezvoltări ce pot mări gradul de utilizare și complexitate al acestui sistem.

4.2 Detecția de gen

4.2.1 Descriere

Aplicația demo de detectare de gen este un program scris în limbajul Java ce preia o înregistrare audio a utilizatorului și returnează genul acestuia. O populație de indivizi poate fi partiționată în diverse subcategorii pe baza mai multor criterii, lucru ce interesează din multiple motive. Detecția automată a

genului prezintă mai multe aplicații: în contextul detecției vorbitorului, genul îmbunătățește performanțele prin limitarea spațiului de căutare, reprezintă o unealtă utilă în sistemele de analiză a semnalelor, în sistemele de indexare, cunoașterea genului vorbitorului este o indicație folosită în notare.

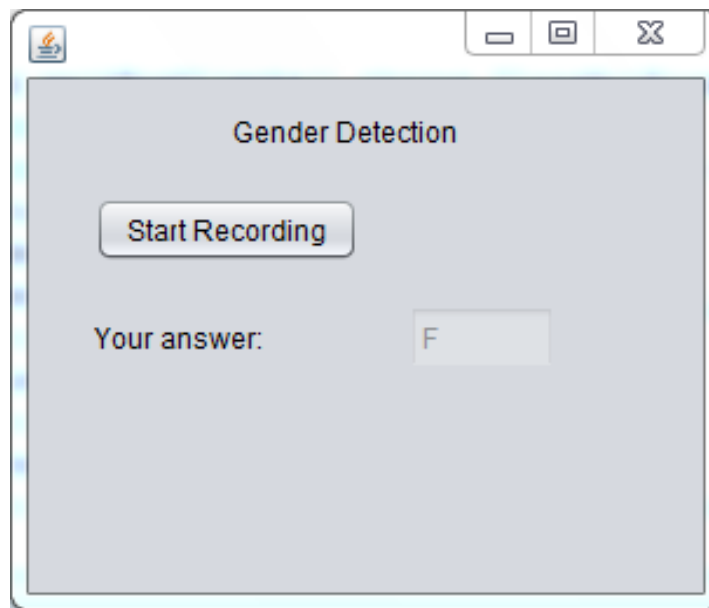


Figura 4.2 Interfața aplicației demonstrative

4.2.2 Interfața utilizatorului

Aplicația, în mod asemănător cu cea anterioară, are o interfață grafică simplă construită cu ajutorul cadrului de lucru swing javax, ce are la bază o conexiune de tip client-server. Apăsând butonul „Start Recording”, începe înregistrarea vocii utilizatorului, eticheta se schimbă în „Stop Recording”, ce ne permite să oprim înregistrarea în momentul în care ne dorim acest lucru și să vedem rezultatul returnat de aplicație. Butonul revine în stare activă de „Start Recording”, astfel încât putem să ne înregistrăm de mai multe ori consecutiv, fără a fi nevoie să pornim întreg procesul. Fișierul audio rezultat este transmis serverului, care îl prelucrează și returnează răspunsul detectat, ce va fi afișat în caseta de text corespunzătoare, precedată de eticheta „Your answer:”, figura 4.2.

4.2.3 Implementarea

Pentru dezvoltarea proiectului de recunoaștere a genului am utilizat o conexiune client-server, în cadrul căreia fiecare clasă are rolul său. Figura 4.3 prezintă o diagramă de secvență simplificată a aplicației:

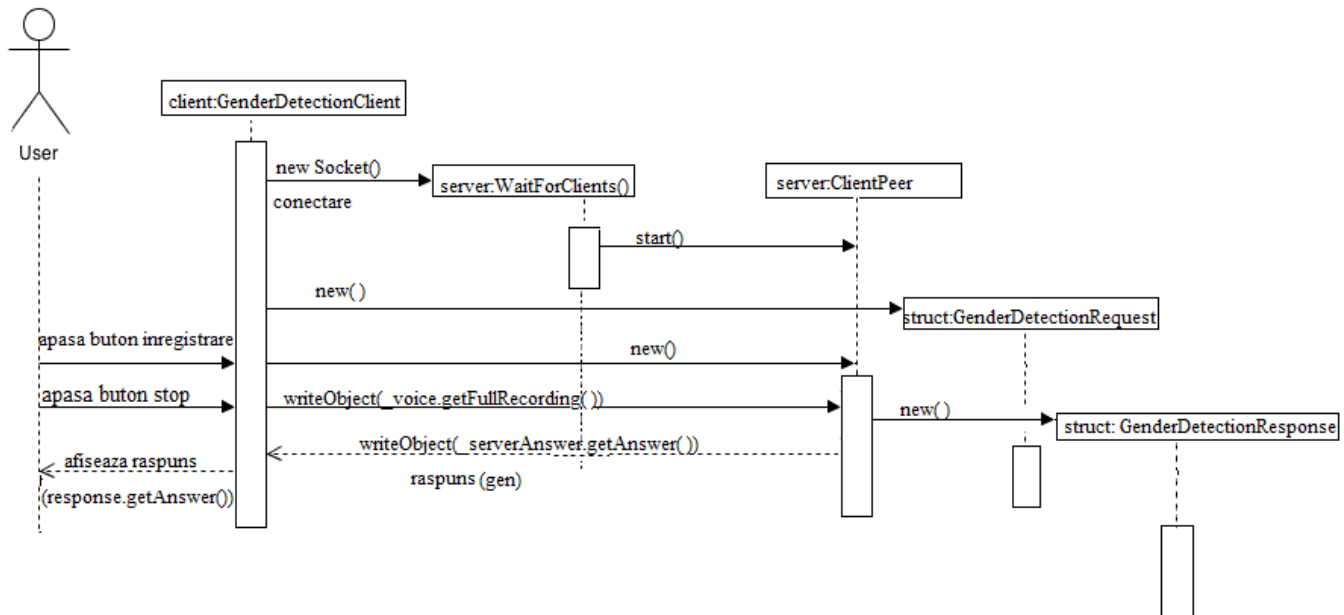


Figura 4.3 Diagrama de secvență a aplicației demonstrative

Principalele clase folosite pentru executarea sarcinii de recunoaștere a genului sunt prezentate în figura 4.4 :

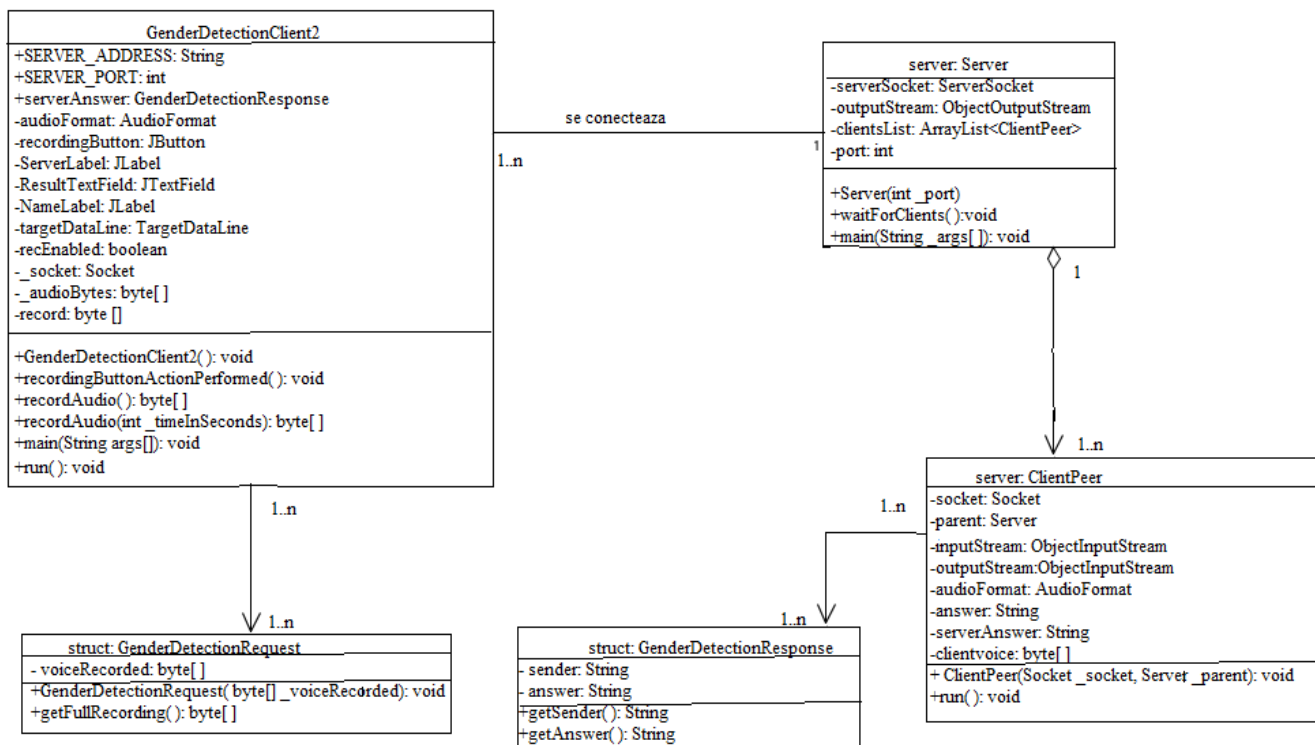


Figura 4.4 Diagrama UML (de clase) a aplicației demonstrative

Figura descrie structura internă a sistemului prin identificarea claselor, a atributelor și a relațiilor dintre clase. Fiecare clasă este reprezentată sub forma unui dreptunghi cu trei compartimente pentru: numele clasei, compartimentul din mijloc afișează atributele clasei, iar cel inferior operațiile efectuate.

4.2.3.1 Clientul

Clientul se conectează la server, înregistrează vocea sub forma unei înșirui de octeți, îi transmite serverului, recepționează și afișează răspunsul. Toate aceste acțiuni au loc în mod transparent utilizatorului, prin intermediul librăriei SwingWorker, descrisă în capitolul anterior 4.1.3, ce furnizează actualizarea interfeței în timp ce rulează, în spate, diverse sarcini. În acest timp, serverul așteaptă conexiunile clienților, recepționează fișierele audio trimise de aceștia, rulează partea de detectare și transmite răspunsul înapoi la client.

Pentru a realiza o aplicație client se va folosi clasa Socket ce permite stabilirea unei conexiuni cu un server existent precum și folosirea conexiunii pentru schimb bidirecțional de date, de tip TCP. Porturile și soclurile reprezintă mecanismul prin care se realizează legătura cu un server. Porturile reprezintă o poartă la care sunt primiți clienții pentru un anumit serviciu. Clientul trebuie să știe numele mașinii serverului pe care rulează precum și numărul portului pe care serverul ascultă. La nivelul său, socket-ul se realizează specificând adresa IP și portul serverului.

De exemplu la aceeași adresă se pot oferi diferite servicii, acestea oferindu-se la porturi diferite. Același calculator (cu o singură adresă IP) poate să ne ofere oricâte servicii dorește. Clienții care apelează la serviciile acestui calculator vor utiliza aceeași adresă, indiferent la care serviciu apelează, și toți clienții care doresc utilizarea unui anumit serviciu vor utiliza același port. Un soclu este de fapt un nod abstract de comunicație, reprezentând o interfață de nivel scăzut pentru comunicarea în rețea. Ele permit comunicarea între procese aflate pe același calculator sau pe calculatoare diferite din rețea. Java suportă trei tipuri de socluri. Clasa *Socket* utilizează un protocol orientat pe conexiune (TCP) și fluxuri de date (din engleză stream) pentru a trimite și a recepționa mesaje.

Un client ce dorește o comunicare cu serverul, trebuie să cunoscă trei lucruri:

- adresa server
- portul server utilizat pentru comunicare
- protocolul de comunicație utilizat de server

4.2.3.2 Serverul

În aplicația Server, pentru început, se crează un obiect de tip *ServerSocket* care ascultă pe un anumit port. Constructorul pentru *ServerSocket* aruncă o excepție dacă nu poate asculta pe portul respectiv. Dacă însă, se conectează pe portul său, acceptă conexiunea de la client și creează un fir de execuție nou, căruia îi încredințează socketul de legătură. Este necesară adăugarea referinței către fiecare client în parte într-o listă, pentru a putea transmite răspunsul către adresa corespunzătoare. Serverul continuă să asculte cererile celorlalți clienți care vor să se conecteze și îi acceptă cu ajutorul metodei *accept()*.

Pe fiecare fir de execuție (numit în aplicația curentă *ClientPeer*), în mod independent, este creat un flux de intrare pe care se primește șirul de octeți corespunzător vocii clientului și un flux de ieșire pentru a

returna răspunsul. Fluxul primit îl transformăm într-un fișier audio și pe baza acestuia este format un fișier .seg ce conține informații precum: numele înregistrării, canalul utilizat, începutul cadrului, durata în milisecunde, genul utilizatorului (în cazul în care deținem informații de referință ce pot fi folosite la verificare), banda (8kHz sau 16kHz) și mediu în care a avut loc înregistrarea. Extragerea coeficienților cepstrali MFCC se poate realiza prin metodele puse la dispoziție de LIUM sau de Sphinx-4 (sphinx_fe), când se precizează și caracteristice pe care le dorim. Următorul pas îl reprezintă alegerea modelelor de gen folosite.

LIUM estimează modelele de gen prin una din metodele: EM sau MAP pornind de la un set de înregistrări și etichetele respective. În cazul modelării folosind metoda MAP (din engleza maximum a posteriori), copiem modelul UBM-ul (din engleza universal background model), pentru fiecare vorbitor, obținut din mixarea GMM-urilor dependente de gen și banda. Pentru fiecare cluster, doar mediile sunt adaptate, conform formulei (2.3.3).

Decodarea va genera la ieșire un fișier cu rezultate ce conține informații precum: numele înregistrării, detecția sistemului, scorurile obținute pentru cele patru modele de gen și banda și referința (în cazul în care dispunem de ea pentru o verificare a performanțelor). După calcularea rezultatului, serverul este capabil să trimită răspunsul său clientului, iar acesta va fi afișat în câmpul de text rezervat. Utilizatorul are acum permisiunea de a își reînregistra vocea pentru o noua detecție.

4.2.4 Dezvoltări ulterioare

Există mai multe direcții în care aplicația poate fi îmbunătățită. Acestea se pot reflecta fie în interfață prin prezentarea mai multor informații sau prin permiterea unui control mai mare al utilizatorului, fie în codul pe care se bazează.

Unul din pașii doriți constă în modificarea și testarea caracteristicilor folosite în antrenarea modelelor vorbitorilor și în decodare, pentru a observa care combinație obține cele mai bune rezultate. Totodată, se pot investiga și cauzele erorilor de detecție obținute.

O alta îmbunătățire adusă aplicației ar fi detecția genului în diverse condiții de fundal ale înregistrării precum: zgomot, reflexii ale vocii, ce pot îngreuna decizia sistemului.

4.2.5 Concluzii

Scopul principal al acestei aplicații a fost acela de a recunoaște genul vorbitorului și de a observa cum arată o astfel de aplicație, ce poate fi folosită în diverse domenii reale. Un exemplu ar consta în adaptarea muzicii de așteptare la telefon sau furnizarea unor informații statistice pe baza distribuțiilor de gen. Am folosit sistemul LIUM, care este de nouă generație în domeniul vorbirii, bazat pe segmentare urmată de o clasificare. Scorurile obținute pentru cele patru categorii de combinații gen-bandă au prezentat diferențe foarte mici, lucru ce a condus la unele detecții de gen eronate.

CAPITOLUL 5

CONCLUZII

5.1 Concluzii generale

Scopul principal al acestei lucrări a fost acela de a prezenta diversele forme sub care sistemele de recunoaștere a vorbirii pot influența. Această sarcină a fost împărțită în trei sub-pați: un sistem de recunoaștere a cifrele în limba română, o căutare a termenilor roștiți într-un context și un sistem de detecție a genului vorbitorului, fiecare cu subdiviziunile specifice. Prima parte a fost îndeplinită cu succes, pentru cea de-a doua am obținut rezultate pe baza experimentelor, în timp ce ultima încă mai prezintă o rată de detecție greșită destul de semnificativă.

Lucrarea prezintă pașii urmați pentru îndeplinirea fiecărei atribuții. În capitolul 1 a fost descrisă motivația unor astfel de aplicații, în timp ce Capitolul 2 a prezentat detaliat cele trei părți componente ale lucrării.

Secțiunea 2.1 s-a axat pe prezentarea sistemele de recunoaștere automată a vorbirii, împreună cu uneltele puse la dispoziția utilizatorilor pentru prelucrarea datelor și obținerea rezultatelor.

Subcapitolul 2.2 a detaliat sarcina Spoken Web Search, de găsim a unui termen, împreună cu planul de evaluare NIST.

Ultime secțiune 2.3 a dezvoltat problema recunoașterii genului, parcurgând pașii necesari pentru rezolvarea acestuia și modalități de evaluare a performanțelor.

Capitolul 3 s-a focusat pe rezultatele experimentale obținute în competiția de recunoaștere a termenilor roștiți MediaEval 2013 (Spoken Web Search) și în detecția de gen. Tot în cadrul acestui capitol, am prezentat metodologia algoritmului de căutare DTW pentru MediaEval și rezultatele optime obținute prin varierea parametrilor α și β în algoritmul DTWSS (din engleză Dynamic Time Warping String Search). De asemenea, experimentele pentru recunoașterea de cifre rostite în limba română m-au ajutat în găsirea celor mai bune modele acustice și de limbă, ce au dat cele mai bune rate de eroare.

Capitolul 4 a prezentat două aplicații demonstrative cu interfața grafică pentru utilizatori. Prima a fost dezvoltată pentru îndeplinirea sarcinii de recunoaștere a cifrelor și a presupus folosirea setului de instrumente CMU Sphinx, în timp ce a doua prezintă recunoașterea genului și detaliază o arhitectură de timp client-server pe care s-a bazat.

Aceste două aplicații pot fi folosite pentru a demonstra performanțele unor sisteme de recunoaștere.

5.2 Contribuții personale

Atribuțiile personale ale autorului ce au dus la îndeplinirea obiectivelor acestei lucrări pot fi sumarizate astfel:

1) Proiectul recunoașterii de cifre

- a) Crearea unui sistem de recunoaștere a cifrelor dependent de vorbitor.
- b) Crearea unui sistem de recunoaștere a cifrelor independent de vorbitor.
- c) Efectuarea experimentelor de optimizare a configurației unui sistem de recunoaștere automată a cifrelor din limba română.
- d) Realizarea unei aplicații cu interfața grafică pentru recunoașterea cifrelor captate într-o înregistrare cu ajutorul microfonului.

2) Proiectul de găsim a termenilor vorbiți (MediaEval 2013)

- a) Crearea unor fișiere ce conțin numele înregistrărilor în care se efectuează căutările și respectiv a înregistrărilor cu termenii căutați;
- b) Efectuarea experimentelor de optimizarea a configurației unui sistem de recunoaștere automată a vorbirii pentru limba română;
- c) maparea fonemelor din limbile în care au fost realizate înregistrările cu foneme din limba română;
- d) adaptarea modelului acustic în limba română la datele acustice furnizate;
- e) crearea modelelor de limbă trigrame;
- f) realizarea experimentelor pentru optimizarea parametrilor din cadrul algoritmului DTWSS;
- g) eliminarea termenilor formați din mai puțin de patru caractere, pentru a obține rezultate
- h) Rularea scriptului de evaluare furnizat de NIST STD

i) Reprezentarea curbelor DET pentru toate combinațiile configurate ale parametrilor algoritmului de căutare DTWSS.

3) Proiect de recunoaștere a genului

- a) Crearea unor fișiere folosite în procesul de antrenare și evaluare a metodelor de detectare a genului.
- b) Antrenarea cu ajutorul unei baze de înregistrări a unor modele pentru genul masculin și feminin.
- c) Realizarea unor experimente având ca scop comparația diverselor modele de gen și metode de decodare.
- d) Crearea unei aplicații cu interfața grafică de tip client-server pentru detectarea genului utilizatorului.
- e) Dezvoltarea sistemului se bazează pe modelele de gen furnizate de LIUM.

BIBLIOGRAFIE

- [ASR] <http://speed.pub.ro/speed3/wp-content/uploads/2013/04/Indrumar-de-proiect-PCDTV-v4.pdf> (Accesat la 12.05.2014)
- [B01] Buzo, A., "Automatic Speech Recognition over Mobile Communication Networks," Teza de doctorat, Universitatea Politehnica din Bucuresti, România, 2011, <http://speed.pub.ro/academic/research-and-development-project-in-spoken-language-technology>.
- [C01] S. S. Chen and P. S. Gopalakrishnan, Clustering via the Bayesian information criterion with applications in speech recognition, in Proc. IEEE Int. Conf. Acoustics Speech, and Signal Processing, vol. 2, pp. 645 - 648, Seattle, USA, May 1998.
- [C02] Cucu, H., "Towards a speaker-independent, large-vocabulary continuous speech recognition system for Romanian," Teza de doctorat, Universitatea Politehnica din Bucuresti, România, 2011.
- [D01] N. Deshmukh, A. Ganapathiraju, J. Hamaker, J. Picone, and M. Ordowski, "A public domain speech-to-text system," in *Proceedings of the 6th European Conference on Speech Communication and Technology*, vol. 5, Budapest, Hungary, Sept.1999, pp. 2127–2130.
- [Diarizare] <http://archives.limsi.fr/RS2005/chm/tlp/tlp1/index.html> (Accesat la 21.06.2014)
- [Gender] <http://www.itl.nist.gov/iad/mig/publications/proceedings/darpa98/html/bn30/bn30.html> (Accesat la 15.04.2014)
- [H01] Huang, X., Acero, A., Wuen-Hon, H., *Spoken Language Processing – A Guide to Theory, Algorithm, and System Development*, Prentice Hall, 2001.
- [H02] Hwang, M.Y., X.D. Huang, "Acoustic Classification of Phonetic Hidden Markov Models," *Eurospeech 1991*, 1991.
- [J01] Daniel Jurafsky, James H. Martin, *Speech and Language Processing- An Introduction to Natural Language Processing, Computational Linguistics and Speech Recognition*, 1999.
- [JSGF] Java Speech Grammar Format, <http://java.sun.com/products/java-media/speech/forDevelopers/JSGF/index.html>. (Accesat la data de 01.06.2014)
- [L01] X. X. Li, Y. Zhao, X. Pi, L. H. Liang, and A. V. Nefian, "Audio-visual continuous speech recognition using a coupled hidden Markov model," in *Proceedings of the 7th International Conference on Spoken Language Processing*, Denver, CO, Sept. 2002, pp. 213–216.
- [P01] Poritz, A., "Hidden Markov models: a guided tour," *ICASSP 1988*, pp. 7–13.
- [R01] 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.
- [S01] George Saon, Jen-Tzung, „Large-Vocabulary Continuous Speech Recognition System” în *IEEE Signal Processing Magazine*, nr.6/2012, pp.19-25.

- [Speech] http://en.wikipedia.org/wiki/Speech_recognition (Accesat la 15.06.2014)
- [Sphinx] Toolkitul CMU Sphinx si tutorialele asociate, <http://cmusphinx.sourceforge.net/wiki/> (Accesat la 21.06.2014)
- [STD] <http://www.multimediaeval.org/mediaeval2013> (Accesat la 20.06.2014)
- [W01] Willie Walker, Paul Lamere, Philip Kwok, Bhiksha Raj, Rita Singh, Evandro Gouvea, Peter Wolf, and Joe Woelfel, „Sphinx-4: A Flexible Open Source Framework for Speech Recognition”, November 2004
- [Y01] S. Young, “The HTK hidden Markov model toolkit: Design and philosophy,” Cambridge University Engineering Department, UK, Tech. Rep. CUED/F-INFENG/TR152, Sept. 1994.