

APLICAȚII SOFTWARE ÎN DOMENIUL TEHNOLOGIEI VORBIRII

PROIECT DE DIPLOMĂ

Prezentat ca cerință parțială pentru obținerea titlului de *Inginer*
în domeniul *Electronică, Telecomunicații și Tehnologia Informației*
Programul de studii *Calculatoare și Tehnologia Informației*

Conducători științifici

Ș.I. Dr. Ing. Horia CUCU

Prof. Dr. Ing. Corneliu BURILEANU

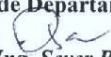
Absolvent

Alexandru Lucian GEORGESCU

București
2016

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

Aprobat Director de Departament :


Prof. Dr. Ing. Sever PAȘCA

TEMA PROIECTULUI DE DIPLOMĂ
a studentului Georgescu E. Alexandru Lucian, grupa 443A

1. Titlul temei: **Aplicații software în domeniul tehnologiei vorbirii**

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

Serviciu web de autentificare prin voce

- Aplicație client: interfață grafică ce oferă funcționalități de autentificare și înregistrare a unor utilizatori noi în sistem. Aceasta va fi implementată folosind tehnologii Web: HTML, CSS, Javascript, Web Audio API.
- Baza de date: stochează informații despre utilizatorii aplicației.
- Server de autentificare: serviciu de tip REST utilizând tehnologie Java ce realizează agregarea și procesarea datelor, interacționând cu aplicația client, cu baza de date și cu serviciile de recunoaștere de vorbire și vorbitor (ultimile două dezvoltate în cadrul unui alt proiect).

Aplicație pentru detecția cuvintelor cheie

- Crearea unui sistem de recunoaștere automată a vorbirii folosind platforma de dezvoltare Kaldi și evaluarea performanței sistemului și a resurselor de calcul necesare
- Implementarea unei aplicații client-server de detecție a cuvintelor cheie folosind sistemul de recunoaștere a vorbirii menționat mai sus. Aplicația este scrisă în limbajul de programare C și folosește biblioteca de funcții Gstreamer.

3. Proiectul se bazează pe cunoștințe dobândite în principal la următoarele 3-4 discipline: Programare Obiect-Orientată, Tehnologii de Programare în Internet, Proiectarea bazelor de date, Interfețe om-mașină.

4. Proprietatea intelectuală asupra proiectului aparține: Laborator Cercetare Speed, student Georgescu E. Alexandru Lucian

5. Locul de desfășurare a activității: UPB

6. Realizarea practică rămâne în proprietatea: Laborator Cercetare Speed, student Georgescu E. Alexandru Lucian

7. Data eliberării temei: 01.10.2015

CONDUCĂTOR LUCRARE:
Ș.l. Dr. Ing. Horia CUCU



STUDENT:
Alexandru Lucian GEORGESCU



DECLARAȚIE DE ONESTITATE ACADEMICĂ

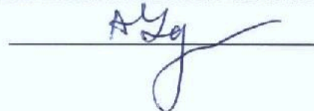
Prin prezenta declar că lucrarea cu titlul „*Aplicații software în domeniul tehnologiei vorbirii*”, prezentată în cadrul Facultății de Electronică, Telecomunicații și Tehnologia Informației a Universității „Politehnica” din București ca cerință parțială pentru obținerea titlului de Inginer în domeniul Electronică, Telecomunicații și Tehnologia Informației, programul de studii Calculatoare și Tehnologia Informației 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 resursele 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, 08.09.2016

Absolvent *Alexandru Lucian* GEORGESCU



CUPRINS

Cuprins	7
Listă de Figuri	9
Listă de Tabele	11
Listă de Acronime	13
INTRODUCERE	15
Motivația lucrării	15
Obiective	16
Structura lucrării	16
CAPITOLUL 1 Domeniul tehnologiei vorbirii	19
1.1 Structura unui sistem de recunoaștere automată a vorbirii	19
1.1.1 Producerea semnalului vocal	20
1.1.2 Percepția semnalului vocal	21
1.1.3 Parametri acustici	22
1.1.4 Modele acustice	23
1.1.5 Modele de limbă	25
1.1.6 Modelul fonetic	26
1.1.7 Metrice de evaluare	26
1.1.8 Utilitare pentru recunoașterea automată a vorbirii	26
1.2 Structura unui sistem de detecție a cuvintelor cheie	28
1.2.1 Căutarea prin lattice	28
1.2.2 Căutarea prin transcriere text	29
1.2.3 Metrice de evaluare	29
CAPITOLUL 2 Construcția unor sisteme de recunoaștere automată a vorbirii	33
2.1 Crearea sistemului "Rodigits"	33
2.1.1 Baza de date audio "Rodigits"	34
2.1.2 Antrenarea modelelor în vederea antrenării folosind utilitarul CMU Sphinx	34
2.1.3 Experimente și rezultate	35
2.1.4 Prepararea modelelor în vederea antrenării folosind utilitarul Kaldi	41
2.1.5 Experimente și rezultate. Compararea utilităților	43
2.2 Crearea sistemului "Molina"	44
2.2.1 Baza de date audio "Molina"	45

2.2.2	Antrenarea modelelor	45
2.2.3	Experimente și rezultate	50
CAPITOLUL 3 Construcția sistemului de detecție a cuvintelor cheie		55
3.1	Analiza asupra modelelor de limbă folosite pentru detecția cuvintelor cheie	55
3.1.1	Modelul de limba 3-gram	55
3.1.2	Alternativă la modelul de limba 3-gram.....	57
3.2	Crearea referinței.....	59
3.3	Experimente și rezultate.....	59
CAPITOLUL 4 Serviciul web de autentificare prin voce		63
4.1	Descriere și arhitectură.....	63
4.2	Funcționalitate din perspectiva utilizatorului.....	65
4.3	Scenarii de utilizare.....	68
4.4	Tehnologii utilizate	69
4.5	Implementarea aplicației client	73
4.6	Implementarea bazei de date.....	76
4.7	Implementarea serviciului de autentificare	78
4.8	Concluzii și dezvoltări ulterioare	83
CAPITOLUL 5 Aplicația de detecție a cuvintelor cheie		85
5.1	Descriere și arhitectură.....	86
5.2	Funcționalitate din perspectiva utilizatorului.....	86
5.3	Tehnologia Gstreamer.....	87
5.4	Descrierea algoritmului.....	88
5.5	Verificarea unui sistem interactiv cu răspuns vocal.....	90
5.6	Implementare	91
5.7	Concluzii și dezvoltări ulterioare	94
Concluzii finale		95
Contribuții personale		97
Bibliografie		99

LISTĂ DE FIGURI

Figura 1.1 Arhitectura unui sistem ASR. Sursa: [1]	20
Figura 1.2 Modelul sursă-filtru al producerii vorbirii. Sursa:[3]	21
Figura 1.3 Scara Mel.....	22
Figura 1.4 Metodă de extragere a MFCC. Sursa:[4].....	22
Figura 1.5 Reprezentarea în timp și frecvență a semnalului. Sursa:[5]	23
Figura 1.6 Bancul de filter Mel. Sursa:[6]	23
Figura 1.7 Structura unui HMM	24
Figura 1.8 Mixtură de gaussiene. Sursa:[7]	25
Figura 1.9 Arhitectură CMU Sphinx. Sursa:[8].....	27
Figura 1.10 Arhitectură Kaldi. Sursa:[9]	27
Figura 1.11 Exemplu de lattice	28
Figura 1.12 Schemă de evaluare a sistemului KWS. Sursa:[10]	30
Figura 2.1 Modelul de limbă de tip gramatică Rodigits (CMU Sphinx). Sursa:[11].....	35
Figura 2.2 Raport WER - Densități gaussiene	38
Figura 2.3 Raport SER - Densități Gaussiene.....	38
Figura 2.4 Evaluarea vorbitorilor necunoscuți (WER)	39
Figura 2.5 Evaluarea vorbitorilor cunoscuți 1 (WER).....	39
Figura 2.6 Evaluarea vorbitorilor cunoscuți 2 (WER).....	40
Figura 2.7 Evaluarea vorbitorilor cunoscuți 3 (WER).....	40
Figura 2.8 Modelul de limbă 1-gram Rodigits (Kaldi)	42
Figura 2.9 Modelul de limbă de tip gramatică Rodigits (Kaldi).....	43
Figura 2.10 Graficul corespunzător modelului de limbă de tip gramatică Molina.....	49
Figura 3.1 Evaluarea sistemului de detecție a cuvintelor cheie	60
Figura 4.1 Arhitectura serviciului web de autentificare prin voce.....	64
Figura 4.2 Formularul de înrolare	66
Figura 4.3 Formularul de autentificare	67
Figura 4.4 Pagina contului personal.....	67
Figura 4.5 Model de Context Audio (Web Audio API). Sursa:[15].....	71
Figura 4.6 Structura aplicației client.....	74
Figura 4.7 Tabelele bazei de date	76

Figura 4.8 Lista procedurilor stocate	77
Figura 4.9 Diagramă UML a claselor	79
Figura 5.1 Arhitectura aplicației de detecție a cuvintelor cheie	86
Figura 5.2 Rularea aplicației de detecție a cuvintelor cheie	87
Figura 5.3 Pipeline-ul GStreamer	88
Figura 5.4 Diagramă UML de stare a aplicației de detecție a cuvintelor cheie.....	89

LISTĂ DE TABELE

Tabelul 2.1 Experiment 1. Rezultatele modelului independent de context	36
Tabelul 2.2 Experiment 1. Rezultatele modelului dependent de context (WER).....	36
Tabelul 2.3 Experiment 1. Rezultatele modelului dependent de context (SER)	36
Tabelul 2.4 Experiment 2. Rezultatele modelului independent de context	36
Tabelul 2.5 Experiment 2. Rezultatele modelului dependent de context (WER).....	36
Tabelul 2.6 Experiment 1. Rezultatele modelului dependent de context (SER)	36
Tabelul 2.7 Experiment 3. Rezultatele modelului independent de context	37
Tabelul 2.8 Experiment 3. Rezultatele modelului dependent de context 1 (WER).....	37
Tabelul 2.9 Experiment 3. Rezultatele modelului dependent de context 1 (SER)	37
Tabelul 2.10 Experiment 3. Rezultatele modelului dependent de context 2 (WER).....	37
Tabelul 2.11 Experiment 3. Rezultatele modelului dependent de context 2 (SER)	37
Tabelul 2.12 Variații ale densităților gaussiene și senonelor în raport cu WER	38
Tabelul 2.13 Experiment 1. Rezultate Rodigits (Kaldi)	44
Tabelul 2.14 Experiment 2. Rezultate Rodigits (CMU Sphinx).....	44
Tabelul 2.15 Caracteristicile bazei de date TED-LIUM.....	46
Tabelul 2.16 Rezultatele antrenării bazei de date TED-LIUM.....	46
Tabelul 2.17 Experiment 1. Rezultate Molina (LM1-v1)	50
Tabelul 2.18 Evaluarea consumului de resurse (LM1-v2).....	52
Tabelul 2.19 Evaluarea consumului de resurse (LM2).....	53
Tabelul 3.1 Rezultatele evaluării sistemului în limba română.....	56
Tabelul 3.2 Consumul de memorie în cazul modelului de limbă 3-gram.....	57
Tabelul 3.3 Consumul de memorie în cazul modelului de limbă alternativ	58

LISTĂ DE ACRONIME

ASR - Automatic Speech Recognition
KWS - Keyword Spotting
IVR - Interactive Voice Response System
MFCC - Mel Frequency Cepstral Coefficients
PLP - Perceptual Linear Prediction
HMM - Hidden Markov Models
GMM - Gaussian Mixture Model
FSG - Finite State Grammar
SER - Sentence Error Rate
WER - Word Error Rate
JSGF - Java Speech Grammar Format
FST - Finite State Transducer
LDA - Linear Discriminant Analysis
MLLT - Maximum Likelihood Linear Transform
SAT - Speaker Adaptive Training
MMI - Maximum Mutual Information
CPU - Central Processing Unit
ECF - Experiment Control Files
RTTM - Rich Transcription Time Marked
TWV - Term Weighted Value
DET - Detection Error Tradeoff
HTML - HyperText Markup Language
CSS - Cascade Style Sheet
API - Application Programming Interface
REST - Representational State Transfer
HTTP - HyperText Transfer Protocol
XML - Extensible Markup Language
RTF - Real Time Factor

INTRODUCERE

MOTIVAȚIA LUCRĂRII

Comunicarea verbală este o abilitate pe care omul o deprinde în primii ani de viață. Deși originea acestei particularități nu a fost încă elucidată și este un subiect de dezbatere printre specialiști, dezvoltarea ei a fost cu siguranță un pas extrem de important în evoluția umană. Există studii ce atestă faptul că și alte viețuitoare posedă capacitatea de a comunica, dar delimitarea dintre acestea și om este trasată de precizia și complexitatea pe care numai comunicarea umană le prezintă.

În mod uzual, comunicarea este definită ca o acțiune ce permite schimbul de mesaje, idei și gânduri, fiind esențială în relațiile interumane. Pentru ca acest proces să aibă loc, este necesară existența unor componente: emițătorul, codul, canalul, informația și receptorul. Până în istoria recentă a omenirii, acestea au impus unele limitări. Distanța dintre emitor și receptor nu putea să fie prea mare, limba folosită trebuia să fie cunoscută de către ambele părți, iar canalul era reprezentat de aer. Avansul tehnologic, reprezentat mai ales de apariția telefonului și a radioului, au adus pentru prima dată modificări în acest scenariu clasic de comunicare, limitările fiind eliminate. Problema distanței a fost depășită, iar mesajele au devenit semnale electrice.

Dezvoltarea ulterioară a sistemelor de calcul a atras după sine un concept nou, acela de interacțiune om-mașină. Necesitatea acestui domeniu provine din dorința de a profita cât mai mult de avantajele utilizării calculatoarelor, asigurând în același timp un mod cât mai facil de a realiza acest lucru. Un exemplu în acest sens este apariția interfețelor grafice, acestea au înlesnit activitatea utilizatorilor mai puțin experimentați, oferind o interacțiune mult mai intuitivă decât interfețele bazate pe comenzi text.

Având în vedere aceste considerente, nici comunicarea verbală nu a fost lăsată deoparte. Deoarece vorbirea este un proces natural, ce se desfășoară fără depunerea unui efort considerabil, s-a dorit implementarea interacțiunii bazate pe voce. Sistemele de calcul au devenit capabile să interpreteze limbajul uman sau chiar să îl reproducă folosind tehnologia sintezei vocale.

O ramură importantă din acest domeniu este recunoașterea automată a vorbirii (ASR). Aceasta presupune convertirea semnalului vocal în text. Tehnologia actuală a ajuns la un nivel de performanță notabil, transcrierea fiind realizată într-un timp foarte scurt și cu o acuratețe din ce în ce mai bună. Recunoașterea automată a vorbirii este un subiect de mare interes la nivel mondial, în prezent existând mai multe grupuri de cercetare ce au ca scop găsirea unor noi algoritmi sau îmbunătățirea celor existenți. Preocuparea crescută poate fi motivată de diversitatea aplicațiilor ce au la bază recunoașterea vorbirii: comenzi vocale, indexare audio, transcrierea vorbirii spontane, sisteme biometrice, etc.

Am interacționat pentru prima dată cu domeniul recunoașterii vocale în timpul anului 3 de studii, atunci când am aflat de activitatea laboratorului de cercetare Speed (Speech and Dialogue), din cadrul Facultății de Electronică, Telecomunicații și Tehnologia Informației. Deși este un domeniu foarte complex, mi-am dorit să cunosc mai mult, mai ales pentru că am avut ocazia să realizez unele proiecte cu caracter practic, în care noțiunile teoretice își dovedeau din plin aplicabilitatea. Așadar, interesul manifestat față de recunoașterea automată a vorbirii, împreună cu cel pentru dezvoltarea aplicațiilor software, reprezintă motivația temei abordate în această lucrare.

OBIECTIVE

Lucrarea de față are ca scop expunerea sub formă de aplicații software a unor sisteme de recunoaștere automată a vorbirii (ASR) și a unui sistem de detecție a cuvintelor cheie (KWS). Se vor prezenta etapele de creare a unor astfel de sisteme în corelație cu noțiuni teoretice din domeniul tehnologiei vorbirii. Apoi se vor realiza diverse analize de performanță pentru a stabili configurația optimă a sistemelor. În final, acestea vor fi utilizate la implementarea a două aplicații software. În mod concret, pentru realizarea acestor sarcini vor fi îndeplinite următoarele obiective:

1. Crearea modelelor (acustic, fonetic, ligvistic) pentru un sistem de recunoaștere de cifre în limba română. Evaluarea sistemului.
2. Crearea modelelor (acustic, fonetic, ligvistic) pentru un sistem de recunoaștere a frazelor în limba engleză corespunzătoare unui sistem interactiv cu răspuns vocal (IVR). Evaluarea sistemului.
3. Crearea unui sistem de detecție a cuvintelor cheie ce se regăsesc în frazele IVR. Evaluarea sistemului.
4. Implementarea serverului de autentificare ca modul component al unui serviciu web de autentificare prin voce. Crearea aplicației client și a bazei de date ce stochează informații despre utilizatorii serviciului web.
5. Implementarea sistemului de detecție a cuvintelor cheie sub forma unei aplicații client-server.

STRUCTURA LUCRĂRII

Capitolul 1 este o introducere teoretică ce are ca scop prezentarea fundamentelor ce stau la baza tehnologiei vorbirii.

Sunt descrise caracteristicile semnalului vocal, modul în care acesta este produs, dar și perceput, precum și procesul prin intermediul căruia este prelucrat în vederea obținerii parametrilor vocali. Este prezentată arhitectura unui sistem de recunoaștere automată a vorbirii, algoritmi utilizați, metricile de evaluare, precum și utilitățile ASR folosite.

A doua secțiune a capitolului prezintă principiile unui sistem de detecție a cuvintelor cheie. Sunt prezentate câteva metode de căutare ale cuvintelor cheie, precum și metricile de evaluare a sistemului.

Capitolul 2 cuprinde procesul de creare a două sisteme de recunoaștere automată a vorbirii: un sistem de recunoaștere a cifrelor în limba română și un sistem de recunoaștere a frazelor în limba engleză, corespunzătoare unui sistem interactiv cu răspuns vocal. Se vor prezenta experimente și rezultate din care pot fi desprinse concluzii despre performanța sistemelor și a algoritmilor folosiți.

Capitolul 3 descrie procesul de creare a unui sistem de detecție a cuvintelor cheie. Se vor analiza modelele de limbă potrivite pentru o astfel de sarcină, iar la final se va realiza evaluarea sistemului.

Capitolul 4 este dedicat serviciului web de autentificare prin voce. Se va prezenta arhitectura serviciului, precum și modul de funcționare din punctul de vedere al unui utilizator. Apoi se vor sumariza tehnologiile folosite și se vor oferi detalii de implementare. Capitolul se încheie cu o descriere a dezvoltărilor ulterioare ce pot fi efectuate, însoțite de concluzii.

Capitolul 5 descrie aplicația de detecție a cuvintelor cheie. În mod similar cu capitolul precedent, se va prezenta arhitectura aplicației și modul de funcționare din perspectiva utilizatorului, iar ulterior se vor aduce la cunoștință tehnologiile folosite și detaliile de implementare. În încheiere, se vor prezenta unele concluzii și direcții viitoare de dezvoltare.

CAPITOLUL 1

DOMENIUL TEHNOLOGIEI VORBIRII

1.1 STRUCTURA UNUI SISTEM DE RECUNOAȘTERE AUTOMATĂ A VORBIRII

Recunoașterea automată a vorbirii (ASR) este o tehnologie ce permite unui sistem de calcul să transforme vorbirea în text. Intrarea unui astfel de sistem este reprezentată de un semnal vocal, în timp ce ieșirea sistemului este constituită dintr-o secvență de cuvinte corespunzătoare transcrierii.

Din punct de vedere probabilistic, acest proces implică determinarea celei mai probabile secvențe de cuvinte W^* , dat fiind mesajul vorbit X :

$$W^* = \operatorname{argmax} P(W | X).$$

Conform regulii lui Bayes rezultă:

$$W^* = \operatorname{argmax} P(W | X) = \operatorname{argmax} \frac{P(X | W) * P(W)}{P(X)} = \operatorname{argmax} P(X | W) * P(W)$$

Probabilitatea mesajului vorbit $P(X)$ este independentă de secvență de cuvinte W și poate fi ignorată. În acest fel, sarcina inițială poate fi descompusă în alte două sarcini mai mici:

- Estimarea probabilității mesajului vorbit, dată fiind secvență de cuvinte pronunțată $P(X|W)$.
- Estimarea probabilității apriori a secvenței de cuvinte $P(W)$.

Prima probabilitate poate fi determinată folosind un **model acustic**, în timp ce a doua probabilitate se obține folosind un **model de limbă**.

Figura de mai jos ilustrează arhitectura unui sistem ASR:

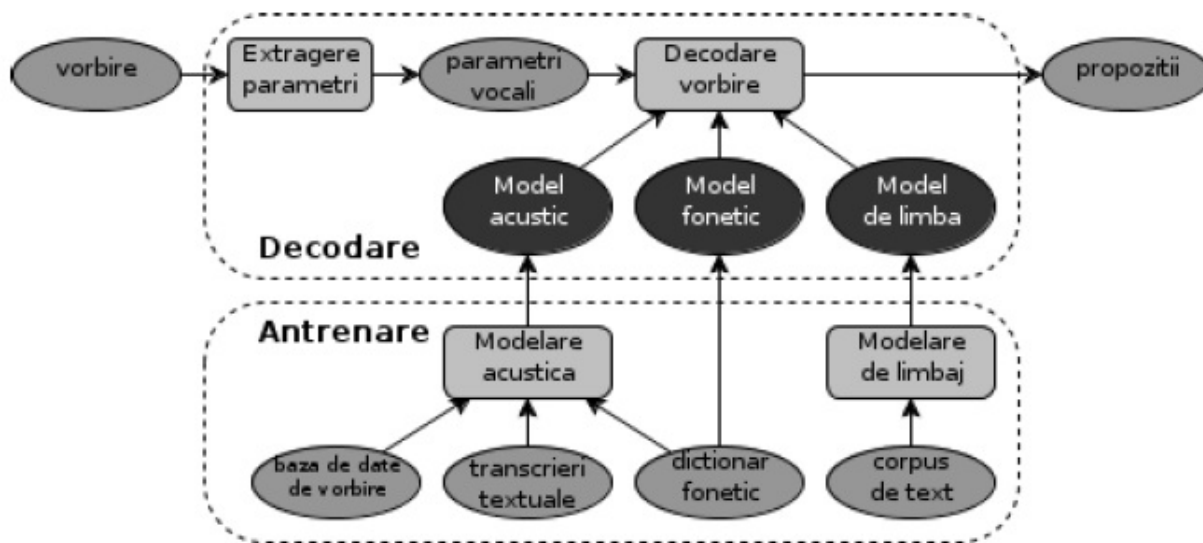


Figura 1.1 Arhitectura unui sistem ASR. Sursa: [1]

Un sistem ASR implică două faze distincte: o fază de antrenare și o fază de decodare. Recunoașterea vorbirii se realizează în faza de decodare, folosind parametrii vocali extrași din semnalul vocal și modelele (acustic, fonetic, lingvistic) dezvoltate în faza de antrenare.

Modelul acustic estimează probabilitatea unui mesaj vorbit, dată fiind secvență de cuvinte pronunțată. Acesta este construit pe baza unui set de înregistrări audio cât mai mari pentru care este cunoscută transcrierea text. El oferă o reprezentare statistică a fiecărui fonem. Fonemul este unitatea sonoră cea mai mică a limbii.

Modelul de limbă contribuie la îmbinarea cuvintelor astfel încât să poată fi obținute propoziții cu sens. În mod uzual se folosesc două tipuri de modele de limbă. Modelele probabilistice, obținute prin calculul probabilităților de apariție a cuvintelor sau de succesiune a lor, observând o cantitate cât mai mare de text. Al doilea tip este reprezentat de modelele de tip gramatică, acestea fiind modele bazate pe reguli, ce consistă într-un set de reguli ce definesc ordinea de apariție a cuvintelor.

Modelul fonetic are rolul de a realiza legătura între modelul acustic și modelul de limbă. Acesta este sub forma unui dicționar ce specifică pronunția cuvintelor. Fiecărui cuvânt din vocabular îi este asociată o secvență de foneme.

1.1.1 Producerea semnalului vocal

Cunoașterea aparatului fonator uman oferă unele informații importante în vederea proiectării sistemelor de recunoaștere automată a vorbirii. Astfel, vorbirea este produsă de aerul eliminat din plămâni, modulat mai departe de către componentele tractului vocal, rezultând astfel un semnal în gama de frecvență audio. În 1960, Gunnar Fant a definit producerea vorbirii cu ajutorul unui model sursă-filtru. Aparatul respirator, prin aerul expirat, formează semnalul de excitație. Acesta ajunge în trahee, la capătul căruia se află laringele, cu rol de a modula presiunea aerului. Aici se găsesc corzile vocale, care prin depărtare formează o deschidere numită glotă. Sunetele vocalizate apar datorită vibrației cvasi-periodice a corzilor, în timp ce în cazul sunetelor nevocalizate nu există această vibrație, ele fiind cauzate de turbulențe la nivelul tractului vocal. Tractul vocal, ce are rolul unui filtru, cuprinde elementele sistemului fonator aflate deasupra corzilor vocale. Acesta este format din diferite cavități rezonatoare (faringiană,

nazală, bucală) și organele articulatorii (limba, vâl palatin, palatul dur, dinții, buzele). Periodicitatea semnalului de excitație se numește frecvență fundamentală. Acesta este diferită în funcție de caracteristicile fiecărei persoane. La bărbați, frecvența fundamentală medie este de 132 Hz, în timp ce la femei se află în jurul valorii de 223 Hz. [2]

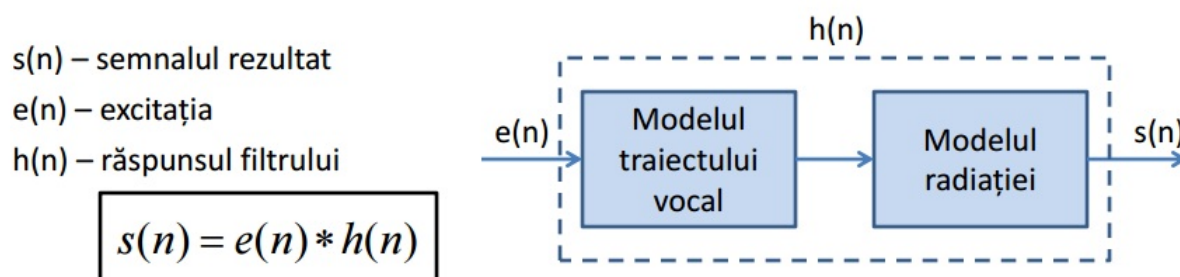


Figura 1.2 Modelul sursă-filtru al producerii vorbirii. Sursa:[3]

În funcție de modul în care sunt articulate, fonemele se împart în vocale și consoane (fricative, nazale, stopate). Vocalele sunt caracterizate de intensitatea crescută a energiei în raport cu alte foneme. Producerea lor are loc fără obstrucții la nivelul tractului vocal. În schimb, producerea consoanelor este condiționată de obstrucționarea aerului de către diferite părți componente ale tractului vocal.

1.1.2 Percepția semnalului vocal

Așa cum se poate observa în schema arhitecturii unui sistem ASR, semnalul vocal nu este utilizat în forma lui brută. Premergător etapei de antrenare, cât și a celei de decodare, semnalul vocal este supus mai multor operații ce au ca scop obținerea unui set de parametri vocali. Acești parametri conțin informații relevante pentru sarcina de recunoaștere automată a vorbirii, iar cei mai utilizați sunt coeficienții mel-cepstrali (MFCC).

Așa cum un sistem de sinteză de vorbire imită producerea semnalului vocal, un sistem de recunoaștere automată a vorbirii imita modul în care urechea umană percepe semnalul vocal. Pentru extragerea parametrilor vocali este necesar să se cunoască câteva detalii despre percepție.

Sistemul auditiv uman interpretează frecvențele într-un mod neliniar. Tonul fundamental este perceput liniar în gama de frecvențe 0-1000 Hz [4]. Peste această valoare, scara devine logaritmică. În acest sens a fost dezvoltată scara mel, definită de următoarea formulă:

$$F_{mel} = \frac{1000}{\log(2)} \times \left[1 + \frac{F_{Hz}}{1000} \right].$$

F_{mel} este frecvența rezultată pe scara mel, iar F_{Hz} este frecvența normală.

Prin aplicarea scării mel, parametrii vocali vor fi grupați în jurul frecvențelor joase, în zona tonului fundamental perceput. Astfel, semnalul va fi mult mai bine reprezentat din punct de vedere al sistemului auditiv uman.

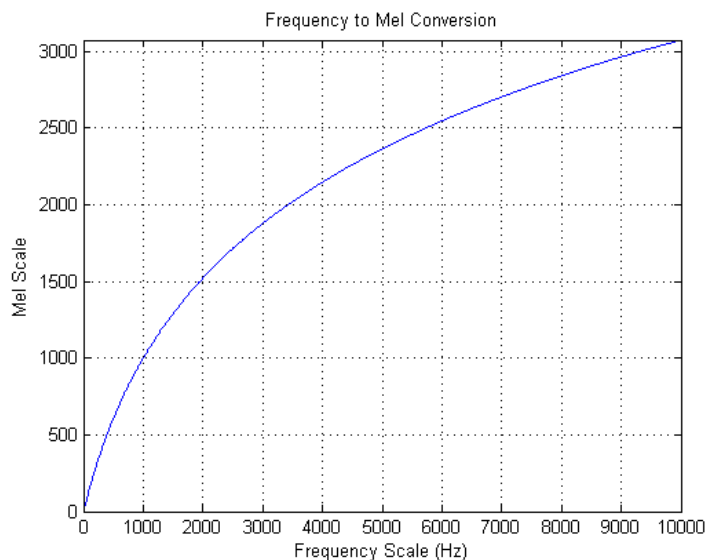


Figura 1.3 Scara Mel [5]

1.1.3 Parametri acustici

Figura de mai jos conține procesul de extragere al parametrilor MFCC, aceștia sunt în număr de 7. Se va face o descriere pentru fiecare parametru.

- A. Într-o primă fază are loc procesul de preaccentuare care are rol de a crește energia semnalului la frecvență înaltă.

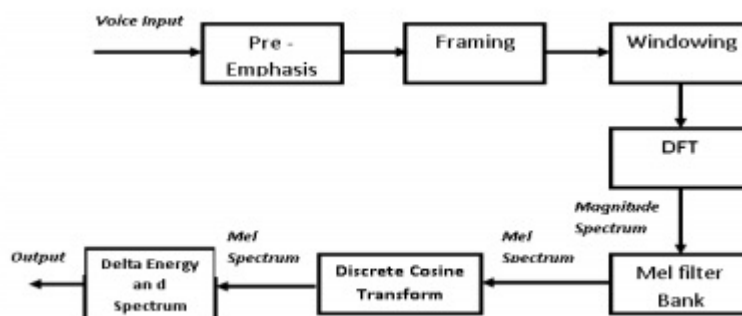


Figura 1.4 Metodă de extragere a MFCC. Sursa:[6]

- B. Deoarece semnalul vocal nu este staționar, analiza spectrală nu poate fi realizată pe întreg semnalul, ci pe porțiuni (ferestre) mai mici, cvasi-stationare cu o durată de 20-30 milisecunde. Ferestrele vor avea o perioadicitate de 10-15 milisecunde [7]. Ferestrele vor avea un timp de suprapunere pentru evitarea pierderii unor zone de semnal între acestea. Ferestrele pot fi de mai multe tipuri, iar cea mai uzuală este detaliată la punctul următor. Orice tip de fereastră are un spectru de tip trece jos (un lob principal la frecvențe joase și mai mulți lobi secundari atenuați diferit).
- C. Fiecare cadru obținut la pasul anterior se va multiplica cu o fereastră Hamming. Scopul acestui pas este de corectare a discontinuităților care pot să apară.

$$W[n] = \begin{cases} 0,54 - 0,46 \cdot \cos \frac{2\pi n}{N-1}, & 0 \leq n \leq N \\ 0, & \text{in rest} \end{cases}$$

- D. Așa cum am descris la secțiunea 1.1, semnalul vocal este determinat de convoluția în timp a semnalului de excitație dat de aerul din plămâni și răspunsul în timp al filtrului reprezentat de tractul vocal. Acesta din urmă fiind relevant pentru extragerea parametrilor MFCC. Deoarece cele două componente nu pot fi separabile în domeniul timp se va aplica transformata Fourier rapidă pentru a fi translatată în domeniul frecvență. Prin această transformare se ajunge la o relație care poate fi logaritmată putându-se astfel ajunge la o relație liniară de tip sumă care este separabilă.

$$S(n) = E(n) * h(n) \xLeftrightarrow{T.F.} S(\omega) = E(\omega) \cdot H(\omega)$$

$$S(\omega) = E(\omega) \cdot H(\omega) \xRightarrow{\log} \log(S(\omega)) = \log(E(\omega)) + \log(H(\omega))$$

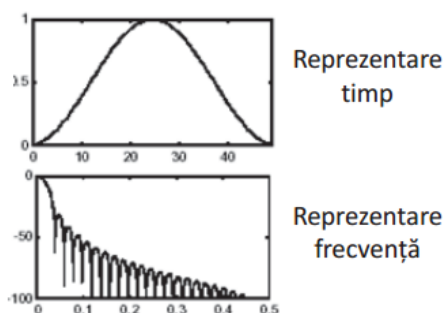


Figura 1.5 Reprezentarea în timp și frecvență a semnalului. Sursa:[8]

- E. Spectrul rezultat în urma aplicării transformatei Fourier se va netezi cu ajutorul unui banc de filtre mel triunghiulare. Este calculat spectrul mediu în jurul frecvenței centrale a fiecărei benzi, iar fiecare filtru va ocupa o bandă mai largă în funcție de ordinea din filtru, așa cum este prezentat în figura următoare:

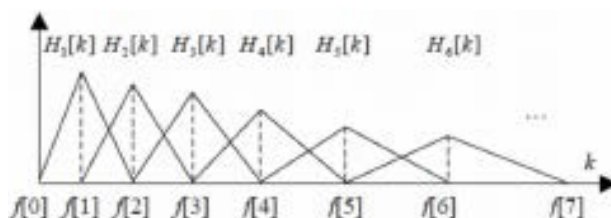


Figura 1.6 Bancul de filtre Mel. Sursa:[9]

- F. Următorul pas este convertirea spectrului logaritmic pe scara MEL înapoi în domeniul timp folosind transformata cosinus discretă. Rezultatul acestei conversii fiind chiar coeficienții MFCC. Transformata cosinus discretă are capacitatea de a concentra informația spectrală într-un număr mai mic de parametri și de a decorela aceste valori.
- G. Dintr-o fereastră de analiză se extrag de obicei 13 parametrii MFCC dar uneori acest vector acustic este extins fiind adăugate și derivatele funcție de timp de ordinul I și II (parametrii delta, delta+delta numiți și coeficienți de viteză și accelerație). În final se ajunge la un vector acustic cu 39 parametri.

1.1.4 Modele acustice

Modelul acustic are rolul de a estima probabilitatea că mesajul vorbit să fi fost generat de o secvență de cuvinte. În acest caz ar fi necesar să se realizeze observații pe baza unui set de cuvinte de la mai mulți vorbitori creând astfel un model matematic pentru fiecare cuvânt în

parte. Totuși această abordare nu este una eficientă. O unitate de vorbire de bază trebuie să aibă 3 caracteristici [10]:

- Precizie: este nevoie ca unitatea să rezeze realizarea acustică în contexte diferite;
- Antrenabilitate: datele de care dispunem trebuie să fie îndeajuns pentru o estimare cât mai exactă a unității;
- Generalizabilitate: posibilitatea construirii oricărui cuvânt având la bază un set de unități.

Deoarece cuvintele nu sunt antrenabile și nici generalizabile se impune folosirea unor unități de bază sublexicale numite foneme. Acestea pot îndeplini cu ușurință cele trei caracteristici enumerate mai sus. Dar există totuși un inconvenient: având în vedere sistemul articulator uman un fonem este dependent de contextul în care apare, el depinzând de fonemele vecine. Pentru înlăturarea acestui inconvenient s-a ajuns la concluzia că fiind necesară trecerea la analiza fonemelor luate în context cu cel anterior și următor ajungându-se astfel la noțiunea de trifoneme.

Deoarece există foarte multe trifoneme ce conduc la un număr foarte mare de parametri ce trebuie estimați, pentru a se ajunge la un optim între precizie și antrenabilitate se procedează la gruparea trifoanelor funcție de stări. Fiecare astfel de grup de stări (stări Markov) au fost denumite senone, acestea fiind în momentul de față cele mai utilizate unități acustice de bază.

Având în vedere definiția de mai sus a unui trifonem ca un fonem ce depinde de context cu stările din vecinătatea sa el poate fi modelat folosind platforma HMM-GMM.

Modelul Markov Ascuns (HMM) este un automat cu stări finite, el este considerat ascuns deoarece observatorul nu cunoaște secvența de stări ci doar vectorul de proprietăți acustice al fiecărei stări (MFCC) generat de o funcție de densitate de probabilitate.

În figura de mai jos este prezentată structura unui HMM.

Următorii parametri caracterizează un HMM:

Un set de stări:

$$S = S_1 S_2 S_3 \dots$$

Un set de probabilități de tranziție între stări de forma:

$P(s_i | s_j)$ - reprezentând probabilitatea de tranziție de la starea i la starea j ;

Observații probabilistice:

$P(X | s_i)$ - fiind probabilitatea că observația x să fi fost generată de starea i .

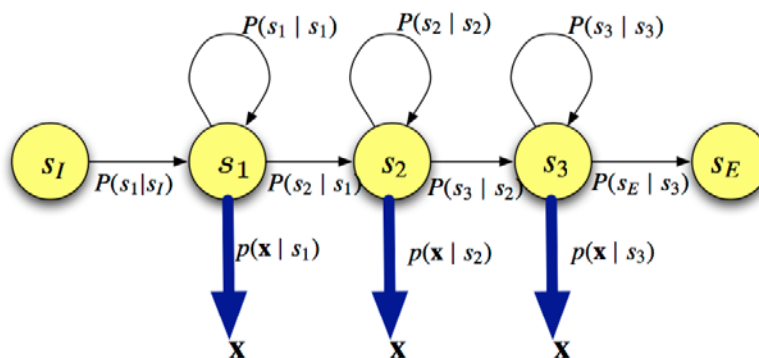


Figura 1.7 Structura unui HMM

Un GMM (model de mixturi Gaussiene) este probabilitatea ca observația x să fi fost generată de starea i . El este format dintr-o sumă de funcții Gaussiene, fiecare caracterizat de o medie și o dispersie proprie.

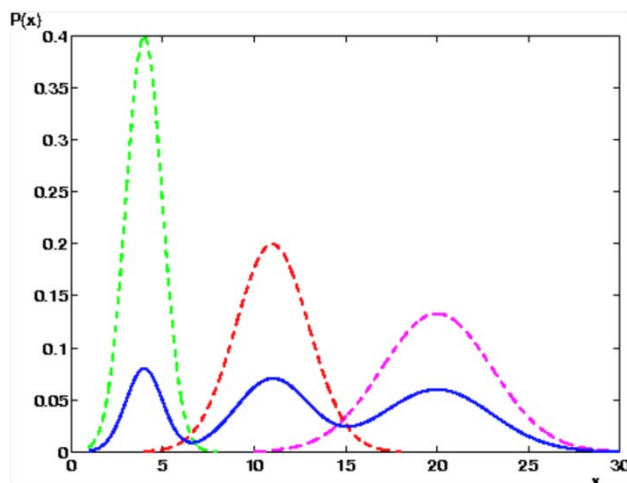


Figura 1.8 Mixtură de gaussiene. Sursa:[11]

Densitatea de probabilitate reprezentată cu albastru în figură de mai sus este dată de suma celorlate trei.

$$P(x|\mu, \sigma^2) = N(x; \mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$$

1.1.5 Modele de limbă

Scopul final de a obține propoziții cu sens din îmbinarea cuvintelor este dat de modelul de limbă. În mod uzual se folosesc două tipuri de modele de limbă.

- a. Model probabilistic (de tip n-gram) care poate fi obținut prin calculul probabilităților de apariție a cuvintelor sau de succesiune a lor prin observarea unui volum cât mai mare de text. Probabilitatea este calculată în funcție de istoricul cuvântului și anume de secvența numărului finit de cuvinte care îl preced.

$$P(w) = P(w_1, w_2, \dots, w_n) = P(w_1) * P(w_2|w_1) * P(w_n|w_1, w_2, \dots, w_{n-1})$$

Astfel, probabilitatea unei secvențe de cuvinte este descompusă în estimarea probabilităților unui singur cuvânt, dată fiind secvența de cuvinte anterioară acestuia.

Cele mai uzuale sunt modelele de tip 3-gram și 2 gram unde este necesară o istorie de două și respectiv un cuvânt. În cazul unui model 2-gram, calculul probabilității unei perechi de cuvinte se calculează în felul următor:

$$P(w_j|w_i) = \frac{\text{count}(w_i, w_j)}{\sum \text{count}(w_i, w)}$$

Probabilitatea de apariție a perechii de cuvinte (w_i, w_j) este dată de numărul de apariții ale cuvântului w_i urmat de cuvântul w_j , în raport cu numărul de apariții ale aceluiași cuvânt w_i , urmat de alte cuvinte.

Pentru obținerea unor probabilități cât mai corecte este necesară analiza unei cantități foarte mari de text.

- b. Model bazat pe reguli, de tip gramatică cu stări finite (FSG). Constă într-un set de reguli ce definesc ordinea de apariție a cuvintelor. Un astfel de model are forma unui graf, unde nodurile sunt cuvinte, iar arcele sunt tranziții între cuvinte. Fiecare arc poate fi

caracterizat de o probabilitate de tranziție. Drumurile grafului reprezintă secvențe de cuvinte permise.

1.1.6 Modelul fonetic

Modelul fonetic are rolul de a conecta modelul acustic cu modelul de limbă. Acest model are forma unui dicționar de pronunții. Fiecărui cuvânt din vocabular îi este atribuită o secvență de foneme ce reprezintă pronunția acelui cuvânt. În cazul existenței a mai multor variante de pronunție pentru un singur cuvânt, acestea vor fi specificate fiecare în parte.

Crearea unui model fonetic este o sarcină realizată de obicei manual, mai ales atunci când vocabularul este unul redus. Sarcina devine însă una complicată, dacă se dorește crearea unui sistem de recunoaștere pe întregul vocabular al unei limbi. Se pot folosi unele aplicații software denumite *fonetizatoare*, ce au implementate reguli de pronunție ale literelor în funcție de contextul în care apar.

1.1.7 Metrici de evaluare

Evaluarea unui sistem de recunoaștere automată a vorbirii este posibilă prin transcrierea unei baza de date de evaluare. Pentru această bază de date sunt cunoscute transcrierile reale ale clipurilor audio, numite transcrieri de referință. Prin compararea acestora cu transcrierile generate de sistemul de recunoaștere automată a vorbirii, sunt calculate ratele de eroare la nivel de propoziție și la nivel de cuvânt.

Rata de eroare la nivel de propoziție (SER) este dată de raportul dintre numărul de propoziții transcrise eronat și numărul total de propoziții. O propoziție este eronată dacă ea conține cel puțin un cuvânt greșit.

$$SER[\%] = \frac{\# \text{propozitii eronate}}{\# \text{total propozitii in referinta}} \times 100$$

Această rată este concludentă în situațiile în care transcrierea greșită a unui singur cuvânt conduce la anularea caracterului valabil al întregii propoziții (de exemplu, transcrierea greșită a unei cifre dintr-un număr de telefon va face acel număr inutil, în timp ce transcrierea greșită a unui cuvânt dintr-o știre difuzată la radio nu va afecta prea mult inteligibilitatea propoziției).

Rata de eroare la nivel de cuvânt se calculează prin utilizarea unui algoritm de aliniere a transcrierii ipotetice cu transcrierile de referință. Pot apărea erori de inserție (apar cuvinte în plus față de referință), erori de substituție (sunt transcrise alte cuvinte decât cele din referință) și erori de ștergere (nu sunt transcrise cuvinte ce apar în referință). Formula de calcul este următoarea:

$$WER[\%] = \frac{\# \text{Erori de inserție} + \# \text{Erori de substituție} + \# \text{Erori de ștergere}}{\# \text{Total cuvinte in referinta}} \times 100$$

1.1.8 Utilitare pentru recunoașterea automată a vorbirii

CMU Sphinx

CMU Sphinx este un utilitar cu sursă deschisă, scris în limbajul Java, utilizat pentru crearea sistemelor de recunoaștere automată a vorbirii. Acesta a fost dezvoltat în cadrul Carnegie Mellon University. Este o aplicație modulară, fiind formată din 3 blocuri principale:

- *Frontend*
- *Decoder*
- *Knowledge database*

Această arhitectură poate fi observată în imaginea următoare:

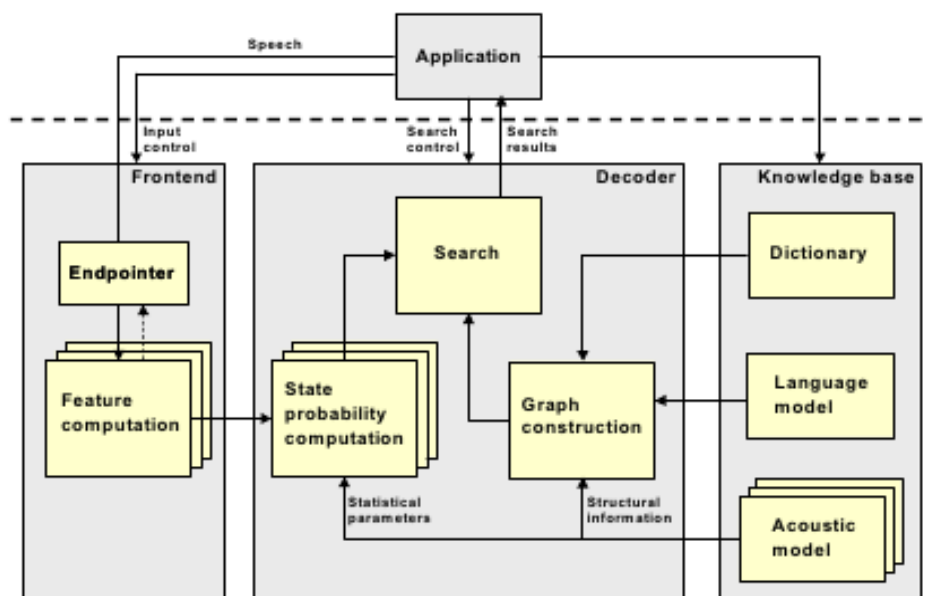


Figura 1.9 Arhitectură CMU Sphinx. Sursa:[12]

Blocul *Frontend* are rol în extracția parametrilor vocali (MFCC sau PLP). Blocul *Knowledge base* conține resursele acustice, fonetice și ligvistice necesare sarcinii de recunoaștere. Blocul *Decoder* realizează transcrierea semnalului audio în text.

Modelele acustice folosite în CMU Sphinx sunt antrenate folosind platforma HMM-GMM. Modelarea fonetică permite utilizarea atât a fonemelor independente de context, cât și a celor dependente (trifoneme). Modelele de limbă pot fi modele probabilistice, dar și modele de tip gramatică, scrise în formatul JSGF (Java Speech Grammar Format).

Kaldi

Un alt utilitar folosit în crearea sistemelor de recunoaștere automată a vorbirii este Kaldi. Acesta a fost scris în limbajul C++, fiind de asemenea cu sursă deschisă.

Arhitectura Kaldi poate fi observată în următoarea imagine:

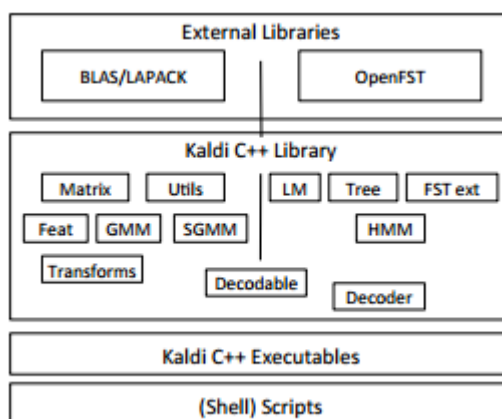


Figura 1.10 Arhitectură Kaldi. Sursa:[13]

Utilitarul depinde de două librării externe, OpenFST (pentru implementarea platformelor cu stări finite) și BLAS/LAPACK (o colecție de biblioteci matematice). Librăria Kaldi propriu-zisă conține module de extragere a parametrilor (MFCC sau PLP), antrenarea modelelor acustice (bazate pe GMM și mai nou pe rețele neuronale), compilarea modelelor de limbă și bineînțeles decodarea audio.

Kaldi permite crearea mai multor tipuri de sisteme acustice:

- *Monophone*: folosește foneme independente de context;
- *Triphone (TRI1)*: folosește foneme cu 3 stări, ce depind de contextul în care apar;
- *LDA+MLLT (TRI2)*: Este o alternativă la folosirea derivatelor coeficienților cepstrali. LDA (Linear Discriminant Analysis) are ca scop reducerea spațiului de caracteristici, fiind păstrate numai cele care conțin informație discriminatorie.
- *LDA+MLLT+SAT (TRI3)*: În plus față de modelul precedent, acesta realizează și adaptarea la vorbitor (Speaker Adaptive Training).
- *LDA+MLLT+SAT+MMI (TRI3_MMI)*: MMI (Maximum Mutual Information) tinde să maximizeze probabilitatea a posteriori a secvenței rostite.

1.2 STRUCTURA UNUI SISTEM DE DETECȚIE A CUVINTELOR CHEIE

Procesarea informației a devenit o preocupare majoră a zilelor noastre, iar procesarea limbajului vorbit este una dintre aceste provocări.

Detecția cuvintelor cheie are ca scop determinarea existenței unor cuvinte cheie într-o secvență de vorbire. Este o sarcină necesară în situația în care se dorește indentificarea unor termeni roștiți într-o cantitate foarte mare de date audio.

Un studiu din 2014 al companiei Domo, numit "Data Never Sleeps 2.0" [14], a arătat printre altele că în fiecare minut al unei zile, pe Youtube sunt încărcate 72 ore de conținut audiovizual, după ce în 2012 cantitatea era de 48 ore de conținut audiovizual. În numai doi ani, valoarea a crescut de 1.5 ori. Așadar, în momentul actual, dacă tendința s-a păstrat, este posibil ca un număr de aproximativ 100 de ore de conținut să fie urcate pe Youtube în fiecare minut. Așadar, dacă ne raportăm la aceste cifre, căutarea termenilor vorbiți în conținutul audio este imposibilă să fie realizată prin ascultarea propriu-zisă a înregistrărilor. Folosind tehnica de procesare modernă, sistemele de detecție a cuvintelor cheie sunt singurele care pot îndeplini această sarcină în timp util.

1.2.1 Căutarea prin lattice

Pentru înțelegerea acestei metode este mai întâi necesar să se cunoască definiția unei lattice. Latticea se poate defini ca un graf $G(N,A)$, unde N reprezintă nodurile și A sunt arcele. În contextul recunoașterii automate a vorbirii, rezultatul decodării poate fi structurat sub forma unei lattice, unde fiecare arc are o anumită probabilitate și reprezintă o secvență alternativă de cuvinte rostite. Alegând calea cu cea mai mare probabilitate se va obține cea mai bună ipoteză a transcrierii, în acest fel obținându-se transcrierea textului în sine.

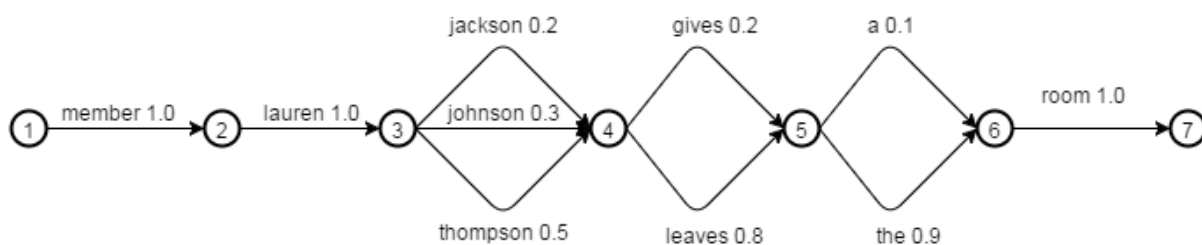


Figura 1.11 Exemplu de lattice

Această metodă este implementată sub forma unui modul KWS în utilitarul Kaldi. Kaldi realizează sarcină de KWS în mod offline în doi pași:

- vorbirea este decodată cu un sistem ASR generând astfel o lattice de cuvinte;

- cuvintele cheie sunt căutate prin lattice.

În continuare se va detalia acest pas. Mai întâi este creată o listă de cuvinte cheie ce se dorește a fi identificate apoi acestora li se asociază câte un ID și sunt folosite pentru a se crea un automat cu stări finite (FST). De asemenea secvențele audio rostite în care se realizează căutarea le sunt asociate un ID. Cu latticele obținute la primul pas în felul următor:

ID-ul secvenței rostite | momentul începerii | momentul de final | probabilitatea posterioară

În acest fel pentru fiecare secvență de cuvinte se va cunoaște secvență audio corespunzătoare, momentele de timp între care se regăsește și probabilitatea. În final căutarea este realizată folosind automatul cu stări finite ce conține cuvintele cheie, aplicată latticeelor indexate. Principalul avantaj al acestei metode este reprezentat de posibilitatea de a seta un prag al probabilităților pe baza căruia cuvintele cheie găsite pot fi acceptate sau respinse. Fixarea unui prag jos va duce la acceptarea multor ipoteze crescând astfel șansa de a găsi cuvintele cheie căutate, dar în același timp probabilitatea unei alarme false este mărită. În schimb, fixarea unui prag înalt va asigura o rată scăzută a alarmelor false fiind posibil în același timp că aparițiile unor cuvinte cheie să nu fie contabilizate.

1.2.2 Căutarea prin transcriere text

Această metodă este una destul de simplă neavând de a face cu lattice și cu probabilitățile arcelor pentru a lua o decizie referitoare la apariția cuvintelor cheie. Este de asemenea o metodă care se realizează în doi pași:

- secvența audio este decodată folosind un sistem ASR, generând o transcriere text;
- cuvintele cheie sunt căutate direct în transcrierea text.

Spre deosebire de metoda precedentă la acesta există doar două posibilități: cuvântul cheie să se regăsească sau nu în secvența rostită. Performanța acestei metode depinde foarte mult de acuratețea sistemului ASR, aceasta fiind influențată puternic de către modelul de limbă și modelul acustic folosit.

Pe baza rezultatelor experimentale se poate observa că un model de limbă foarte bogat produce apariția unor probleme legate de resursele mașinii de calcul: memorie, CPU și timpul de decodare.

1.2.3 Metrici de evaluare

Acest subcapitol prezintă unele metrice folosite pentru a măsura performanța unui sistem KWS. Un sistem perfect trebuie să fie capabil să detecteze în cadrul secvenței rostite, locul exact de apariție al fiecărui cuvânt cheie, fără a fi generate false detecții. Pentru a realiza evaluarea unui sistem KWS, este necesară mai întâi rularea sistemului de detecție a cuvintelor cheie. Pe baza înregistrărilor audio existente în baza de date, a fișierului ECF (ce conține indexarea fișierelor audio) și a fișierului KWList (ce conține cuvintele cheie căutate), este generat KWSList (fișierul cu detecțiile rezultate). Pentru fiecare cuvânt cheie sunt afișate fișierele audio în care acesta se găsește, momentul de timp al apariției, durată, o probabilitatea de apariție (cu o valoare normată între 0 și 1) și o decizie de interpretare a acestei probabilități (da/nu). De asemenea este necesar un fișier de referință (RTTM) care să conțină o listă cu toate aparițiile cuvintelor cheie în toate secvențele audio rostite. Imaginea de mai jos prezintă schema de evaluare a sistemului KWS:

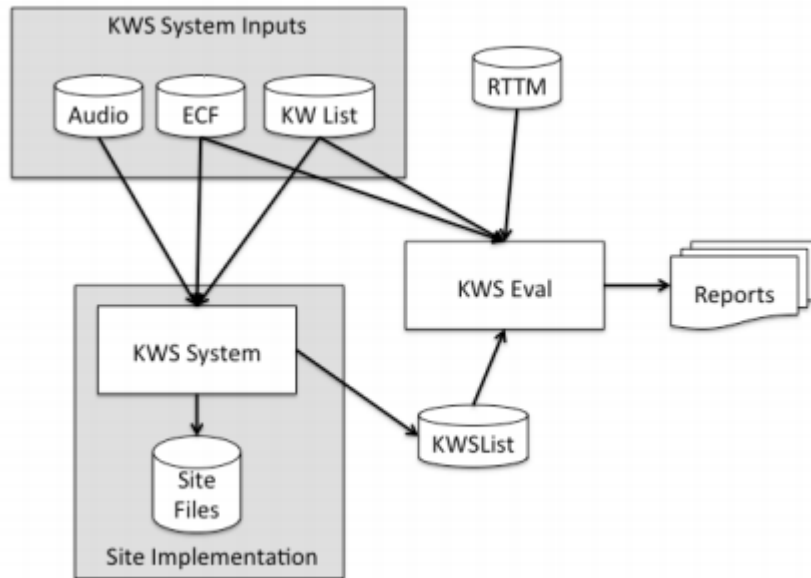


Figura 1.12 Schemă de evaluare a sistemului KWS. Sursa:[15]

Dacă existența unui cuvânt cheie în secvența rostită conduce la decizia "DA" a sistemul KWS, acest lucru va fi contabilizat ca o detecție corectă altfel va fi contabilizat că o eroare de neidentificare. Dacă sistemul KWS ia decizia "DA" pentru un cuvânt cheie care în realitate nu apare în secvența rostită acest lucru va fi contabilizat că o eroare de alarmă falsă. Aceste decizii ale sistemului sunt bazate pe un prag: dacă probabilitatea de apariție este mai mare sau egală decizia este "DA", altfel este "NU".

Probabilitatea erorii de neidentificare pentru un cuvânt cheie "q" și un prag "θ"

$$P_{miss}(q, \theta) = \frac{N_{miss}(q, \theta)}{N_{act}(q)}$$

unde:

$N_{miss}(q, \theta)$ = număr de erori de neidentificare pentru un cuvânt cheie "q" și un prag "θ"

$N_{act}(q)$ = numărul real de apariții ale cuvântului "q"

Probabilitatea erorii de alarmă falsă pentru un cuvânt cheie "q" și un prag "θ"

$$P_{fa}(q, \theta) = \frac{N_{fa}(q, \theta)}{N_{nt}(q)}$$

$N_{fa}(q)$ = numărul de alarme false corespunzătoare unui cuvânt cheie "q" și un prag "θ"

$N_{nt}(q)$ = numărul maxim de alarme false

$N_{nt}(q) = n_{tps} * T_{audio} - N_{act}(q)$; $n_{tps}=1$

T_{audio} = durata totală a secvențelor vorbite (durata fișierelor audio)

Probabilitățile medii de eroare peste toate cuvintele cheie sunt date de formulele

$$P_{miss}(\Theta) = \frac{1}{|Q|} \sum P_{miss}(q, \Theta)$$

$$P_{fa}(\Theta) = \frac{1}{|Q|} \sum P_{fa}(q, \Theta)$$

Term-Weighted Value (TWV) este o combinație ponderată a probabilităților celor două erori de neidentificare respectiv de alarmă falsă, calculat cu următoarea relație:

$$TWV(\Theta) = 1 - (P_{miss}(\Theta) + \beta * P_{fa}(\Theta))$$

Unde factorul β depinde de costul erorilor de neidentificare respectiv alarmă falsă.

Un alt parametru important este entropia normalizată C_{nxe} care măsoară cunoștințele apriori pe care sistemul KWS le are asupra referinței. Un sistem perfect ar trebui să aibă un parametru C_{nxe} care tinde la zero.

Curba de detecție a erorii (DET curve) este un grafic unde axele sunt reprezentate de cele două probabilități de eroare ($P_{miss}(\Theta)$, $P_{fa}(\Theta)$), aceasta fiind un compromis între acestea.

CAPITOLUL 2

CONSTRUCȚIA UNOR SISTEME DE RECUNOAȘTERE AUTOMATĂ A VORBIRII

2.1 CREAREA SISTEMULUI "RODIGITS"

"Rodigits" este un sistem de recunoaștere automată a vorbirii continue cu vocabular redus în limba română. Scopul acestuia este de a realiza transcrierea unor clipuri audio a câte 12 cifre fiecare. Au fost dezvoltate două mari versiuni ale sistemului "Rodigits", atât cu ajutorul utilitarului CMU Sphinx, cât și cu ajutorul utilitarului Kaldi.

Crearea acestui sistem a dus la îndeplinirea a două obiective propuse:

1. Obținerea unei acurateți cât mai mari pentru a putea folosi ulterior sistemul în cadrul aplicației descrisă în capitolul 4, "Serviciul web de autentificare prin voce".
2. Realizarea unei comparații directe între cele două utilitare pentru recunoașterea automată a vorbirii.

2.1.1 Baza de date audio "Rodigits"

Această bază de date de clipuri audio a fost folosită atât pentru antrenarea modelelor acustice ale sistemului, cât și pentru evaluarea sa. La crearea înregistrărilor audio au participat studenți din cadrul Facultății de Electronică, Telecomunicații și Tehnologia Informației, fiecare furnizând un set de 100 de clipuri audio a câte 12 cifre fiecare.

2.1.2 Antrenarea modelelor în vederea antrenării folosind utilitarul CMU Sphinx

Modelul fonetic

Acesta a fost creat prin transcrierea fonetică a fiecărei cifre din sistemul zecimal. Utilitarul CMU Sphinx folosește modele fonetice în format text, iar manualul acestuia recomandă utilizarea de foneme specifice pentru fiecare cuvânt. Fiecare cuvânt va trebui să modeleze în mod diferit fonemele în funcție de poziția la care acestea apar în cuvânt. Fișierul ce conține modelul fonetic arată în felul următor:

```
zero z_zero e_zero r_zero o_zero
unu u_unu1 n_unu u_unu2
doi d_doi o_doi i3_doi
trei t_trei r_trei e_trei i3_trei
patru p_patru a_patru t_patru r_patru u_patru
cinci k1_cinci1 i_cinci1 n_cinci k1_cinci2 i_cinci2
șase s1_șase a_șase s_șase e_șase
șapte s1_șapte a_șapte p_șapte t_șapte e_șapte
opt o_opt p_opt t_opt
nouă n_nouă o1_nouă w_nouă a1_nouă
```

Modelul de limbă

Deoarece acest sistem de recunoaștere automată a vorbirii are un vocabular foarte redus, modelul de limbă folosit este unul bazat pe reguli, de tip gramatică cu stări finite (FSG).

Pentru a obține gramatica într-un format compatibil cu utilitarul CMU Sphinx, se va crea un fișier în formatul Java Speech Grammar (JSGF):

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

În antetul gramaticii se poate observa identificatorul ce indică faptul că fișierul conține o definiție în format JSGF, precum și numele gramaticii, "rodigits". În corpul gramaticii se regăsește regula ce definește variabila "numbers". Separarea prin caracterul de disjuncție logică "|" a cuvintelor ce reprezintă cifrele din sistemul zecimal, semnifică faptul că variabilei "numbers" îi poate fi atribuită una dintre aceste valori. Caracterul asterisc indică posibilitatea cuvintelor de a se repeta.

Automatul cu stări finite rezultat în urma compilării modelului de limbă va fi sub forma unui graf, unde nodurile reprezintă cuvintele. Arcele grafului reprezintă tranzițiile între cuvinte și indică modul în care este permisă parcurgerea grafului, implicit înlănțuirea cuvintelor.

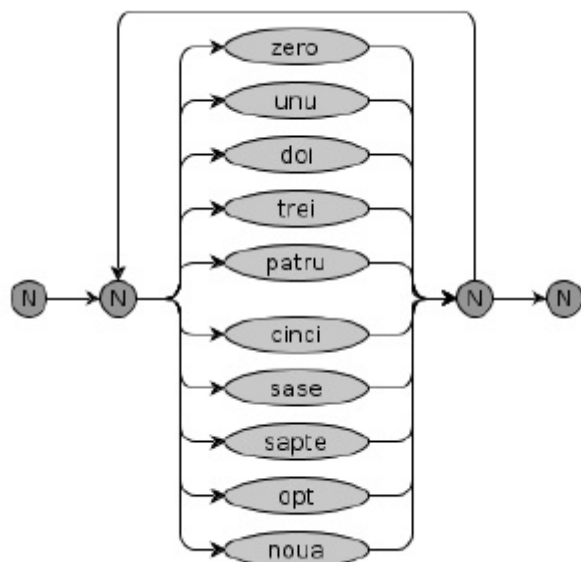


Figura 2.1 Modelul de limbă de tip gramatică Rodigits (CMU Sphinx). Sursa:[16]

Nodurile notate cu "N" sunt noduri de infrastructură. Ele asigură intrarea și de ieșirea din gramatica, precum și tranzițiile înapoi.

Modelul acustic

Pentru crearea modelului acustic este necesară existența următoarelor resurse: baza de date de clipuri audio, transcrierea text a clipurilor audio și modelul fonetic creat anterior. Pentru o mai bună acuratețe, acesta poate fi îmbunătățit prin adăugarea unor elemente ce modelează fonemele nonverbale: liniște, zgomot, etc.

Deoarece s-a dorit crearea unui sistem de recunoaștere automată a vorbirii cât mai robust, au fost create mai multe modele acustice ce au fost testate și optimizate treptat, până la obținerea unor rezultate satisfăcătoare. Pentru a înțelege mai bine necesitatea parcurgerii acestor pași intermediari, va fi prezentat pe scurt modul în care CMU Sphinx realizează antrenarea acustică. Folosind configurația standard a utilitarului, modelul acustic este construit cu unități de vorbire dependente de context de tip trifonem. Acestea sunt implementate cu ajutorul modelelor Markov ascunse cu 3 stări emise, unde fiecărei stări îi corespunde o distribuție de probabilitate de ieșire implementată cu ajutorul unei mixturi de Gaussiene. Datorită asemănării unor trifoneme, ele pot avea stări-model comune. Profitând de acest fapt și având în vedere numărul foarte mare de stări, s-a recurs la gruparea lor pe bază de similaritate, un astfel de grup fiind numit senone. În mod implicit, CMU Sphinx folosește 200 de senone, fiecare senonă modelând parametrii acustici de ieșire cu ajutorul a 8 densități de probabilitate Gaussiene.

Prin modificarea configurației standard, modelele dependente de context au suferit variații în ceea ce privește numărul de senone utilizat, precum și numărul de densități Gaussiene. De asemenea, au fost efectuate teste și cu modele independente de context, ce folosesc o singură densitate Gaussiană pentru fiecare fonem.

2.1.3 Experimente și rezultate

Experiment 1

Acest prim experiment constă în crearea unui **model acustic dependent de vorbitor**. Au fost selectate din baza de date "Rodigits" numai cele 100 de fișiere audio proprii. Acestea au fost împărțite în felul următor: pentru etapa de antrenare au fost folosite 80 de fișiere (primele 50 și ultimele 30), iar pentru etapa de testare s-au folosit restul de 20 de fișiere.

A. Rezultatele modelului independent de context:

WER[%]	22.9
SER[%]	100

Tabelul 2.1 Experiment 1. Rezultatele modelului independent de context

B. Rezultatele modelului dependent de context:

WER[%]		Densități Gaussiene per senone			
		1	2	4	8
Senone	100	5.4	0.8	0	0
	200	14.2	7.1	10.4	12.1

Tabelul 2.2 Experiment 1. Rezultatele modelului dependent de context (WER)

SER[%]		Densități Gaussiene per senone			
		1	2	4	8
Senone	100	35	10	0	0
	200	70	50	70	70

Tabelul 2.3 Experiment 1. Rezultatele modelului dependent de context (SER)

Se observă că modelul independent de context oferă o acuratețe cu certitudine mai slabă decât modelul dependent de context. Acesta din urmă realizează o transcriere perfectă în situația folosirii a 100 de senone și un număr de densități gaussiene cel puțin egal cu 4.

Experiment 2

Acest experiment presupune **evaluarea modelului dependent de vorbitor cu date de la alți 5 vorbitori necunoscuți**, fiind alese câte 20 de înregistrări de la fiecare dintre aceștia.

A. Rezultatele modelului independent de context:

WER[%]	59.6
SER[%]	99

Tabelul 2.4 Experiment 2. Rezultatele modelului independent de context

B. Rezultatele modelului dependent de context:

WER[%]		Densități Gaussiene per senone			
		1	2	4	8
Senone	100	69.8	63.2	65.2	67.6
	200	61.5	70.4	78.2	85.1

Tabelul 2.5 Experiment 2. Rezultatele modelului dependent de context (WER)

SER[%]		Densități Gaussiene per senone			
		1	2	4	8
Senone	100	100	88	85	91
	200	99	99	100	100

Tabelul 2.6 Experiment 1. Rezultatele modelului dependent de context (SER)

Așa cum era previzibil, rezultatele sunt destul de slabe. Acest lucru semnifică faptul că modelul dependent de vorbitor nu are capacitatea să se adapteze în cazul unor vorbitori necunoscuți.

Experiment 3

Având în vedere concluzia de la experimentul precedent, s-a trecut la crearea unui **model independent de vorbitor**. Au fost selectați din baza de date "Rodigits" 50 de vorbitori. Cele 100 de fișiere ale fiecăruia au fost împărțite în felul următor: pentru etapa de antrenare au fost folosite câte 80 de fișiere de la fiecare vorbitor (primele 50 și ultimele 30), în timp ce restul fișierelor au fost folosite în etapa de testare. Așadar, **testarea se va face cu vorbitori cunoscuți** sistemului, aceștia fiind identici cei implicați în procesul de antrenare.

A. Rezultatele modelului independent de context:

WER[%]	36.5
SER[%]	95.3

Tabelul 2.7 Experiment 3. Rezultatele modelului independent de context

B. Rezultatele modelului dependent de context:

WER[%]		Densități Gaussiene per senone			
		1	2	4	8
Senone	100	16.9	10.9	5.1	4.3
	200	21	12.9	9	8.5

Tabelul 2.8 Experiment 3. Rezultatele modelului dependent de context 1 (WER)

SER[%]		Densități Gaussiene per senone			
		1	2	4	8
Senone	100	65.7	52.2	33.2	29.7
	200	77	65.8	51.6	45

Tabelul 2.9 Experiment 3. Rezultatele modelului dependent de context 1 (SER)

Se observă că în acest caz al antrenării și testării folosind fișierele a 50 de vorbitori, modelul optim este cel cu 100 senone și 8 densități gaussiene. Având în vedere rezultatele de până la acest punct, este evident faptul că modelele dependente de context oferă o acuratețe superioară celor independente de context. Deoarece scăderea ratei de eroare se produce concomitent cu creșterea numărului de densități gaussiene, s-a trecut la testarea cu mai multe densități gaussiene. De asemenea, deoarece modelele ce folosesc 100 de senone au rezultate mai bune decât cele cu 200 de senone, testele următoare s-au efectuat folosind numai 100 senone.

WER[%]		Densități Gaussiene per senone			
		16	32	64	128
Senone	100	3.3	2.7	2.2	1.7

Tabelul 2.10 Experiment 3. Rezultatele modelului dependent de context 2 (WER)

SER[%]		Densități Gaussiene per senone			
		16	32	64	128
Senone	100	27.3	22.7	18.5	13.6

Tabelul 2.11 Experiment 3. Rezultatele modelului dependent de context 2 (SER)

Următoarele grafice prezintă tendința de scădere a ratelor de eroare în raport cu creșterea densităților gaussiene.

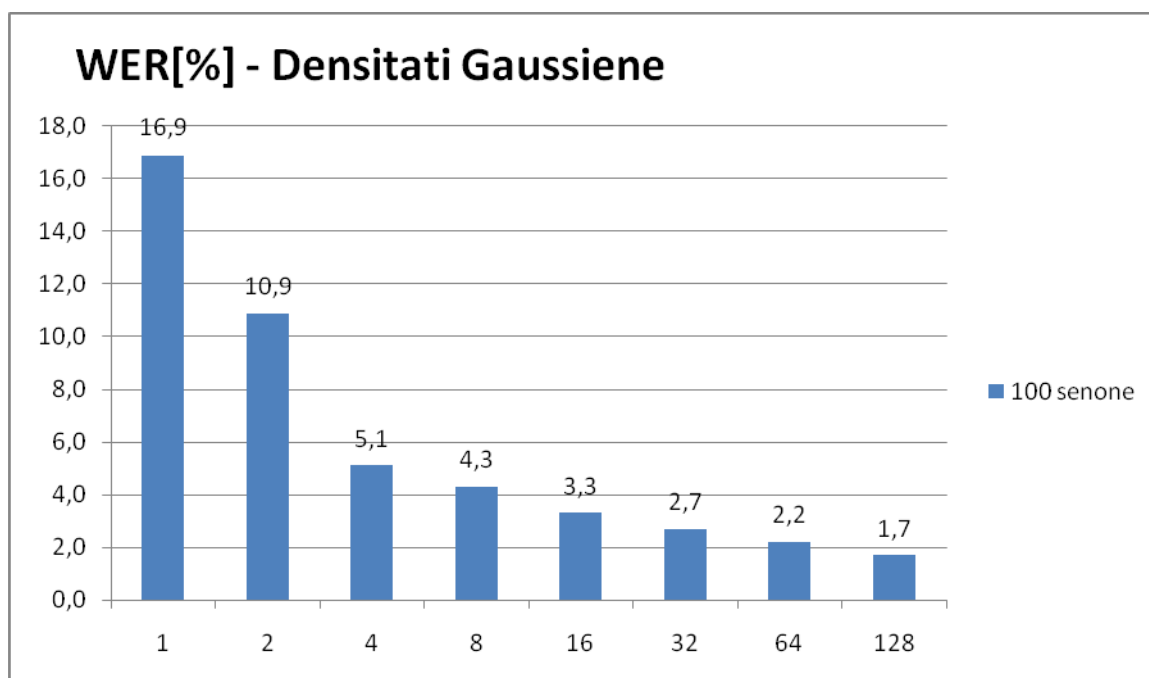


Figura 2.2 Raport WER - Densități gaussiene

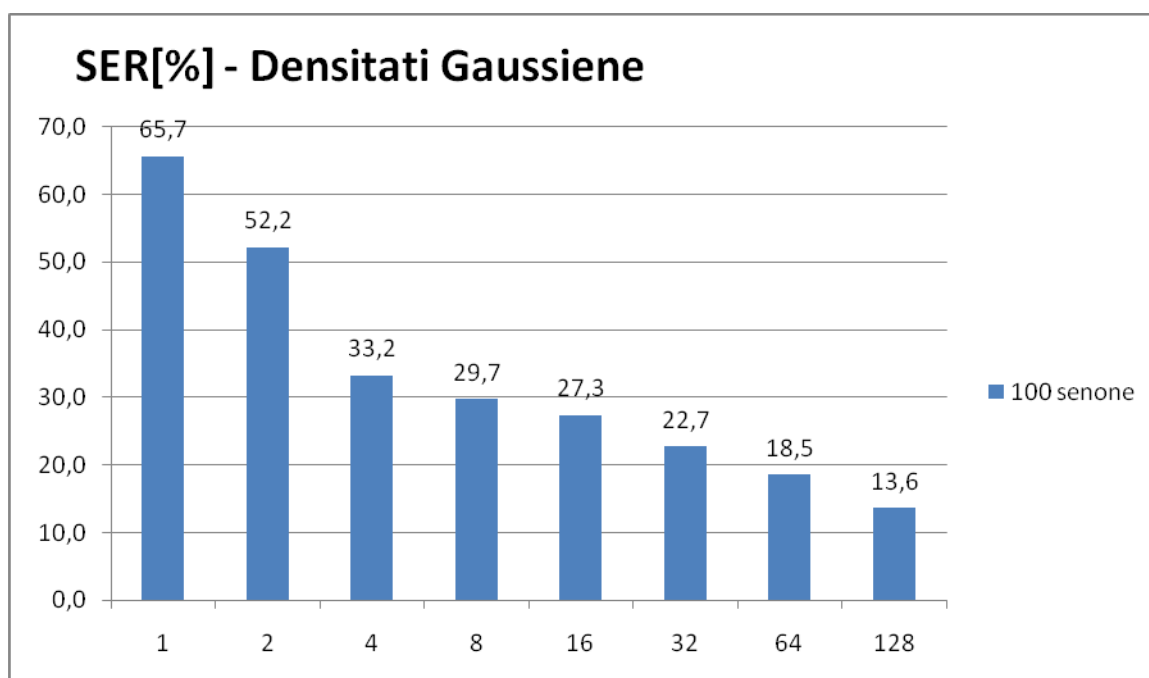


Figura 2.3 Raport SER - Densități Gaussiene

Mai departe, având în vedere performanțele crescute ale modelelor cu 100 de senone în comparație cu cele cu 200 de senone, s-a încercat o nouă scădere a acestei valori. Însă, așa cum se poate observa din tabelul următor, acest test nu a condus la o îmbunătățire a rezultatelor.

WER[%]		Densități Gaussiene per senone				
		1	2	4	8	128
Senone	100	16.9	10.9	5.1	4.3	1.7
	50	16.9	10.9	5.1	3.3	2

Tabelul 2.12 Variații ale densităților gaussiene și senonelor în raport cu WER

În urma acestor experimente, s-a constatat faptul că **modelul optim este cel cu 100 senone și 128 densități gaussiene**. Pentru a evidenția caracterul independent de vorbitor al modelului, s-a efectuat testarea sa cu fișierele a **20 de vorbitori necunoscuți**. Au fost obținute următoarele rate de eroare:

- **WER: 1.1%;**
- **SER: 11%.**

Rezultatele obținute sunt destul de interesante, mai ales din prisma faptului că eroarea pentru vorbitorii necunoscuți (1.1%) este mai mică decât cea pentru vorbitorii cunoscuți (1.7%). Pentru a afla cauza, s-a trecut la analiza erorii la nivel de cuvânt pentru fiecare vorbitor.

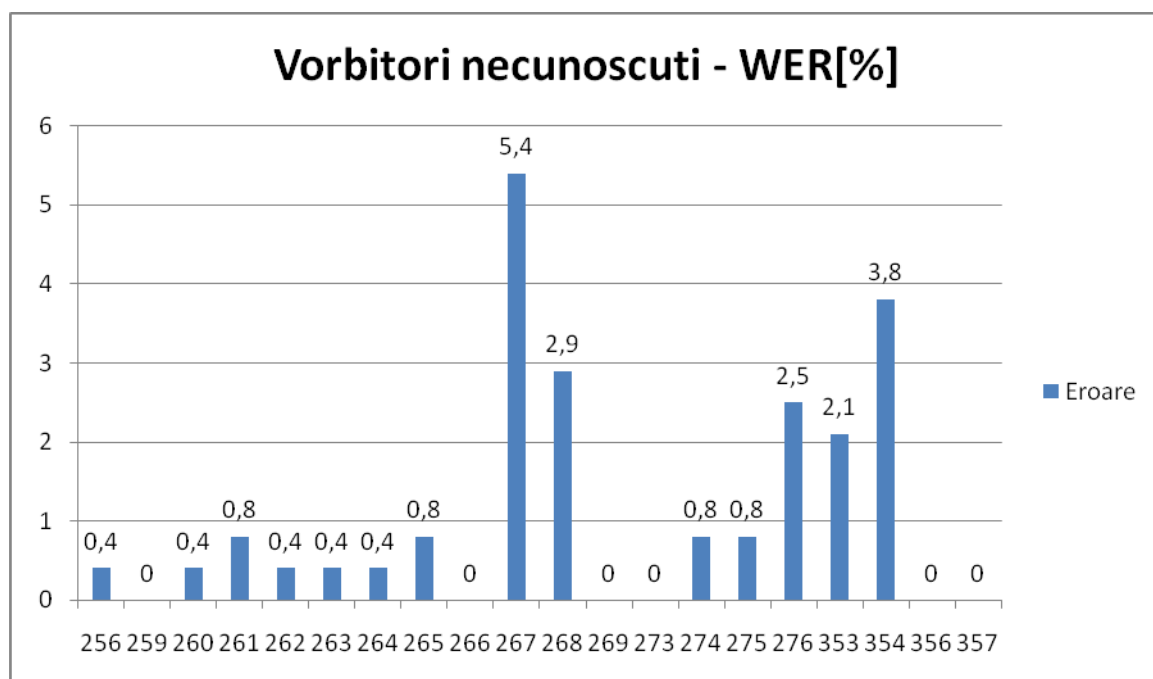


Figura 2.4 Evaluarea vorbitorilor necunoscuți (WER)

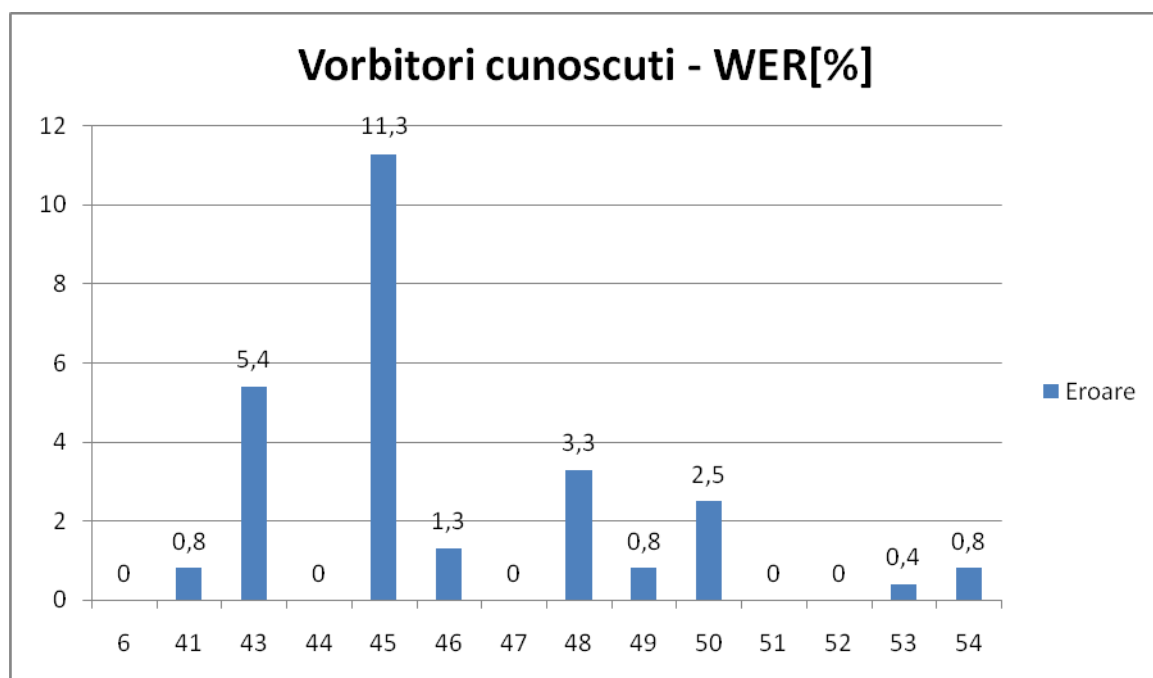


Figura 2.5 Evaluarea vorbitorilor cunoscuți 1 (WER)

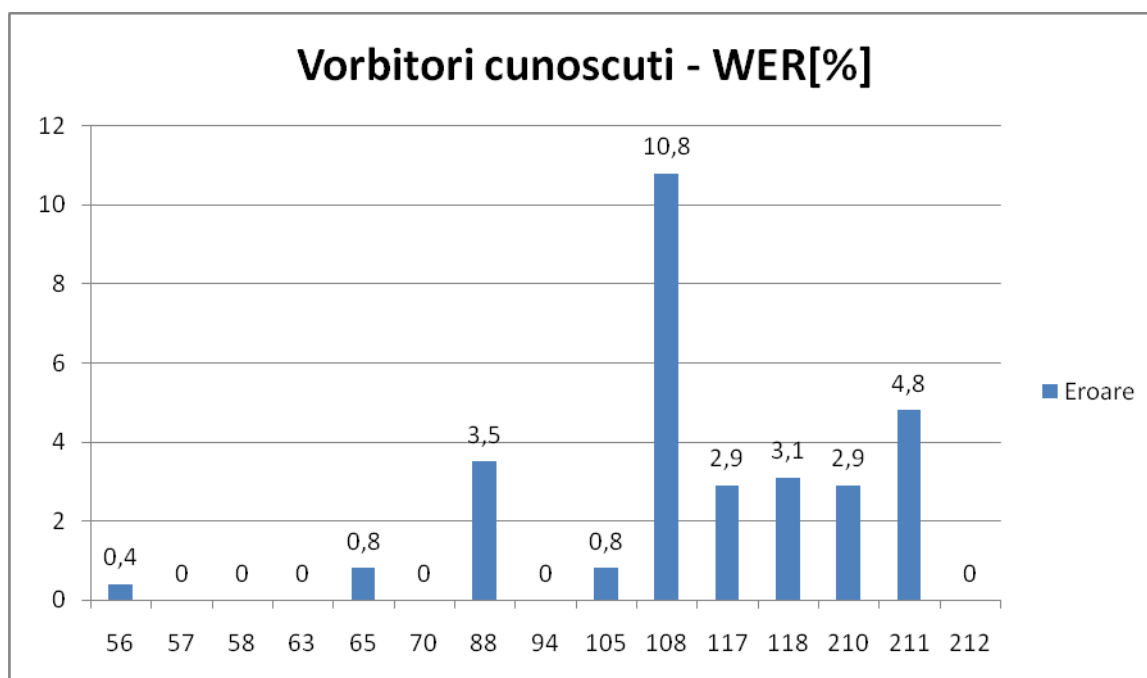


Figura 2.6 Evaluarea vorbitorilor cunoscuți 2 (WER)

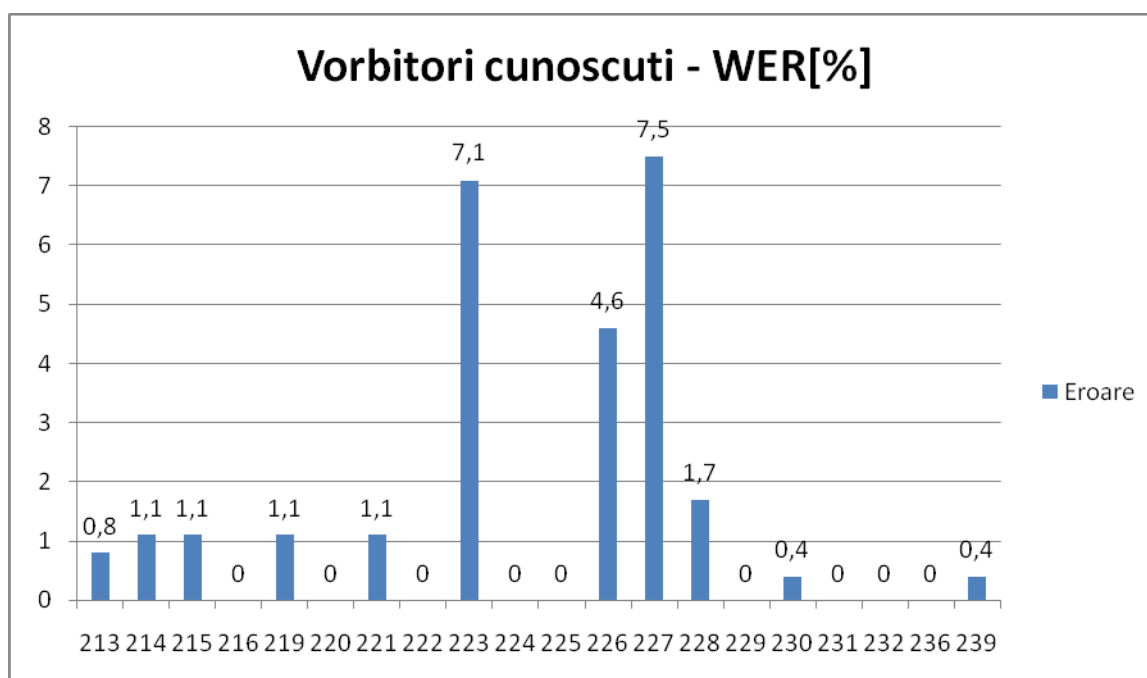


Figura 2.7 Evaluarea vorbitorilor cunoscuți 3 (WER)

La vorbitorii necunoscuți, cea mai mare rată de eroare este de 5.4% per vorbitor, în timp ce la vorbitorii cunoscuți, apar doi vorbitori rate de eroare cu valori puțin peste 10%. În acest moment, rezultatele anterioare sunt explicabile. Prin ascultarea înregistrărilor audio de la cei doi vorbitori cu valori peste 10%, s-a constatat:

Vorbitorul 45 (masculin)

- Înregistrarea 45: rostește prima cifră înainte de începerea înregistrării;
- Înregistrarea 58: rostește o singură cifră;
- Pronunția este foarte accentuată, nenaturală, iar silabele sunt alungite.

Vorbitorul 108 (feminin)

- Înregistrarea 24: lipsește prima cifră;
- Înregistrarea 92: sunt rostite primele 4 cifre din grup, restul lipsesc;
- Pronunță încet, uneori neclar, fiind foarte asemănător modul în care rostește cuvintele *șase* și *șapte*.

2.1.4 Prepararea modelelor în vederea antrenării folosind utilitarul Kaldi

Modelul fonetic

Acesta a fost realizat prin transcrierea fonetică a cuvintelor, similar cu cel creat în cadrul proiectului folosind utilitarul CMU Sphinx. Se poate preciza faptul că de data aceasta nu a fost necesar să se precizeze poziția fonemului în cuvânt.

```
<SIL> SIL
zero z e r o
unu u n u
doi d o i3
trei t r e i3
patru p a t r u
cinci k1 i n k1 i
șase s1 a s e
șapte s1 a p t e
opt o p t
nouă n o1 w a1
```

Modelul de limbă

Au fost create două modele de limbă separate:

1. Un model de limbă **probabilistic de tip n-gram** realizat cu ajutorul utilitarului SRI-LM. Acesta a primit la intrare un corpus de text (în acest caz lista cuvintelor ce desemnează cifrele în limba română) și a generat probabilitatea de apariție a fiecărui cuvânt sau grup de cuvinte:

```
\data\
ngram 1=11

\1-grams:
-1 zero
-1 unu
-1 doi
-1 trei
-1 patru
-1 cinci
-1 șase
-1 șapte
-1 opt
-1 nouă
-99 <s>
-1 </s>

\end\
```

Înainte de fiecare cuvânt se regăsește probabilitatea să de apariție, exprimată ca logaritm în baza 10. Modelul a fost compilat sub forma unui automat cu stări finite (FST) cu ajutorul utilitarului OpenFST, iar graful corespunzător poate fi vizualizat mai jos:

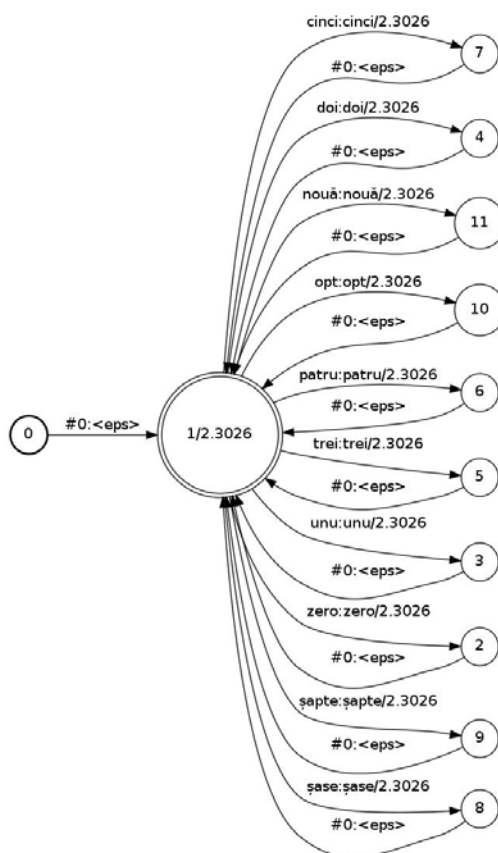


Figura 2.8 Modelul de limba 1-gram Rodigits (Kaldi)

Așa cum se poate observa din imagine, există un nod de intrare în graf, urmat de un nod al stării inițiale și 10 noduri finale, corespunzătoare fiecărui cuvânt. Rostirea unui cuvânt va determina o tranziție din nodul inițial într-un nod final. Arcele inverse fac posibilă întoarcerea în nodul inițial, fiind astfel realizabilă o nouă tranziție. Fiecare arc are asociată o probabilitate de tranziție. În cazul curent ele sunt echiprobabile.

2. Un **model bazat pe reguli de tip gramatică cu stări finite** realizat cu ajutorul utilitarului Thrax.

```
numbers= "[2]" | "[3]" | "[4]" | "[5]" | "[6]" | "[7]" | "[8]" | "[9]" | "[10]" | "[11]";
export rodigits=Optimize[ numbers+ ];
```

Kaldi nu lucrează direct cu cuvinte în format text, ci folosește o tabelă de indexare a acestora. Simbolurile "[2]"..."[11]" reprezintă indecșii cuvintelor și sunt separate prin caracterul de disjuncție logică "|", semnificând faptul că variabilă "numbers" poate lua oricare dintre aceste valori. Simbolul plus indică posibilitatea cuvintelor de a se repeta. În urma compilării a rezultat graful următor:

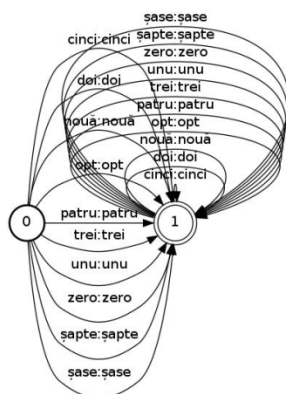


Figura 2.9 Modelul de limba de tip gramatică Rodigits (Kaldi)

O prima rostire va determina o tranziție din starea #0 în starea #1. Următoarele rostiri vor avea ca efect tranziții repetate din starea #1 în starea #1.

Modelul acustic

Au fost folosiți mai mulți algoritmi disponibili în Kaldi, fiecare generând un tip de sistem acustic diferit. Au fost păstrate valorile standard ale parametrilor acustici.

2.1.5 Experimente și rezultate. Compararea utilitatelor

Pentru efectuarea experimentelor, au fost folosiți 85 de vorbitori din baza de date "Rodigits", fiecare având un set de 100 de fișiere. Acestea au fost împărțite în felul următor:

- Setul de antrenare:
 - o 60 de vorbitori, 80 de fișiere de la fiecare;
- Setul de testare:
 - o Vorbitori cunoscuți: cei 60 folosiți la antrenare, restul de 20 de fișiere de la fiecare;
 - o Vorbitori necunoscuți: 25 de vorbitori, 20 de fișiere de la fiecare.

Experiment 1

În urma antrenării și testării folosind utilitarul Kaldi au fost obținute următoarele rezultate:

		Model de limba N-gram		Model de limbă FST	
Tipul modelului acustic	Parametri acustici	WER[%]		WER[%]	
		Vorbitori cunoscuți	Vorbitori necunoscuți	Vorbitori cunoscuți	Vorbitori necunoscuți
MONOPHONE		1.5	1.68	2.31	2.5
TRI1	LeavesTri1=100	2.21	2.05	3.46	3.12
	GaussTri1=256				
TRI2 LDA+MLLT	LeavesMLLT=2500	0.56	1.37	0.76	2.25
	GaussMLLT=15000				
TRI3 LDA+MLLT+SAT	LeavesSAT=100	1.13	0.62	1.53	0.97
	GaussSAT=256				
SGMM2	GaussUBM=128	0.43	0.28	0.7	0.73
	LeavesSGMM=100				
	GaussSGMM=128				

Tabelul 2.13 Experiment 1. Rezultate Rodigits (Kaldi)

Se observă că anumite modele acustice din Kaldi obțin o acuratețe mai bună decât celelalte, SGMM2 fiind în această situație cel mai bun. Același lucru se poate spune despre modelul de limba n-gram, în comparație cu cel de tip gramatică. Deși ambele modele de limbă oferă probabilități egale de tranziție între cuvinte, totuși rezultatele sunt diferite. În momentul actual nu se cunoaște cauza, acest lucru fiind investigat.

Experiment 2

Utilizând în mod identic setul de antrenare, precum și cele de testare descrise anterior, experimentul a fost replicat folosind utilitarul CMU Sphinx, parametrii acustici fiind cei determinați ca optimi în secțiunea 3.1.3.

		Model de limbă FST	
Tipul sistemului acustic	Parametri acustici	WER[%]	
		Vorbitori cunoscuți	Vorbitori necunoscuți
HMM-GMM	Densități gaussiene=128 Senone=100	0.8	1.5

Tabelul 2.14 Experiment 2. Rezultate Rodigits (CMU Sphinx)

Comparând rezultatele din cele două experimente, cel puțin pentru configurațiile experimentale prezentate, Kaldi pare să ofere o acuratețe mai bună în comparație cu CMU Sphinx.

2.2 CREAREA SISTEMULUI "MOLINA"

"Molina" este un sistem de recunoaștere automată a vorbirii continue cu vocabular relativ redus (215 cuvinte unice) în limba engleză. Scopul acestuia este de a realiza transcrierea unor clipuri audio ce conțin fraze care imită răspunsul unui sistem interactiv cu răspuns vocal. În continuare, se vor prezenta detalii referitoare la achiziția și structura bazei de date de evaluare, procedurile

de creare a modelelor componente (acustic, fonetic, lingvistic), dar și rezultate privind acuratețea și performanțele sistemului.

2.2.1 Baza de date audio "Molina"

Aceasta reprezintă baza de date de evaluare a sistemului ASR. Pentru crearea ei a fost selectată o listă de 75 de fraze, similare sau chiar identice cu cele ale sistemului IVR Molina [17]. Aceste fraze simulează răspunsul unui sistem IVR și au ca scop ghidarea utilizatorilor acestuia, oferindu-le informațiile dorite. Mai jos se regăsesc câteva exemple:

0001: *Welcome to the Interactive Voice Response System for Molina Healthcare of Wisconsin. Please listen carefully because our options have changed.*

0002: *Press 1 for Eligibility Information.*

0003: *Press 2 for Claim Status.*

0004: *Press 3 for an authorization Status.*

0005: *If you would like to hear these choices again, press 9.*

Este important de specificat faptul că frazele au un format specific, acestea fiind alcătuite din **părți statice** și **părți dinamice**. Frazele care diferă numai prin părțile lor dinamice, în timp ce părțile statice sunt identice, se consideră că aparțin aceluiași **tip de frază**:

Fraza 1: *If your credit is five thousand dollars, press one.*

Fraza 2: *If your credit is seventy three dollars, press two.*

Tip frază: *If your credit is NUMBER dollars, press UNITS.*

În exemplul dat, cele două fraze conțin câte două părți dinamice, marcate prin culori diferite. Atât "five thousand" cât și "seventy three" aparțin părții dinamice denumite în mod generic "NUMBER" (ce poate reprezenta orice număr), în timp ce "one" și "two" aparțin părții dinamice denumite în mod generic "UNITS" (ce poate reprezenta orice cifră). În baza de date "Molina" se disting 42 de tipuri de fraze.

La crearea înregistrărilor audio din această bază de date au participat 5 membri ai Laboratorului de Cercetare "Speed" din cadrul Facultății de Electronică, Telecomunicații și Tehnologia Informației, fiecare furnizând un set de 75 fraze. S-a folosit o aplicație web ce afișează fiecare frază și permite utilizatorului să înregistreze, să asculte și să reînregistreze în cazul în care este necesar. Așadar, baza de date conține un număr de 375 fraze și are o durată totală de 42 minute.

2.2.2 Antrenarea modelelor

Modelul acustic

A fost obținut prin antrenarea unei baze de date în limba engleză numită "TED-LIUM"[11] ce conține atât fișierele audio cât și transcrierile de la conferințele "TedTalks" [13]. Tabelul de mai jos conține caracteristicile acestei baze de date:

Caracteristică	Antrenare	Dezvoltare
Număr de conversații	774	19
Durata totală	118 h, 4 m, 48 s	4 h, 12 m, 55 s

• Vorbitori masculini	81 h, 53 m, 7s	3 h, 13 m, 57 s
• Vorbitori feminini	36 h, 11 m, 41 s	58 m, 58 s
Durata medie	9 m, 9 s	13 m, 18 s
Număr de vorbitori unici	666	19
Număr de segmente	56.6 K	2 K
Număr de cuvinte	1.7 M	47 K

Tabelul 2.15 Caracteristicile bazei de date TED-LIUM

Sistemul de recunoaștere automată a vorbirii "TED-LIUM" este unul din proiectele oferite ca exemplu de către creatorii utilitarului Kaldi. Prin descărcarea acestuia, utilizatorului îi sunt oferite toate componentele necesare pentru crearea sistemului ASR (baze de date de antrenare, de dezvoltare și de testare împreună cu transcrierile aferente, modelul fonetic și modelul de limba). De asemenea, este pus la dispoziție un fișier cu diverse comenzi asupra modulelor din utilitar, iar prin simpla rulare a lui se va realiza antrenarea sistemului, rezultând mai multe modele acustice, fiecare model fiind obținut cu ajutorul unui algoritm ce corespunde unui tip de sistem acustic diferit. La final este realizată decodarea fișierelor de test și se evaluează performanța sistemului.

Modelul de limbă corespunzător sistemului "TED-LIUM" este unul de tip probabilistic, creat pe baza transcrierilor conferințelor TedTalks și conține probabilitățile de apariție pentru:

- 1-gram: 19794 termeni;
- 2-gram: 1377200 termeni;
- 3-gram: 3178194 termeni.

Tabelul de mai jos prezintă o comparație între rezultatele oficiale (obținute de creatorii utilitarului) și rezultatele personale ale erorii la nivel de cuvânt.

Tipul sistemului acustic	Setul de date	Rezultate WER[%]	
		Rezultate oficiale	Rezultate personale
TRI1	dev	36.0	36.0
	test	34.7	34.6
TRI2	dev	32.1	31.8
	test	29.8	29.9
TRI3	dev	27.4	27.4
	test	24.6	24.9
TRI3 MMI	dev	24.0	24.3
	test	21.6	21.3

Tabelul 2.16 Rezultatele antrenării bazei de date TED-LIUM

Având în vedere că valorile obținute prin rulare programului sunt foarte apropiate de cele ale creatorilor utilitarului, acest lucru indică faptul că proiectul oferit ca exemplu a rulat corespunzător. În felul acesta au fost obținute câteva modele acustice pentru limba engleză ce vor fi folosite mai departe în sistemul Molina ASR.

Modelul fonetic

A fost creat prin transcrierea fonetică manuală a fiecărui cuvânt ce se găsește în baza de date de evaluare. Modelul poate conține mai multe transcrieri ale unui cuvânt dacă pentru acesta există mai multe tipuri de pronunție.

Mai jos este un extras din transcrierea fonetică creată:

a AH
 access AE K S EH S
 after AE F T ER
 again AH G EH N
 amount AH M AW N T
 an AE N
 and AE N D
 and AH N D
 another AH N AH DH ER
 any EH N IY
 are ÂÂ R
 ashley AE SH L IY
 asked AE S K T
 asked AO S K T

De asemenea, în compunerea modelului fonetic se regăsește o listă a tuturor fonemelor existente precum și o listă a fonemelor nonverbale.

Fonemele limbii engleze	
AA	L
AE	M
AH	N
AO	NG
AW	OW
AY	OY
B	P
CH	R
D	S
DH	SH
EH	T
ER	TH
EY	UH
F	UW
G	V
HH	W
IH	Y
IY	Z
JH	ZH
K	

Fonemele nonverbale (liniște, zgomot, sunete necunoscute, etc):

SIL
 BRH

CGH
NSN
SMK
UM
UHH

Modelul de limbă

Deoarece aceste sisteme de recunoaștere automată a vorbirii au un vocabular relativ redus, modelul de limbă utilizat este unul creat special pentru această sarcină, fiind un model bazat pe reguli, de tip gramatică cu stări finite (FSG).

Acesta a fost creat folosind utilitarul Thrax, iar mai jos este exemplificată procedura de creare.

Modelul de limbă de tip gramatică LM1-v1
<pre>f1="[214][201][197][121][211][180][195][100][144][110][155][219][166][135][64][52][163][160][109][67]"; f2="[170][2][100][88][119]"; f3="[170][20][100][72][194]"; f4="[170][27][100][41][49][194]"; f75="[115][224][76][123][35][85][170][2][115][150][170][20]"; fragment=f1 f2 f3 f4 f5 f75; export molina=Optimize[fragment];</pre>

Variabilele f1, f2, f3 ...f75 iau ca valoare o serie de simboluri numerice ce reprezintă indexul fiecărui cuvânt din vocabularul bazei de date de evaluare. Astfel, aceste variabile compun frazele bazei de date. Caracterul "|" reprezintă operatorul de disjuncție logică, acesta indică faptul că fraza decodată de sistem poate fi oricare dintre frazele f1, f2, ...f75.

Acest model de limbă a fost compilat sub forma unui automat de stări finite (FST) folosind următoarele comenzi specifice Thrax:

```
thraxmakedep molina.grm
make
farextract Molina.far
```

Automatul cu stări finite rezultat poate fi vizualizat sub forma unui graf. Deoarece acesta este destul de complex, în imaginea de mai jos se regăsește un fragment ce reprezintă primele 5 fraze:

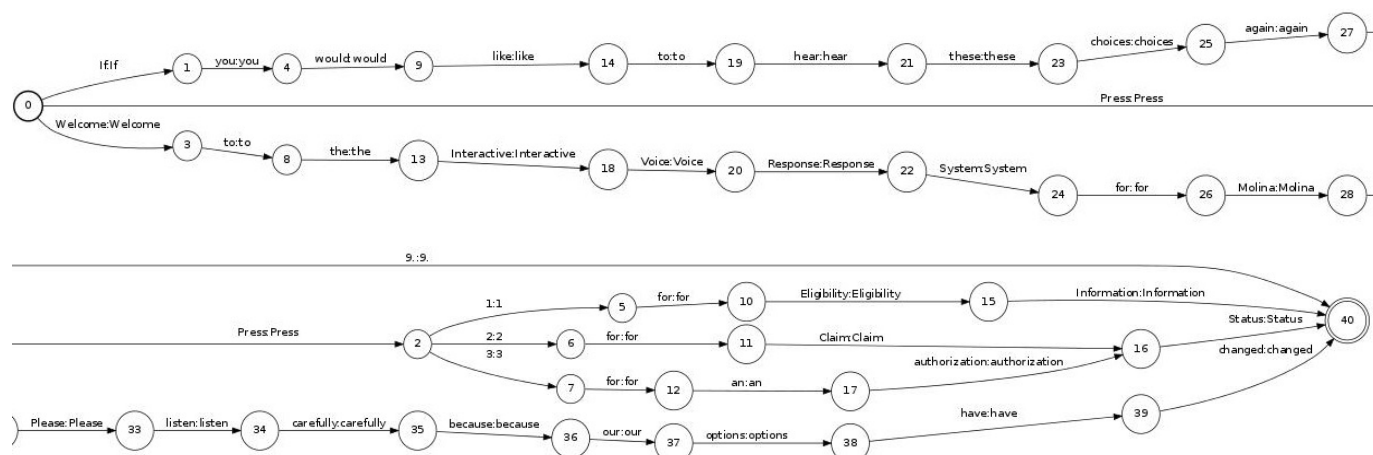


Figura 2.10 Graful corespunzător modelului de limbă de tip gramatică Molina

Optimizarea modelului de limbă precedent

Modelul de limbă creat poate fi considerat ca fiind unul "static", unde fiecare frază este definită prin șirul de cuvinte componente. Mai mult, fiecare cuvânt este considerat de sine stătător, fără a se lua în calcul faptul că unele cuvinte sau grupuri de cuvinte pot fi obținute prin concatenarea altor unități mai mici. Ținând cont de acest fapt, dar și de formatul specific al frazelor, s-a realizat trecerea la o nouă versiune a modelului de limbă bazat de reguli, un model dinamic, unde numerele, datele calendaristice, numele proprii să fie obținute prin concatenarea cuvintelor componente. Mai precis, au fost create subgramatici pentru părțile dinamice ale frazelor, iar mai departe, frazele au fost definite prin concatenarea părților dinamice cu părțile statice. În tabelul următor sunt prezentate subgramaticile obținute: O astfel de subgramatica pentru cazul numerelor formate din două cifre este prezentată mai jos:

```
units="[217]" "[137]" "[195]" "[190]" "[75]" "[69]" "[172]" "[166]" "[53]" "[125]";
tens="[194]" "[186]" "[74]" "[67]" "[174]" "[168]" "[55]" "[127]";
twoDigits=tens units;
```

Variabilă "units" poate fi orice cuvânt reprezentat de cifrele din intervalul [0;9], în timp ce variabilă "tens" poate fi orice cuvânt reprezentat de numerele din mulțimea [20; 30; 40; 50; 60; 70; 80; 90] (cuvintele sunt reprezentate prin simbolul echivalent din tabela de indexare a cuvintelor specifică utilitarului Kaldi). Variabilă "twoDigits" rezultă prin concatenarea celorlalte două variabile și va corespunde tuturor numerelor cuprinse în intervalul [20;99]. În mod similar au fost realizate subgramatici pentru toate părțile dinamice existente. Tabelul următor reprezintă un extras din noua versiune de gramatică:

Modelul de limbă de tip gramatică LM1-v2
<pre>f2 = "[104]" UNITS "[44]" "[38]" "[63]" ; f3 = "[104]" UNITS "[44]" "[20]" "[122]" ; f4 = "[104]" UNITS "[44]" "[33]" "[66]" "[122]" ; f19 = "[88]" FIRST_NAME LAST_NAME "[124]" DATE "[63]" "[30]" "[38]" "[44]" "[103]" "[104]" "[103]" "[93]" PLAN_NAME ; f42 = "[62]" "[141]" "[293]" "[63]" NUMBER "[34]" "[104]" UNITS "[62]" "[96]" "[104]" UNITS ;</pre>

```

phrases = f1 | f2 | f3 | f4 | f5 | f6 | f7 | f8 | f9 | f10 | f11 | f12 | f13 | f14 | f15 | f16 | f17 | f18 | f19 |
f20 | f21 | f22 | f23 | f24 | f25 | f26 | f27 | f28 | f29 | f30 | f31 | f32 | f33 | f34 | f35 | f35 | f36 | f37 |
f38 | f39 | f40 | f41 | f42 ;
export molina_lm = Optimize [ phrases+ ] ;

```

Concluzie: Așa cum se va observa în secțiunea următoare, modelul de limba LM1-v1 va oferi o acuratețe mai bună decât modelul de limba LM1-v2. Acest lucru este normal, deoarece primul model impune o foarte mare strictețe din punct de vedere al îmbinării cuvintelor în fraze, în timp ce al doilea permite o variabilitate mai mare. Totuși, avantajul celui de-al doilea model este conferit de posibilitatea de a modifica sau de a adăuga o nouă subgramatica într-un mod mult mai facil.

2.2.3 Experimente și rezultate

Experiment 1

Folosind modelele explicitate mai sus și baza de date de evaluare Molina s-a obținut următoarele rezultate de evaluare ale sistemului:

	Modelul de limba de tip gramatică LM1-v1	
Tipul sistemului acustic	WER[%]	SER[%]
MONOPHONE	1.56	3.73
TRI1	2.75	4.8
TRI2 LDA+MLLT	1.81	3.2
TRI3 LDA+MLLT+SAT	1.1	2.67
TRI3 MMI	0.84	2.13

Tabelul 2.17 Experiment 1. Rezultate Molina (LM1-v1)

Se poate observa că cele mai bune rezultate s-au obținut pentru modelul TRI3 MMI (triphone maximum mutual information) ce urmărește să maximizeze probabilitatea aposteriori a rostirii corecte.

În concluzie, frazele din baza de date de evaluare Molina au fost decodate cu o eroare la nivel de cuvânt de 0,84%, în timp ce eroarea la nivel de frază fiind de 12-13%.

Experiment 2

Au fost realizate teste de performanță ale sistemului de calcul în timpul decodării bazei de date Molina folosind sistemul ASR Molina creat. Testele au fost create folosind două modele de limbă diferite. Primul este un model de limbă bazat pe reguli atât pentru fraze cât și pentru părțile dinamice din acestea (a se vedea secțiunea de mai sus). Iar al doilea este un model de limba 3-gram pentru limba engleză împrumutat din proiectul exemplu Tedlium. Pentru aceste două modele s-a variat numărul de procese și numărul de fire de execuție, au fost create unele

proceduri pentru măsurarea consumului de memorie, a timpului de decodare, valoarea media de solicitare a CPU, factorul de timp real, accelerarea și consumul de energie.

Pentru calculul valorii medii a solicitării CPU și a memoriei a fost creat un script care prin realizarea unor apeluri la comandă "top", din sistemul de operare Linux, la fiecare secundă, pe toată durata procesului de decodare, oferă gradul de utilizare al memoriei și al CPU. Aceste valori au fost însumate și împărțite la durata totală a decodării.

Factorul de timp real (RTF) este un raport între durata decodării și durata secvenței audio decodate. Accelerarea arată de câte ori este mai rapidă sarcină de decodare atunci când sunt folosite mai multe fire de execuție în comparație cu cazul în care este folosit unul singur.

Consumul de energie este calculat folosind utilitarul PowerAPI. În timpul decodării au fost monitorizate toate procesele corespunzătoare acestei sarcini și au fost extrase eșantioane la interval de o secundă ce indică puterea instantanee consumată. A fost calculat consumul mediu de energie precum și valorile instantanee maxime atinse.

Observație: La începutul sarcinii de decodare în timp ce decodorul din Kaldi este inițializat (sunt încărcate modelele acustic, fonetic și de limbă) consumul de energie crește gradual. Apoi în timpul decodării propriu-zise are un caracter constant, iar când decodarea se termină și latticele obținute sunt evaluate, consumul de energie începe să scadă.

Sumar experiment:

Baza de date folosită: *Molina*

Durata totală a bazei de date: *2557 secunde, aproximativ 42 minute*

Model acustic folosit: *TRI3 MMI antrenat pe bază de date TED-LIUM*

Model de limbă folosit: *LM1-v2 (model bazat pe reguli special creat pentru baza de date Molina) și LM2 (3-gram împrumutat din proiectul exemplu TED-LIUM)*

WER[%] folosind LM1-v2: *1.76*

WER[%] folosind LM2: *60.99*

LM1-v2	Număr procese	Număr fire de execuție	Memorie [MB]	Timp decodare [s]	CPU [%]	RTF	Accelerare	Media consumului de putere [W]	Punct maxim de consum [W]
Test nr.									
1	1	1	48	171	4.43	0.067		2.73	2.8
2	1	4	50	48	16.1	0.018	3.56	10.34	11.3
3	1	8	81	28	29.13	0.0109	6.10	18.67	22.5
4	1	16	98	21	51	0.0082	8.14	31.69	45
5	1	24	116	21	67	0.0082	8.14	41.86	62
6	1	32	123	22	61	0.0086	7.77	40.4	61
7	2	1	89	96	8.39	0.0375		5.04	5.5
8	2	4	90	30	28.85	0.0117	3.20	18.16	22.5
9	2	8	148	22	51.3	0.0086	4.36	32.71	45
10	2	16	190	21	70.72	0.0082	4.57	46.87	64
11	2	24	227	21	71	0.0082	4.57	44.5	64.5
12	2	32	258	21	72	0.0082	4.57	45.88	65

Tabelul 2.18 Evaluarea consumului de resurse (LM1-v2)

LM2	Număr procese	Număr fire de execuție	Memorie [GB]	Timp decodare [s]	CPU [%]	RTF	Accelerare	Media consumului de putere [W]	Punct maxim de consum [W]
Test nr.									
1	1	1	4.8	3778	4.37	1.477		2.6	2.8
2	1	4	4.86	1095	16.4	0.428	3.45	10.8	13.8
3	1	8	4.95	644	31.49	0.251	5.86	20.7	25
4	1	16	5.37	466	58	0.182	8.10	38.4	47.2
5	1	24	5.7	441	76	0.172	8.57	51	66
6	1	32	6	447	80	0.174	8.45	53	66
7	2	1	9.54	1919	8.53	0.75		5.45	6.63
8	2	4	9.42	648	31.43	0.253	2.96	20.6	26
9	2	8	9.59	478	59.24	0.186	4.01	40	49
10	2	16	10.47	441	89.75	0.172	4.35	56.17	65
11	2	24	11.22	448	84	0.175	4.28	56.15	66
12	2	32	11.94	455	86	0.177	4.21	57.11	65

Tabelul 2.19 Evaluarea consumului de resurse (LM2)

Concluzie: Se observă că la un anumit moment, chiar dacă numărul de fire de execuție crește, timpul de decodare nu este îmbunătățit. Acest lucru este probabil datorat faptului că peste un anumit număr de fire de execuție creat acestea nu mai sunt necesare. Kaldi folosește un număr optim de fire de execuție în funcție și de volumul datelor ce urmează a fi procesate.

Modelul de limbă bazat pe reguli (LM1-v2) fiind creat special pentru frazele din baza de date Molina a dus la un consum de resurse mult mai bun decât la modelul de limba 3-gram (LM2).

De asemenea, se văd mari diferențe în rata de eroare la nivel de cuvânt. Un model de limbă bazat pe reguli este mult mai indicat dacă se cunoaște formatul frazelor și este relativ ușor de creat atunci când se cunoaște transcrierea bazei de date care se testează.

CAPITOLUL 3

CONSTRUCȚIA SISTEMULUI DE DETECȚIE A CUVINTELOR CHEIE

Capitolul curent are ca scop construcția unui sistem de detecție a cuvintelor cheie (KWS) folosind modulul special creat pentru această sarcină din utilitarul Kaldi. Acesta folosește metoda căutării prin lattice, descrisă în secțiunea 2.1, iar așa cum am menționat în respectiva secțiune este necesar mai întâi existența unui sistem ASR care să genereze aceste lattice. Sistemul KWS astfel creat trebuie să fie capabil să detecteze cuvinte cheie în limba engleză care se pot regăsi în frazele bazei de date Molina.

3.1 ANALIZA ASUPRA MODELELOR DE LIMBĂ FOLOSITE PENTRU DETECȚIA CUVINTELOR CHEIE

3.1.1 *Modelul de limba 3-gram*

O abordare a sarcinii de KWS este de a crea un sistem cu două componente:

- un modul ASR care decodează materialul audio generând lattice;

- un modul care caută cuvintele cheie în latticele indexate sau în mod direct prin transcrierea text.

În acest caz, importanța sistemul ASR este foarte ridicată. O transcriere cât mai precisă a semnalului audio va asigura o bună performanță a întregului sistem KWS. Deoarece s-a dispus de un model de limba 3-gram și un model acustic, ambele dezvoltate pentru un sistem de recunoaștere în limba română a vorbirii continue cu vocabular extins (RVC-VE) s-a decis realizarea unui prim experiment folosind aceste modele. Scopul a fost observarea modului în care Kaldi se comportă atunci când este folosit un model 3-gram de o dimensiune considerabilă.

1-gram - 53700 termeni;

2-gram - 16427743 termeni;

3-gram - 23674351 termeni.

Pentru antrenarea modelului acustic s-a folosit o bază de date de vorbire citită în limba română și o bază de date de vorbire înregistrată din emisiuni TV.

Pentru testarea acestui sistem ASR s-a folosit o bază de date mixtă, unele fișiere provenind din baza de date de vorbire în limba română iar altele din baza de date de vorbire înregistrată din emisiuni TV. Fișierele folosite pentru testare nu au fost cele utilizate la antrenare. În tabelul de mai jos sunt prezentate datele în urma evaluării sistemului.

WER[%] Tipul sistemului acustic	Omega test	News all	News clean	News noise
TRI1	12.37	29.77	23.04	34.97
TRI2 LDA+MLLT	11.34	28.94	22.39	33.74
TRI3 LDA+MLLT+SAT	9.75	27.56	20.47	33.07
TRI3 MMI	9.01	26.49	19.41	32.3

Tabelul 3.1 Rezultatele evaluării sistemului în limba română

"News all" reprezintă toate fișierele corespunzătoare emisiunilor de știri și constă în două părți:

- "news clean" (fișiere cu vorbire curată fără zgomot de fond)
- "news noise" (fișiere cu vorbire cu zgomote de fundal: râsete, muzică, etc.)

Se observă că performanța este mai bună cu cât tipul modelului acustic este mai avansat. Rezultatele fișierelor fără zgomot sunt mai bune decât cele cu zgomot.

Mai departe sunt prezentate performanțele sistemului ASR în limba română cu model de limba 3-gram.

Test nr.	Număr procese	Număr fire de execuție	Memorie	Timp de decodare	RTF	Comentariu
1	1	1	24.3 GB	72 min.	0.218	
2	1	4	24.3 GB	24 min.	0.072	
3	1	24	24.3 GB	12 min.	0.036	
4	2	1	47.04 GB	--	--	Crește memoria swap

Tabelul 3.2 Consumul de memorie în cazul modelului de limba 3-gram

S-a încercat varierea numărului de procese cât și al numărului de fire de execuție și s-au notat valorile corespunzătoare consumului de memorie, al timpului de decodare și al factorului de timp real. S-a putut observa că resursele sunt partajate între firele de execuție ale aceluiași proces dar sunt multiplicare în cazul existenței mai multor procese. Acest lucru se poate observa în tabel unde consumul de memorie este identic pentru primele 3 teste iar pentru ultimul test fiind aproape dublu. În oricare din cazuri consumul de memorie este mult prea mare pentru folosirea modelului de limba 3-gram.

Concluzie: Sarcină de detecție a cuvintelor cheie nu poate fi realizată printr-o metodă care să presupună obținerea transcrierii complete. O transcriere completă folosind modelul de limba 3-gram nu este deloc potrivită din punct de vedere al consumului de memorie. Mai departe s-a dorit găsirea unei alternative mai eficiente.

3.1.2 Alternativă la modelul de limba 3-gram

Având în vedere concluziile din subcapitolul precedent s-a căutat o metodă prin care transcrierea să conțină numai cuvintele cheie căutate. Astfel au fost create câteva modele de limbă noi care conțin doar aceste cuvinte. S-au realizat următoarele teste din tabelul de mai jos:

Baza de date folosită pentru testare: Omega test

Durata totală a bazei de date: 329 minute

Model acustic folosit pentru decodare: Tri2 (antrenat pe limba română)

Număr de procese: 1

Număr de fire de execuție: 24

Test nr.	Model de limbă	Număr de cuvinte	Memorie	Timp de decodare	RTF	Comentariu
1	gramatica: disjuncție logică	10	144 MB	2 min.	0.006	
2	gramatica: disjuncție logică	1000	672 MB	3 min.	0.009	
3	gramatica: disjuncție logică +umplutură	12+ filler	6.24 GB	5 min.	0.38	Au fost decodate mai puține fișiere (13 min.)
4	1-gram	53700	768 MB	10 min.	0.03	

Tabelul 3.3 Consumul de memorie în cazul modelului de limbă alternativ

Pentru primele două teste s-a folosit modele de limbă diferite bazate pe reguli ce conțin 10, respectiv 1000 de cuvinte cheie.

Cuvinte = cuvânt 1 | cuvânt 2 | cuvânt 3 | cuvânt 4 | | cuvânt 10;

Modelul de limbă a fost creat prin înșiruirea cuvintelor cheie utilizând operatorul de disjuncție logică între acestea.

Desigur că transcrierea nu este relevantă în acest caz și nu poate fi calculată o rată de eroare la nivel de cuvânt deoarece toate cuvintele ce se găsesc în baza de date de testare vor fi transcrise printr-unul din cuvintele din modelul de limbă. Nici evaluarea de detecție a cuvintelor cheie nu ar avea sens deoarece probabilitatea erorilor de alarmă falsă va fi foarte mare în timp ce probabilitatea de neidentificare va fi nulă. Singurul scop al acestor teste a fost pentru observarea consumului de memorie care a fost mult mai scăzut decât în cazul modelelor 3-gram.

Cel de-al treilea test este similar cu primele două, diferența constând în folosirea unui cuvânt de "umplutură" pe lângă cuvintele cheie căutate.

În ultimul test s-a folosit un model de limbă n-gram, mai exact s-au păstrat doar termenii 1-gram din modelul 3-gram în limba română folosit anterior. Acest test a avut ca scop compararea consumului de memorie în cazul folosirii modelului 1-gram față de 3-gram. Consumul de memorie scade de aproximativ 30 ori (de la 24,3GB la 768MB), dar în schimb eroare la nivel de cuvânt crește de aproximativ 3 ori (32% față de 11,34%).

Concluzie: se poate spune că modelul de limbă cu "filler" utilizat în cadrul celui de-al treilea test poate fi o alternativă viabilă în cadrul unui sistem KWS. Chiar dacă consumul de memorie este mai mic decât în cazul modelului 3-gram este destul de mare în comparație cu primele două teste din tabel. Acest fapt poate fi explicat prin modul în care modelul de limbă cu "filler" este tratat din punct de vedere fonetic. Cuvântul de "umplutură" apare în modelul fonetic ca fiind transcris prin fiecare fonem în parte. Considerând că un fonem are 3 stări, atunci cuvântul de "umplutură" va fi expandat ajungând la o dimensiune egală cu numărul de foneme la puterea 3. Deoarece limba română are 36 de foneme vor exista $36^3=46656$ versiuni ale cuvântului de "umplutură".

3.2 CREAREA REFERINȚEI

Pentru a evalua acuratețea unui sistem de detecție a cuvintelor cheie este necesar crearea unui model de referință ce indică momentul de timp al apariției pentru fiecare cuvânt cheie din baza de date de evaluare. Această operație s-a realizat manual. A fost aleasă o listă de 20 termeni ce se găsesc în baza de date Molina, au fost ascultate înregistrările audio și pentru fiecare fișier s-a notat momentul de timp și durata apariției fiecărui cuvânt cheie. În acest fel a fost etichetată baza de date și a fost creat un fisier de referință respectând standardul RTTM. Mai jos este prezentat un extract din acest fișier:

```
SPEAKER 006_42_0001 1 0.000 8.6 <NA> <NA> SELF <NA>
LEXEME 006_42_0001 1 1.058 0.824 word-0002 lex SELF <NA>
LEXEME 006_42_0001 1 5.409 0.397 word-0011 lex SELF <NA>
SPEAKER 006_42_0002 1 0.000 3.35 <NA> <NA> SELF <NA>
LEXEME 006_42_0002 1 0.000 3.35 NO_KEYWORD lex SELF <NA>
```

Într-un rând din acest fișier se găsesc următoarele câmpuri: tipul obiect (speaker, lexeme, segment, non-lex și non-speech), numele fișierului audio, numărul de canale, momentul de început, durata, câmp ortografic, numele vorbitorului și scorul de încredere. În exemplul de mai sus în fișierul 006_42_0001 se regăsește conținut audio de la un vorbitor. Acesta a fost înregistrat pe un singur canal, timpul de început fiind 0.000 și are o durată totală de 8.6 secunde. Datele care nu sunt disponibile sunt marcate cu <NA>. Următoarele două rânduri fac de asemenea referire la fișierul 006_42_0001 acesta indică faptul că fișierul audio conține cuvinte cheie și anume word-0002 și word-0011 (acestea fiind ID-uri asociate cuvintelor cheie reale). Momentul de început este la secunda 1.058 respectiv 5.409 și durează 0.824 secunde respectiv 0.397 secunde. Acest fișier RTTM trebuie să conțină informații despre toate fișierele audio din baza de date fie că acestea conțin sau nu cuvinte cheie. Scopul acestui fișier de referință este de a oferi informații precise cu privire la apariția cuvintelor cheie în baza de date. Prin compararea acestui fișier cu fișierul rezultat la ieșirea sistemului KWS este posibilă evaluarea acurateții de detecție a cuvintelor cheie.

Un exemplu de fișier rezultat la ieșirea sistemului KWS este prezentat mai jos:

```
<detected_termlist termid="word-0008" term_search_time="1" oov_count="0">
<term file="006_42_0044" channel="1" tbeg="1.71" dur="1.24" score="1" decision="YES"/>
<term file="266_42_0044" channel="1" tbeg="1.62" dur="0.86" score="1" decision="YES"/>
<term file="356_42_0044" channel="1" tbeg="1.25" dur="0.83" score="1" decision="YES"/>
<term file="300_42_0044" channel="1" tbeg="1.43" dur="0.86" score="0.998178"
decision="YES"/>
<term file="324_42_0044" channel="1" tbeg="1.88" dur="1.29" score="0.627089"
decision="YES"/>
</detected_termlist>
```

Pentru fiecare cuvânt cheie căutat este specificat fișierul audio în care acesta apare împreună cu momentul de timp al apariției, durata și probabilitatea.

3.3 EXPERIMENTE ȘI REZULTATE

Acest subcapitol prezintă evaluarea sistemului KWS având la bază sistemul ASR Molina a cărei construcție a fost prezentată în secțiunea 3.2. Au fost alese 20 de cuvinte cheie ce se doresc a fi

găsite. Pentru această sarcină a fost folosit modulul KWS al utilitarului Kaldi (al cărui mod de funcționare a fost prezentat în secțiunea 2.1). În urma procesului de KWS a rezultat o listă în formatul celei prezentate în secțiunea 4.2. Această listă împreună cu fișierul de referință a fost modificată astfel încât să aibă un format specific corespunzător utilitarului folosit pentru evaluarea cuvintelor cheie în cadrul MediaEval. Prin folosirea acestui utilitar pus la dispoziție de către membrii Laboratorului de Cercetare Speed s-a obținut o curbă DET precum și probabilitățile de neidentificare și de alarmă falsă (concepte prezentate în secțiunea 2.3)

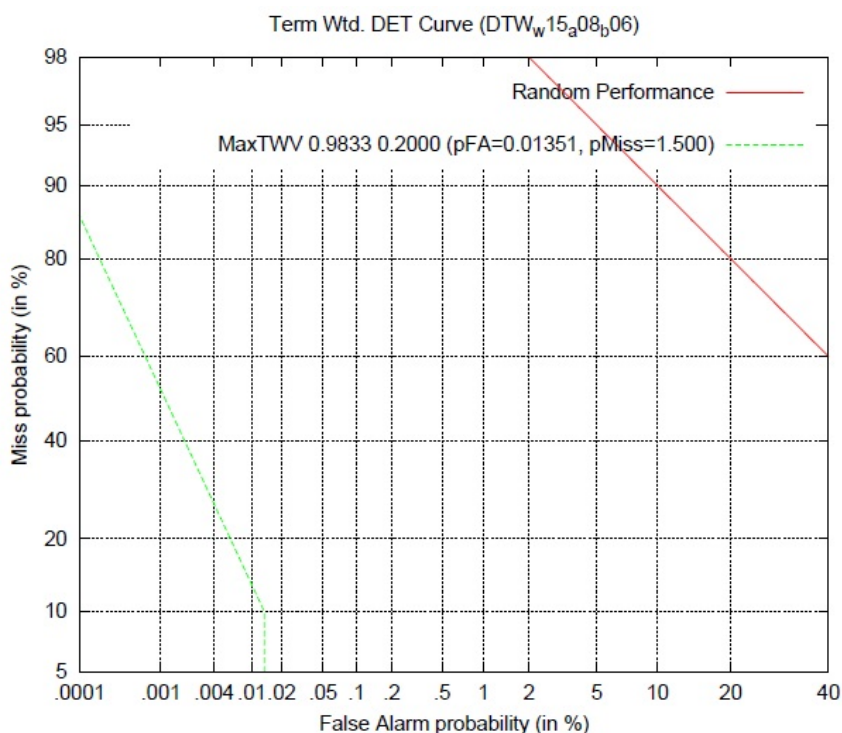


Figura 3.1 Evaluarea sistemului de detecție a cuvintelor cheie

Se poate observa compromisul dintre probabilitatea de neidentificare (1.5%) și cea de alarmă falsă (0.01351 %). În funcție de scopul sistemului KWS obținut se poate selecta un prag care prin nivelul valorii impuse conduce la următoarele două posibilități:

- un prag de nivel scăzut va avea drept rezultat o probabilitate mare de alarme false și una mică pentru cele neidentificate;
- un prag de nivel ridicat va conduce spre o probabilitate mică de alarme false și una mare pentru cele neidentificate.

S-a dorit extinderea testelor pentru evaluarea sistemului KWS prin căutarea a mai mult de 20 cuvinte cheie. Pentru aceasta a fost extins model de limbă bazat pe reguli (LM1), fiind adăugate foarte multe nume proprii ce nu se găsesc în baza de date Molina (aproximativ 700 nume și 500 prenume). În acest fel a fost creată oportunitatea ca sistemul să introducă noi erori de transcriere. Au fost pregătite 3 scenarii de test:

- s-au căutat câte 10 cuvinte cheie (2 dintre acestea existând în transcrieri) iar 8 nu existau, dar ele pot apărea ca erori de transcriere din cauza noului model de limbă;
- s-au căutat câte 100 cuvinte cheie (20 dintre acestea existând în transcrieri) iar 80 nu existau;
- s-au căutat câte 1000 cuvinte cheie (20 dintre acestea existând în transcrieri) iar 980 nu existau;

Din păcate, din cauza faptului că utilitarul de evaluare nu permite căutarea unor cuvinte ce nu apar în referință, aceste teste nu au putut fi efectuate. S-a încercat folosirea unui alt utilitar de evaluare, F4DE, dar și acesta prezintă aceeași problemă. În prezent, problema este investigată.

CAPITOLUL 4

SERVICIUL WEB DE AUTENTIFICARE PRIN VOCE

Serviciul web de autentificare prin voce este o aplicație ce oferă utilizatorilor posibilitatea de înregistrare și autentificare pe bază de parolă și recunoaștere de vorbire și amprentă vocală.

4.1 DESCRIERE ȘI ARHITECTURĂ

Proiectul este compus dintr-un ansamblu de module componente: trei module principale expuse sub forma unor servicii REST (Representational State Transfer), fiecare realizând o sarcină precisă, o aplicație client și o bază de date. Comunicația dintre acestea este realizată prin intermediul mesajelor de tip "cerere/răspuns" specifice protocolului HTTP (Hyper Text Transfer Protocol). Datorită acestui fapt, sistemul poate lucra în mod distribuit.

Componentele aplicației sunt:

- Serviciul de autentificare;
- Serviciul de recunoaștere de vorbitor (nu este parte a acestei lucrări);

- Serviciul de recunoaștere de vorbire (nu este parte a acestei lucrări);
- Aplicația client;
- Baza de date.

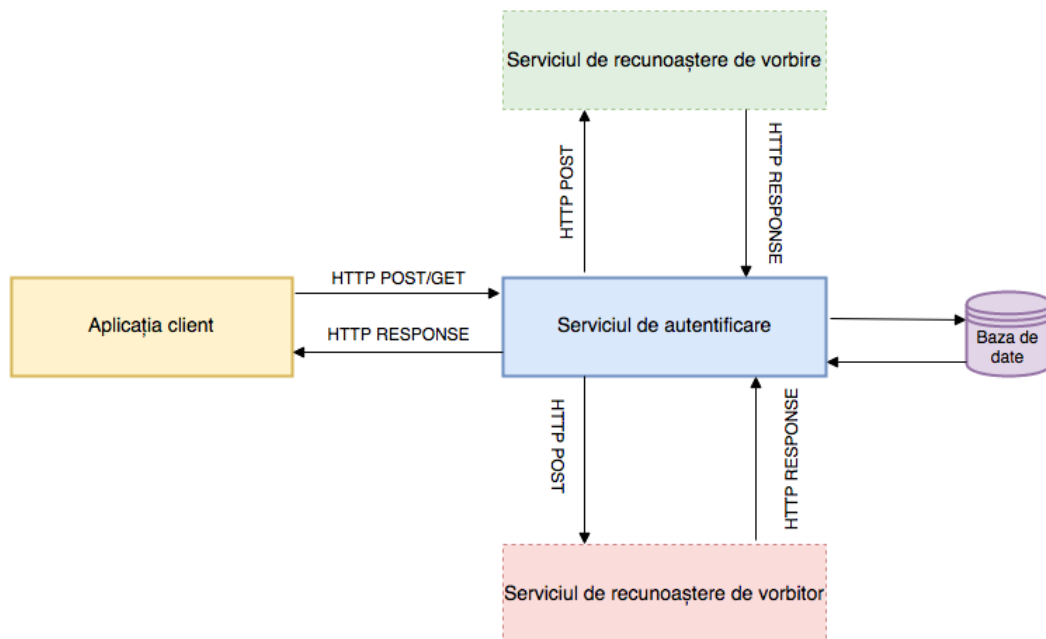


Figura 4.1 Arhitectura serviciului web de autentificare prin voce

Aplicația client

În etapa de înregistrare, utilizatorul trebuie să completeze un formular web cu datele sale personale. Tot în această etapă, trebuiesc create cel puțin 10 înregistrări audio a câte 12 cifre fiecare, șirurile de cifre fiind generate în mod aleator și afișate în aplicația client. Atât datele personale cât și înregistrările audio vor fi trimise către serviciul de autentificare. Clientul va primi un mesaj informativ cu privire la starea contului (dacă a fost creat sau nu).

În etapa de autentificare, utilizatorul va completa un formular cu adresa de poștă electronică și parola și va înregistra o singură secvență audio. Aceste date vor fi de asemenea trimise serviciului de autentificare. Dacă autentificarea a fost realizată cu succes, aplicația client va fi redirectionată spre pagina personală a utilizatorului, în caz contrar primind un mesaj informativ cu privire la eroarea întâmpinată.

Serviciul de autentificare

Este elementul central al aplicației, deoarece realizează agregarea și procesarea datelor, interacționează cu aplicația client, cu baza de date și cu serviciile de recunoaștere de vorbitor și vorbire.

În etapa de înregistrare, datele personale și înregistrările audio primite de la aplicația client vor fi inserate în baza de date. Apoi, aceleași înregistrări audio vor fi trimise către serviciul de recunoaștere de vorbitor, respectiv serviciul de recunoaștere de vorbire. Primul va realiza un model specific vorbitorului, utilizat ulterior la autentificare. Cel de-al doilea va returna transcrierile înregistrărilor audio. Pe baza acestora și a transcrierilor de referiță (cele ce se presupun că au fost citite de către utilizator, afișate în aplicația client la înregistrare), se va calcula rata de eroare la nivel de cuvânt pentru fiecare înregistrare, precum și media erorilor. Baza de date va fi actualizată cu aceste valori. La final, aplicației client îi va fi returnat un mesaj privitor la acțiunea de înregistrare.

În etapa de autentificare, se vor primi de la aplicația client: numele de utilizator (reprezentat de adresa de poștă electronică), parola și o înregistrare audio. Indiferent dacă autentificarea va avea

loc sau nu, tentativa de autentificare va fi stocată în baza de date. Înregistrarea audio va fi trimisă către serviciul de recunoaștere de vorbitor, respectiv serviciul de recunoaștere de vorbire. Se vor realiza 3 verificări:

1. Parola să fie identică cu cea oferită la înregistrare;
2. Serviciul de recunoaștere de vorbitor trebuie să confirme că vocea din înregistrare corespunde modelului utilizatorului pretins;
3. Serviciul de recunoaștere de vorbire va returna transcrierea înregistrării și se va calcula rata de eroare la nivel de cuvânt. Aceasta nu trebuie să depășească un anumit prag sau cel mult să se regăsească în jurul valorii de eroare medie, specifică utilizatorului și calculată la înregistrare.

Serviciul de recunoaștere de vorbitor

În etapa de înregistrare, acesta primește de la serviciul de autentificare setul de fișiere audio ale utilizatorului, pe baza cărora va crea un model specific acestuia, amprenta sa vocală.

În etapa de autentificare, primește de la serviciul de autentificare fișierul audio al utilizatorului, iar prin comparație cu amprenta vocală deja existentă, va decide dacă utilizatorul este cel pretins sau un impostor.

Implementarea acestui serviciu nu este parte a lucrării de față, mai multe detalii pot fi consultate în lucrarea [14].

Serviciul de recunoaștere de vorbire

Atât în etapa de înregistrare, cât și în cea de autentificare, acesta primește de la serviciul de autentificare înregistrările audio ale utilizatorului. Având implementat un sistem de recunoaștere automată a vorbirii, specializat în recunoașterea cifrelor în limba română, acesta întoarce transcrierile fișierelor audio. Rolul acestui serviciu este de a verifica dacă utilizatorul a rostit într-adevăr șirurile de cifre afișate în aplicația client.

Implementarea acestui serviciu nu este parte a lucrării de față, mai multe detalii pot fi consultate în lucrarea [14].

Baza de date

Toate operațiile asupra acestora sunt efectuate de către serviciul de autentificare. Baza de date stochează datele personale ale utilizatorilor înregistrați, fișierele audio împreună cu transcrierile acestora și ratele de eroare calculate. De asemenea, stochează informații referitoare la toate tentativele de autentificare, reușite sau nu.

4.2 FUNCȚIONALITATE DIN PERSPECTIVA UTILIZATORULUI

Etapa de înregistrare

În această etapă, așa cum a fost prezentat și anterior, utilizatorul trebuie să completeze un formular cu datele personale: nume, prenume, adresa de poștă electronică, parola și o fotografie. Seturi a câte 12 cifre sunt afișate în pagină. Pentru înregistrarea audio a unui set, utilizatorul va trebui mai întâi să accepte accesul aplicației la microfon. Apoi, prin apăsarea butonului *Start*, va începe captarea audio, iar utilizatorul va rosti setul de cifre. Pentru încheierea captării audio, se va apăsa din nou butonul *Start* (devenit acum *Stop*, fiind un buton ce poate lua două valori, în funcție de starea sistemului de înregistrare audio: pornit sau oprit). În acest moment, utilizatorul poate asculta înregistrarea audio, iar în cazul în care constată că au intervenit erori de pronunție sau alte probleme, poate recurge la refacerea înregistrării. În momentul în care consideră că înregistrarea audio este corectă, se poate trece la crearea următoarei înregistrări, prin apăsarea butonului *Next*. Totodată, interfața se va actualiza, un nou set de cifre fiind afișat. Apăsarea butonului *Submit* va avea ca efect trimiterea formularului și a înregistrărilor la server. Acesta va

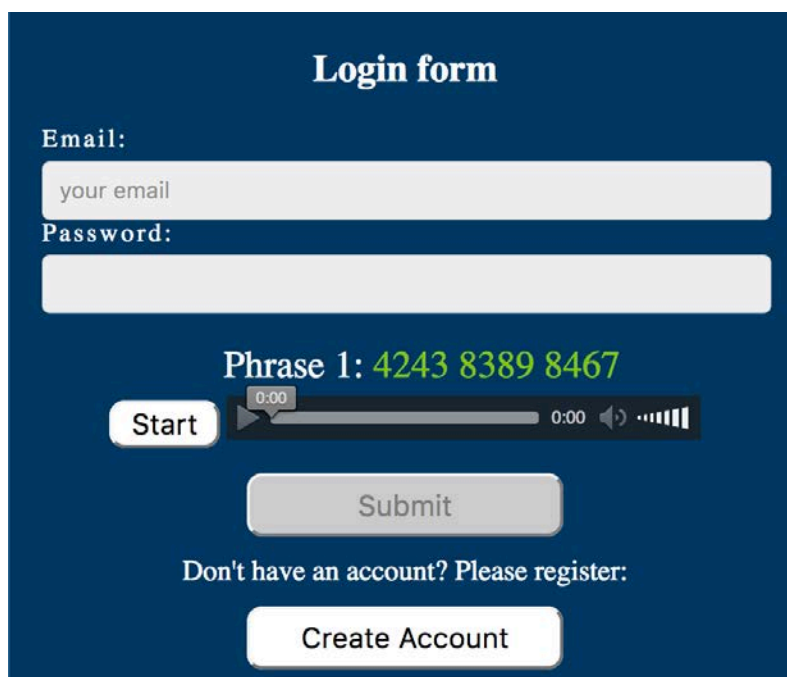
fi disponibil începând cu momentul în care s-a realizat înregistrarea audio a cel puțin 10 seturi de cifre. Pentru a obține o amprentă vocală cât mai puternică, se recomandă crearea a cât mai multor înregistrări audio. La final, utilizatorul va primi un mesaj informativ cu privire la rezultatul acțiunii de înregistrare.

The image shows a 'Register form' interface on a dark blue background. It contains several input fields: 'First name:' with a placeholder 'your first name', 'Last name:' with a placeholder 'your last name', 'Email:' with a placeholder 'your email', 'Password:' (empty), and 'Re-type password:' (empty). Below these is a 'Photo:' section with a file selection button labeled 'Răsfoiește...' and the text 'Niciun fișier selectat.'. At the bottom, there is an audio recording section titled 'Phrase 1: 6459 2588 4266'. It includes a 'Start' button, a progress bar with a play icon and a '0:00' timer, a volume icon, and a 'Next' button. A large 'Submit' button is centered at the very bottom of the form.

Figura 4.2 Formularul de înrolare

Etapă de autentificare

În această etapă, utilizatorul completează un formular cu adresa sa de poștă electronică, parola și înregistrează un singur clip audio. Butonul *Submit* va deveni disponibil, iar datele se vor trimite la serviciul de autentificare unde vor fi realizate cele 3 verificări. În cazul în care una dintre acestea va eșua, utilizatorul va primi un mesaj ce confirmă eroarea întâmpinată. Altfel, va fi redirecționat către contul său personal.



The login form is titled "Login form" in white text on a dark blue background. It contains two input fields: "Email:" with a placeholder "your email" and "Password:". Below these is a phrase "Phrase 1: 4243 8389 8467" in green. A "Start" button is next to a progress bar showing "0:00". A "Submit" button is below the progress bar. At the bottom, there is a link "Don't have an account? Please register:" and a "Create Account" button.

Figura 4.3 Formularul de autentificare

Etapa de accesare a contului personal

În urma autentificării realizate cu succes, aplicația client va afișa contul personal al utilizatorului. Aici se regăsesc datele sale personale oferite în etapa de înregistrare. Totodată, pot fi aduse modificări asupra numelui, parolei și a fotografiei de profil. Pentru părăsirea contului și încheierea sesiunii, se va apăsa butonul *Logout*.



The personal account page is titled "Welcome, Lucian!" in white text on a dark blue background. It features a profile picture of a man with glasses. To the right, there is a "Personal info" section with the name "Lucian", the surname "Georgescu", and the email "lucgeo92@yahoo.com". Below this is a "Logout" button. At the bottom, there are links for "Change name", "Change password", and "Change photo". A "Change your password" section contains three input fields for "Old password:", "New password:", and "Retype new password:", followed by an "Update" button.

Figura 4.4 Pagina contului personal

4.3 SCENARII DE UTILIZARE

Scenariile de utilizare aduc în discuție câteva situații ce pot avea loc în cadrul etapei de autentificare. Se va descrie modul în care aplicația tratează încercarea unui utilizator autorizat de a-și accesa contul, dar și eventualele tentative de acces fraudulos și felul în care acestea sunt prevenite, accentuând rolul serviciilor de verificare de vorbitor și vorbire.

Scenariul 1

Un utilizator autorizat încearcă să își acceseze contul personal creat anterior conform procedurii de înregistrare. În urma completării formularului de autentificare și a creării clipului audio, serviciul de autentificare va constata că parola este cea corectă, serviciul de recunoaștere a vorbitorului va constata că vocea utilizatorului corespunde cu amprenta sa vocală, iar în urma apelării serviciului de recunoaștere de vorbire se va obține transcrierea text a secvenței audio, ce va certifica faptul că a fost rostit setul de cifre cerut.

- Parola corectă ✓
 - Vocea din clipul audio corespunde amprentei vocale ✓
 - Utilizatorul a pronunțat setul de cifre cerut ✓
- } Autentificare reușită ✓

Scenariul 2

Un utilizator neautorizat care cunoaște adresa de poștă electronică și parola unui alt utilizator, încearcă să acceseze fraudulos contul personal al acestuia. În urma completării formularului de autentificare și a creării clipului audio, serviciul de autentificare va constata că parola este cea corectă, în schimb serviciul de recunoaștere a vorbitorului va constata că vocea utilizatorului nu corespunde amprentei sale vocale. Acest fapt va avea ca efect respingerea autentificării, chiar dacă a fost rostit corect setul de cifre cerut.

- Parola corectă ✓
 - Vocea din clipul audio corespunde amprentei vocale ✗
 - Utilizatorul a pronunțat setul de cifre cerut ✓
- } Autentificare reușită ✗

Scenariul 3

Acest scenariu este similar cu cel precedent, singura diferență fiind faptul că utilizatorul neautorizat deține de această dată un fișier audio cu vocea utilizatorului autorizat. Deși amprenta vocală va corespunde cu cea a utilizatorului autorizat, autentificarea va fi din nou respinsă, deoarece fișierul nu va conține rostirea setului de cifre afișat în pagină. Având în vedere faptul că un set nou este generat la fiecare tentativă de autentificare, precum și numărul mare de combinații în care pot apărea cele 12 cifre, șansele sunt extrem de mici că transcrierea returnată de serviciul de recunoaștere de vorbire să fie compatibilă cu setul de cifre afișat.

- Parola corectă ✓
 - Vocea din clipul audio corespunde amprentei vocale ✓
 - Utilizatorul a pronunțat setul de cifre cerut ✗
- } Autentificare reușită ✗

Scenariul 4

Acest scenariu prezintă o metodă prin care utilizatorul neautorizat poate accesa cu succes contul utilizatorului autorizat. Acest fapt reprezintă o vulnerabilitate reală aplicației, fiind indicate totodată câteva soluții ce se pot adopta în această privință.

Similar cu utimele două scenarii anterioare, se presupune că utilizatorul neautorizat cunoaște adresa de poștă electronică, parola aferentă, iar în plus acesta deține un fișier audio cu vocea utilizatorului autorizat, acesta fiind suprins în timp ce rostește un șir de cifre (de exemplu, câteva numere de telefon). Până la acest punct, condițiile inițiale sunt identice cu cele din *Scenariul 3*. Diferența este dată de o nouă abordare în încercarea de a elimina verificarea dată de serviciul de

recunoaștere de vorbire. Fișierul audio ce conține cifre rostite de utilizatorul autorizat va fi împărțit în fișiere mai mici, obținându-se astfel câte un fișier audio în care se va regăsi rostirea unei singure cifre. În acest moment, prin redarea fișierelor în ordinea apariției cifrelor în set, transcrierea returnată de serviciul de recunoaștere de vorbire va corespunde cu setul de cifre afișat, oricare ar fi acesta.

- Parola corectă ✓
 - Vocea din clipul audio corespunde amprenteii vocale ✓
 - Utilizatorul a pronunțat setul de cifre cerut ✓
- } Autentificare reușită ✓

Soluții: În momentul actual, serviciul de recunoaștere de vorbire este capabil să realizeze transcrierea a 10 cuvinte, acestea fiind cifrele sistemului zecimal. Deoarece vocabularul este unul redus, fraudarea sistemului este destul de posibilă, așa cum s-a prezentat în ultimul scenariu. Ca soluție, aplicația client ar trebui configurată astfel încât să afișeze un set de caractere alfanumerice sau chiar grupuri de cuvinte, iar sistemul de recunoaștere de vorbire să fie extins astfel încât să realizeze transcrieri în această nouă situație. Fraudarea ar deveni foarte puțin probabilă, utilizatorul neautorizat va avea nevoie de fișiere audio ce conțin rostirea de către utilizatorul autorizat a cât mai multor cuvinte sau sunete. De asemenea, mixarea acestora într-un singur clip audio ar fi diferită din punct de vedere al articulării vorbirii, în comparație cu un clip audio creat prin utilizarea normală a aplicației.

4.4 TEHNOLOGII UTILIZATE

HTML

HTML (HyperText Markup Language) este un limbaj de marcare utilizat pentru crearea paginilor ce pot fi afișate într-un navigator web. Conținutul unei astfel de pagini poate fi organizat prin folosirea etichetelor și a atributelor.

Etichetele definesc elemente HTML și au un format specific, denumirea elementului regasindu-se între caracterele "<" și ">". De regulă, orice etichetă deschisă atrage după sine și închiderea ei. De exemplu, deschiderea etichetei `<div>` ce definește un bloc, o secțiune a paginii web, va exista în pereche cu etichetă ce marchează închiderea să, `</div>`. Elementele HTML pot fi imbricate, între etichetele de deschidere și închidere ale unui element se pot afla alte elemente. Codul unei pagini web minimale trebuie să fie cuprins între etichetele `<html>` și `</html>`. Apoi, în interiorul acestora, se vor regăsi alte două elemente. Primul element, `<head>` `</head>`, va conține la rândul său elemente ce definesc titlul paginii, precum și elemente ce fac referință la paginile de stil sau la scripturile incluse. Al doilea element, `<body>` `</body>`, va conține elemente ce alcătuiesc pagina propriu-zisă. Alte elemente uzuale sunt cele ce definesc paragrafe, imagini, legături, tabele, etc.

Atributele oferă informații suplimentare asupra elementelor. De regulă, se folosesc atribute ce definesc numele elementului, atribute ce stabilesc un identificator unic al elementului, sau atribute ce plasează elementul într-o clasă de elemente. În funcție de elementul căruia aparțin, atributele pot oferi detalii despre dimensiuni (în cazul blocurilor sau al imaginilor), despre fonturi și culori (în cazul paragrafelor). În exemplul următor, elementul `<p>` `</p>` marchează introducerea unui paragraf, iar `align` și `id` sunt atributele sale:

`<p align="center" id="paragraf">Acesta este un exemplu </p>`.

CSS

De la bun început, scopul elementelor HTML a fost de a structura paginile web. Mai apoi, a devenit necesară formatarea aspectului și au apărut elemente ca `<font>` `</font>` sau atribute ce modifică dimensiunile, culoarea, spațierile, etc. Aceste lucruri însă au făcut din dezvoltarea paginilor un proces dificil.

CSS (Cascading Style Sheets) descrie modul în care elementele HTML ar trebui să fie afișate. CSS a apărut din necesitatea de a formata conținutul paginilor web într-un mod organizat și eficient. La începutul paginii HTML se va introduce în interiorul elementul `<head>` `</head>` un rând de forma: `<link rel="stylesheet" type="text/css" href="style.css">`. Acesta realizează o referință către fișierul `style.css` ce va conține formatarea elementelor. Fiecare element poate fi referit prin identificatorul său unic, clasa din care face parte sau chiar denumirea elementului. În exemplul de mai jos este prezentată formatarea unui element de tip bloc referit prin identificatorul sau, numit *container*:

```
#container{
Margin-top: 20px;
Color: white;
Font-size: 24 px;
}
```

Dintre avantajele folosirii CSS pot fi amintite:

- Eficiență crescută: o singură secvență de cod CSS poate fi aplicată unui grup de elemente HTML;
- Menținanța simplificată: o singură modificare în pagina de stil se va aplica în toate documentele HTML unde este inclusă;
- Datorită partajării codului CSS, paginile se vor încărca mai repede;
- CSS oferă mult mai multe posibilități de stilizare în comparație cu aspectul realizat cu ajutorul atributelor;
- CSS poate realiza compatibilitatea cu diverse dispozitive: același document HTML poate fi prezentat în mod diferit, astfel încât aspectul să fie optim, în funcție de tipul dispozitivului, dimensiunea afișajului, performanțele de calcul, etc.

JavaScript

JavaScript este un limbaj de programare bazat pe obiecte, folosit mai ales pentru introducerea unor funcționalități în paginile web. Este un limbaj de scripting ce nu necesită compilare, codul fiind interpretat. Oferă interactivitate documentelor HTML și poate fi înglobat direct în acestea, regăsindu-se deseori în secțiunea `<head>` `</head>`, cuprins între etichetele `<script type="text/javascript">` `</script>` sau poate fi adăugată o referință către un fișier extern `<script src="javascript_code.js">` `</script>`.

JavaScript interacționează cu evenimente generate de utilizator sau sistem. Un astfel de eveniment poate fi apăsarea unui buton, terminarea încărcării unei pagini sau editarea unui câmp de text. Pot fi utilizate mai multe gestionare de evenimente: *onclick* (declanșează execuția unei funcții atunci când un element HTML este apăsăat), *onload* și *onunload* (declanșează execuția unei funcții atunci când utilizatorul accesează sau părăsește pagina), *onchange* (folosit de obicei odată cu validarea câmpurilor de intrare), *onmouseover*, *onmouseout*, *onmousedown*, *onmouseup* (declanșează execuția unei funcții în funcție de poziția mouse-ului față de un element HTML), *onsubmit* (va apela o funcție atunci când se acționează butonul *Submit* al unui formular). Instrucțiunea `document.write()` este modalitatea prin care poate fi adăugat un conținut variabil în pagină HTML. Instrucțiunea `document.getElementById().value` permite citirea sau modificarea valorii unui element HTML.

Fiind un limbaj ce permite programarea pe obiecte, JavaScript permite atât utilizarea obiectelor deja încorporate, cât și crearea unor obiecte proprii. Fiecare obiect este caracterizat de proprietăți și metode specifice. Câteva obiecte standard ce se regăsesc în limbajul JavaScript: *String*, *Date*, *Array*, *Math*, *RegExp*, *Number*, *Cookie*. Alte obiecte importante sunt: *Window* (reprezintă o fereastră deschisă în navigator), *Screen* (conține informații despre ecranul utilizatorului), *Location* (conține informații despre url-ul curent).

Ajax

Ajax (Asynchronous Javascript and XML) este o tehnologie ce permite actualizarea elementelor dintr-o pagină web în mod dinamic, fără întreruperi pentru încărcări de pagină sau reîmprospătări de ecran. Cererile și răspunsurile serverului nu mai trebuie să corespundă cu anumite acțiuni ale utilizatorului, ci pot avea loc în orice moment convenabil pentru aplicație. Navigatorul permite ca utilizatorul să continue folosirea paginii, în timp ce cererile trimise serverului se vor efectua în mod asincron.

Caracteristic Ajax este obiectul *XMLHttpRequest*, acesta oferind posibilitatea de a trimite cereri HTTP fără ca finalizarea acestora să corespundă neapărat cu încărcarea unei pagini noi. Obiectul este caracterizat de mai multe metode și proprietăți. Metodă *open()* pregătește comunicarea cu serverul, în timp ce metodă *send()* trimite cererea. Proprietatea *readyState* poate lua valori numerice în intervalul [0;4] ce semnifică starea cererii: neinitializată, în încărcare, încărcată, interactivă și finalizată. Proprietatea *onreadystatechange* determina ce manager de eveniment va fi apelat când se modifică proprietatea *readyState* a obiectului. Proprietatea *response* reprezintă datele returnate de către server.

Web Audio API

Web Audio API (Application Programming Interface) este un API Javascript de nivel înalt specializat în prelucrarea și sintetizarea audio în aplicații web. Scopul acestuia este de a include capacități de mixare, procesare și filtrare, atât de specifice aplicațiilor audio moderne.

API-ul este construit în jurul conceptului de *AudioContext*. Acesta este reprezentat de un graf orientat ce definește modul în care fluxul audio trece de la sursă spre destinație. Fiecare nod prin care trece poate inspecta sau modifica proprietățile fluxului. Figură următoare prezintă un astfel de graf, caracterizat de mai multe noduri cu rol de sursă, noduri intermediare cu rol de filtru, precum și un nod destinație.

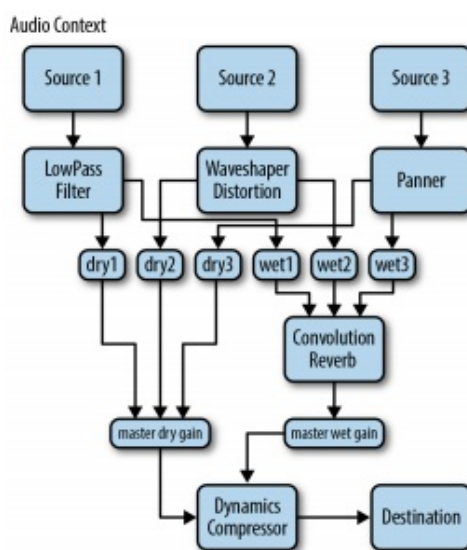


Figura 4.5 Model de Context Audio (Web Audio API). Sursa:[21]

Sursa poate fi reprezentată de un fișier audio sau un microfon, în timp ce destinația constă într-un dispozitiv de redare audio. În cazul mai multor surse, cum este cazul și în imaginea prezentată, semnalele audio vor fi pur și simplu îmbinate. Nodurile oferă posibilitatea de a fi conectate și deconectate, iar asupra fluxului audio se pot aplica diverse filtre și analizoare de semnal.

Java

Java este un limbaj de programare de nivel înalt, obiect-orientat, dezvoltat în special pentru a avea cât mai puține dependențe posibile. Inițiatorul proiectului este James Gosling, iar prima versiune, Java 1.0 a fost lansată public în 1995 în cadrul companiei Sun Microsystems. Acesta a

fost conceput astfel încât să fie scris o singură dată și de a rula oriunde ("write once, run anywhere"). Aplicațiile Java sunt compilate ca bytecode (fișiere de tip clasa) ce pot rula pe orice JVM (Java Virtual Machine), indiferent de arhitectura sistemului de calcul. Din acest punct de vedere, Java este un limbaj portabil.

Java este considerat un limbaj simplu, foarte multe caracteristici ale sale fiind derivate din C/C++ și este bazat în totalitate pe conceptele POO (programării obiect-orientată). A fost conceput pentru aplicații ce lucrează în mod distribuit, fiind utilizate deosebi protocoale de comunicații pentru apelarea la distanță a funcțiilor și accesarea resurselor. O altă caracteristică este dată de posibilitatea de a rula simultan pe mai multe fire de execuție, aplicațiile având un grad de interactivitate crescut. Codul Java este sigur, acesta fiind verificat înainte de execuție, pentru a se confirma faptul că a fost compilat de un compilator standard. Erorile de acces la memorie sunt eliminate, datorită faptului că nu se folosesc pointeri.

În prezent, Java a ajuns la versiunea 8 și este unul dintre cele mai utilizate limbaje de programare, aflându-se la baza multor aplicații și servicii web.

Arhitectura REST

REST (Representational state transfer) este un stil arhitectural care definește un set de principii ce fac mult mai scalabilă și flexibilă dezvoltarea aplicațiilor web. În anul 2000, Roy Fielding a introdus pentru prima dată acest concept, definind un set de constrângeri ce stau la baza serviciilor REST:

- Adresabilitate: datele și informațiile sunt privite în mod abstract sub forma unor resurse, iar fiecare resursă poate fi accesată printr-un URI (Uniform Resource Identifier).
- Interfața uniformă: între clienți și servere există o interfață bine definită și simplificată, bazată pe protocolul HTTP, ce permite manipularea resurselor.
- Lipsa stării: orice cerere HTTP are loc într-o izolare completă, fiind independente față de cererile precedente. Toate informațiile necesare îndeplinirii unei cereri sunt incluse în această.
- Client-server: interfață uniformă ce separă aceste două entități. Fiecare are un rol specific, fiind posibilă dezvoltarea în mod independent, atât timp cât interfața rămâne comună. Clientul nu este preocupat de stocarea datelor, portabilitatea codului său fiind îmbunătățită. În schimb, serverul nu este preocupat de interfațarea cu utilizatorul.
- Sistem stratificat: în orice moment, clientul poate fi conectat direct la serverul final sau la un intermediar. Serverele intermediare îmbunătățesc scalabilitatea sistemului.

Serviciile REST au la bază protocolul HTTP. Acesta presupune un schimb de mesaje de tip cerere/răspuns. Clientul inițiază un mesaj de tip cerere către server în care specifică metoda HTTP folosită, locația resursei, antetul și corpul mesajului. Răspunsul serverului conține un cod de răspuns, un antet și un corp al mesajului. Cele mai importante metode HTTP sunt:

- GET: folosită pentru a interoga serverul în vederea obținerii unei informații specifice. Este o operație idempotentă și sigură; de oricâte ori va fi apelată, rezultatul va fi același iar starea serverului nu va fi influențată.
- PUT: solocita serverului să stocheze coprul mesajului transmis la locația specificată în mesajul HTTP. Este folosită pentru a insera sau modifica date.
- DELETE: folosită pentru a șterge resurse.
- POST: folosită pentru crearea unei resurse. Solicită serverului să accepte și să stocheze datele cuprinse în corpul mesajului.

În urma unei cereri HTTP se generează un cod de răspuns, mai jos fiind amintite principalele clase în care acestea sunt grupate:

- 1xx: clasa codurilor de informare;

- 2xx: clasa specifică cererilor ce au avut loc cu succes;
- 3xx: clasa codurilor de redirectionare. Clientului îi este indicat să realizeze o acțiune suplimentară pentru finalizarea cererii;
- 4xx: clasa codurilor ce indică erori generate de client;
- 5xx: clasa codurilor ce indică erori generate de server.

Jersey

Jersey este o platformă de dezvoltare a serviciilor REST folosind limbajul Java. Jersey oferă programatorilor o serie de adnotări ce pot fi asociate claselor și metodelor cu scopul de a defini resursele și acțiunile ce pot fi efectuate pe baza acestora. Adnotările se scriu înaintea definirii unei clase sau metode. Cele mai întâlnite sunt:

- `@Path()`: valoarea acestei adnotări este un URI relativ ce desemnează calea către clasa sau metoda căreia îi este asociată.

Exemplu: <http://localhost/VoiceApp/login/auth>. Acest URI este format din: protocol, adresa serverului, numele serviciului web, numele clasei și numele funcției ce va prelua cererea HTTP.

- `@GET`, `@POST`, `@PUT`, `@DELETE`: aceste adnotări specifică tipul cererii HTTP ce poate fi prelucrată de către o funcție.
- `@Consumes`: această adnotare specifică tipul media al resursei primite de către server.
- `@Produces`: această adnotare specifică tipul media al resursei ce va fi trimise înapoi clientului sub forma unui răspuns HTTP.

Jersey oferă suport atât pentru crearea aplicațiilor server, cât și pentru aplicațiile client.

4.5 IMPLEMENTAREA APLICAȚIEI CLIENT

Aplicația client oferă utilizatorului o interfață grafică intuitivă cu ajutorul căreia datele sale pot fi preluate sau afișate. A fost dezvoltată cu ajutorul tehnologiilor HTML, CSS și Javascript. Interacțiunea aplicației client cu serverul este realizată cu ajutorul tehnologiei Ajax. Funcționalitatea de înregistrare audio a fost posibilă prin inserarea în aplicație a componentei software RecorderJS [22]. Următoarea diagramă prezintă structura aplicației client, fiind prezentate fișierele componente și legătura dintre acestea.

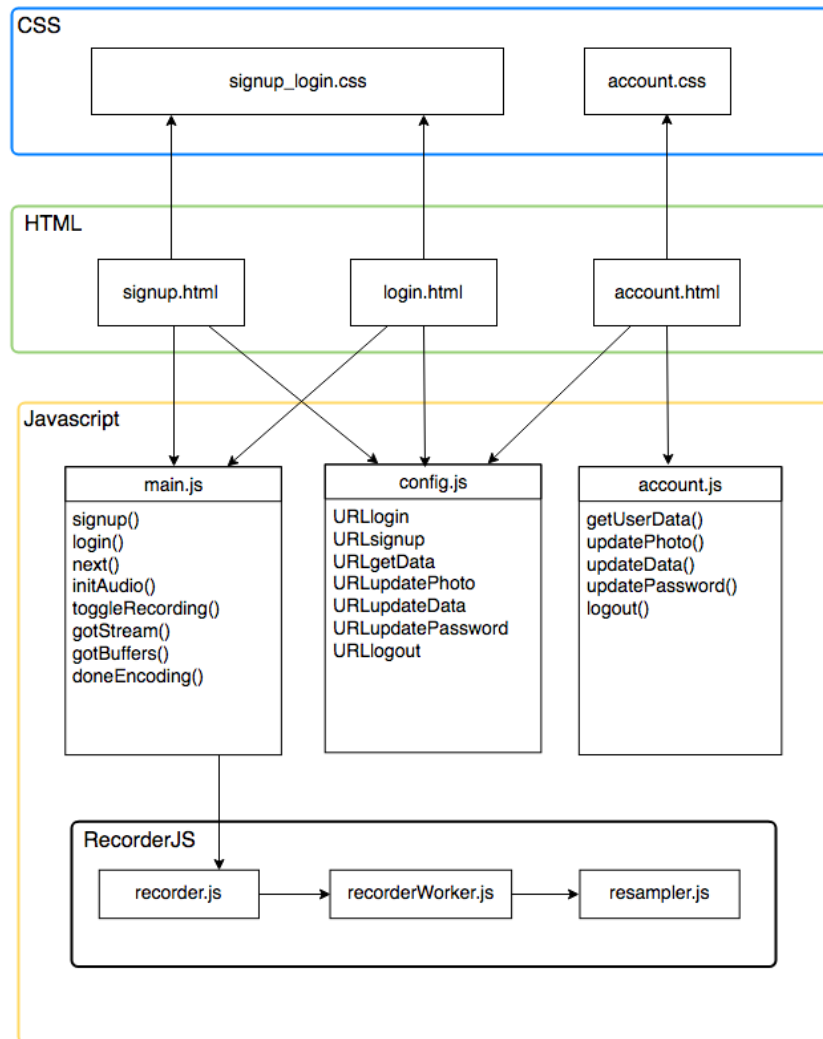


Figura 4.6 Structura aplicației client

HTML

Paginile corespunzătoare procesului de înrolare și autentificare (*signup.html* și *login.html*) conțin o referință către fișierul *signup_login.css* în care este definit aspectul elementelor HTML, precum și fișierele *main.js* și *config.js*. Cele două pagini HTML o structură similară, fiind formate dintr-un formular ce conține câmpuri de intrare, necesare pentru completarea datelor personale. Datele personale, cât și înregistrările audio, vor fi preluate de către funcția *signup()*, respectiv *login()*, definite în fișierul *main.js*.

```
<form name="userData" method="post" onsubmit="signup(); return false;">
  <label>First name: </label><input type="text" name="firstName" id="firstName" placeholder="your first name" required>
  <label>Last name: </label><input type="text" name="lastName" id="lastName" placeholder="your last name" required>
  <label>Email: </label><input type="email" name="email" id="email" placeholder="your email" required>
  <label>Password: </label><input type="password" name="password" id="password" required >
  <label>Re-type password: </label><input type="password" name="retypedPassword" id="retypedPassword" required>
  <label>Photo: </label><input type="file" name="photo" id="photo" size="45" required >
```

În afara formularului, paginile conțin un bloc corespunzător procedurii de creare a clipurilor audio. Aici se regăsesc: butonul de *Start/Stop*, ce activează funcția *toggleRecording()*, butonul *Next*, prin a cărui apăsare se execută funcția *next()*, elemente HTML pentru redare audio. Toate acestea pot fi observate în imaginea următoare:

```

<div id="container">
<div id="wrap">
  <div id="numbers" align="center"> </div>
  <div id="phrase"></div>
</div>
  <div id="controls" align="center">
    <table>
      <tr><td><button type="button" id="toggle" onclick="toggleRecording(this);">Start</button></td>
      <td><audio id="fisier" controls="controls" preload="none"></audio></td>
      <td><button id="nextAudio" type="button" onclick="next()">Next</button></td>
    </table>
    <a id="save" href="#"></a>
    <div id="recordingState"></div>
  </div>
</div>

```

Pagina *account.html* conține câteva blocuri în care vor fi afișate datele utilizatorului autentificat, precum și câmpuri de intrare în care se pot efectua modificări asupra datelor personale (nume, parola, fotografie de profil). După acest pas, prin apăsarea butonului *Update* va fi executată una din funcțiile *updateData()*, *updatePassword()* sau *updatePhoto()*, definite în fișierul *account.js*. Modificările vor fi trimise spre server în vederea actualizării bazei de date.

CSS

Cele două fișiere CSS ale aplicației conțin detalii referitoare la formatarea elementelor HTML. Acestea sunt referite pe baza denumirii elementului, identificatorul unic sau numele clasei din care face parte. Culoarele de fundal, distanțele dintre elemente și dimensiunile lor, aspectul textului sunt câteva dintre proprietățile definite în paginile de stil. Următorul exemplu reprezintă formatarea butonului de trimitere a formularului de înrolare:

```

#submitSignup
{
  border-radius: 10px;
  width:200px;
  height:40px;
  background-color:#CCCCCC;
  font-size:18px;
  margin-bottom:10px;
  margin-top:10px;
}

```

Javascript

Fișierul *main.js* este inclus în paginile *signup.html* și *login.html* și are rolul de a prelua datele din acestea și de a le trimite la server. De asemenea, în fișierul *main.js* este creat un obiect de tip *Recorder*, definit în cadrul modului *RecorderJS*. Acesta captează semnalul audio și îl salvează ca obiect *Blob*.

Funcțiile *login()* și *signup()* citesc valorile elementelor HTML cu ajutorul instrucțiunii *document.getElementById()*. Aceste valori vor fi adăugate într-un formular de tipul *FormData()*. Un obiect *XMLHttpRequest()* va deschide o conexiune către server și va trimite o cerere HTTP ce conține formularul creat. Va fi citit codul de răspuns ce indică starea cererii, precum și răspunsul propriu zis. Secvența următoare conține definiția funcției *login()*:

```

function login()
{
  var email = document.getElementById("email").value;
  var password = document.getElementById("password").value;

  var r=new XMLHttpRequest();
  var form=new FormData();
  form.append("email",email);
  form.append("password",password);
  form.append("audioFile",bloblogin);

```

```

r.open("POST",URLlogin,true);
r.onreadystatechange=function()
{
    if (r.readyState===4)
        if(r.status===200)
        {
            window.location.replace("account.html");
        }
        else
            alert(r.response);
};
r.send(form);
}

```

Fișierul *account.js* este responsabil de opțiunile oferite utilizatorului în momentul în care acesta este autentificat și își vizitează pagina personală. Sunt definite funcții de obținere a datelor personale de la server, precum și de actualizare a acestora. Sunt folosite obiecte de tip XMLHttpRequest(), iar elementele HTML sunt accesate într-un mod similar cu cel prezentat anterior.

Fișierul *config.js* este un fișier de configurare, acesta conține adresele URL corespunzătoare diverselor metode din serviciul de autentificare spre care se îndreaptă cererile HTTP.

RecorderJS este un modul software dezvoltat în Javascript ce utilizează Web Audio API. În fișierul *recorder.js* este definit obiectul *Recorder*, capabil să capteze audio în format WAV. Fișierul *resampler.js* definește o funcție ce permite modificarea frecvenței de eșantionare a semnalului, fiind posibilă și alegerea numărului de canale luate în considerare. În cazul de față, se realizează conversia la 16 KHz și este folosit un singur canal deoarece aceasta este configurația audio necesară serviciilor de recunoaștere de vorbire și vorbitor. Pentru a nu fi afectată performanța paginii și interacțiunea cu utilizatorul, toate aceste operații sunt executate în fundal de către un obiect de tip *Worker()*.

4.6 IMPLEMENTAREA BAZEI DE DATE

Baza de date aferentă serviciului de autentificare rulează în cadrul unui server MySQL și este compusă din 3 tabele: *users*, *audio*, *login*.

users	audio	login
userid INT(11)	index INT(11)	index INT(11)
first_name VARCHAR(45)	userid INT(11)	userid INT(11)
last_name VARCHAR(45)	wav_file MEDIUMBLOB	password VARCHAR(45)
email VARCHAR(45)	displayed_number VARCHAR(45)	wav_file MEDIUMBLOB
password VARCHAR(45)	transcription VARCHAR(45)	displayed_number VARCHAR(45)
photo MEDIUMBLOB	wer FLOAT	transcription VARCHAR(45)
average_wer FLOAT	date_time DATETIME	wer FLOAT
		date_time DATETIME
		valid_password BIT(1)
		valid_wer BIT(1)
		valid_voice BIT(1)

Figura 4.7 Tabelele bazei de date

Primul tabel conține datele personale ale utilizatorilor și valoarea medie a erorii la nivel de cuvânt, obținută în urma transcrierii clipurilor audio create în etapa de înrolare. În următoarea secvență SQL se poate observa instrucțiunea de creare a tabelului, tipurile de date ale câmpurilor, cheile primară și unică.

```
CREATE TABLE `users` (
```

```

`userid` INT(11) NOT NULL AUTO_INCREMENT,
`first_name` VARCHAR(45) NOT NULL,
`last_name` VARCHAR(45) NOT NULL,
`email` VARCHAR(45) NOT NULL,
`password` VARCHAR(45) NOT NULL,
`photo` MEDIUMBLOB NOT NULL,
`average_wer` FLOAT NULL,
PRIMARY KEY (`userid`),
UNIQUE INDEX `email_UNIQUE` (`email` ASC));

```

Al doilea tabel conține clipurile audio provenite de la toți utilizatorii înrolați în sistem. Pentru fiecare clip se cunoaște vorbitorul, secvența de cifre afișată în aplicația client, secvența de cifre rezultată în urma transcrierii realizate de serviciul de recunoaștere de vorbire, rata de eroare la nivel de cuvânt și momentul de timp la care a fost creat.

```

CREATE TABLE `audio` (
  `index` INT NOT NULL AUTO_INCREMENT,
  `userid` INT NOT NULL,
  `wav_file` MEDIUMBLOB NOT NULL,
  `displayed_number` VARCHAR(45) NOT NULL,
  `transcription` VARCHAR(45) NULL,
  `wer` FLOAT NULL,
  `date_time` DATETIME NOT NULL,
  PRIMARY KEY (`index`));

```

Al treilea tabel este dedicat tentativelor de autentificare, reușite sau nu. Acesta conține identificatorul numeric al utilizatorului, parola completată, secvența de cifre afișată și cea transcrisă, rata de eroare la nivel de cuvânt, momentul de timp, precum și 3 câmpuri ce pot lua valori binare în funcție de rezultatul celor 3 verificări efectuate la autentificare.

```

CREATE TABLE `login` (
  `index` INT NOT NULL AUTO_INCREMENT,
  `userid` INT NOT NULL,
  `password` VARCHAR(45) NOT NULL,
  `wav_file` MEDIUMBLOB NOT NULL,
  `displayed_number` VARCHAR(45) NOT NULL,
  `transcription` VARCHAR(45) NULL,
  `wer` FLOAT NULL,
  `date_time` DATETIME NOT NULL,
  `valid_password` BIT(1) NULL DEFAULT b'0',
  `valid_wer` BIT(1) NULL DEFAULT b'0',
  `valid_voice` BIT(1) NULL DEFAULT b'0',
  PRIMARY KEY (`index`));

```

Toate acțiunile asupra bazei de date au loc prin intermediul procedurilor stocate. A fost aleasă această metodă pentru o mai bună organizare a codului aplicației. Au fost create 19 proceduri stocate:

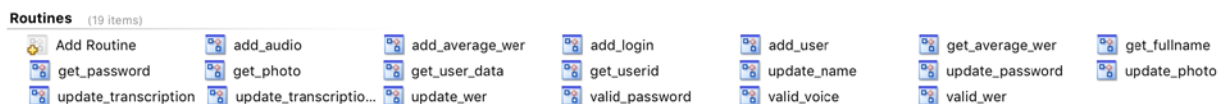


Figura 4.8 Lista procedurilor stocate

Următoarea secvență SQL prezintă instrucțiunea de creare a procedurii stocate corespunzătoare acțiunii de extragere a identificadorului numeric al utilizatorului:

```
DELIMITER $$  
CREATE PROCEDURE `get_userid` (IN arg_email VARCHAR (45), OUT arg_userid INT)  
BEGIN  
    SELECT userid INTO arg_userid from users WHERE email=arg_email;  
END $$  
DELIMITER ;
```

Procedura este definită sub forma unei funcții cu două argumente. Adresa de poștă electronică reprezintă argumentul de intrare, pe baza acestuia se va selecta din tabel identificadorul numeric corespunzător, livrat apoi ca argument de ieșire.

4.7 IMPLEMENTAREA SERVICIULUI DE AUTENTIFICARE

Serviciul de autentificare a fost dezvoltat în mediul de programare Eclipse Marș, utilizând limbajul Java și platforma de dezvoltare Jersey, specifică serviciilor REST. Este compus din mai multe clase organizate în pachete. În următoarea imagine este prezentată diagramă UML a claselor:

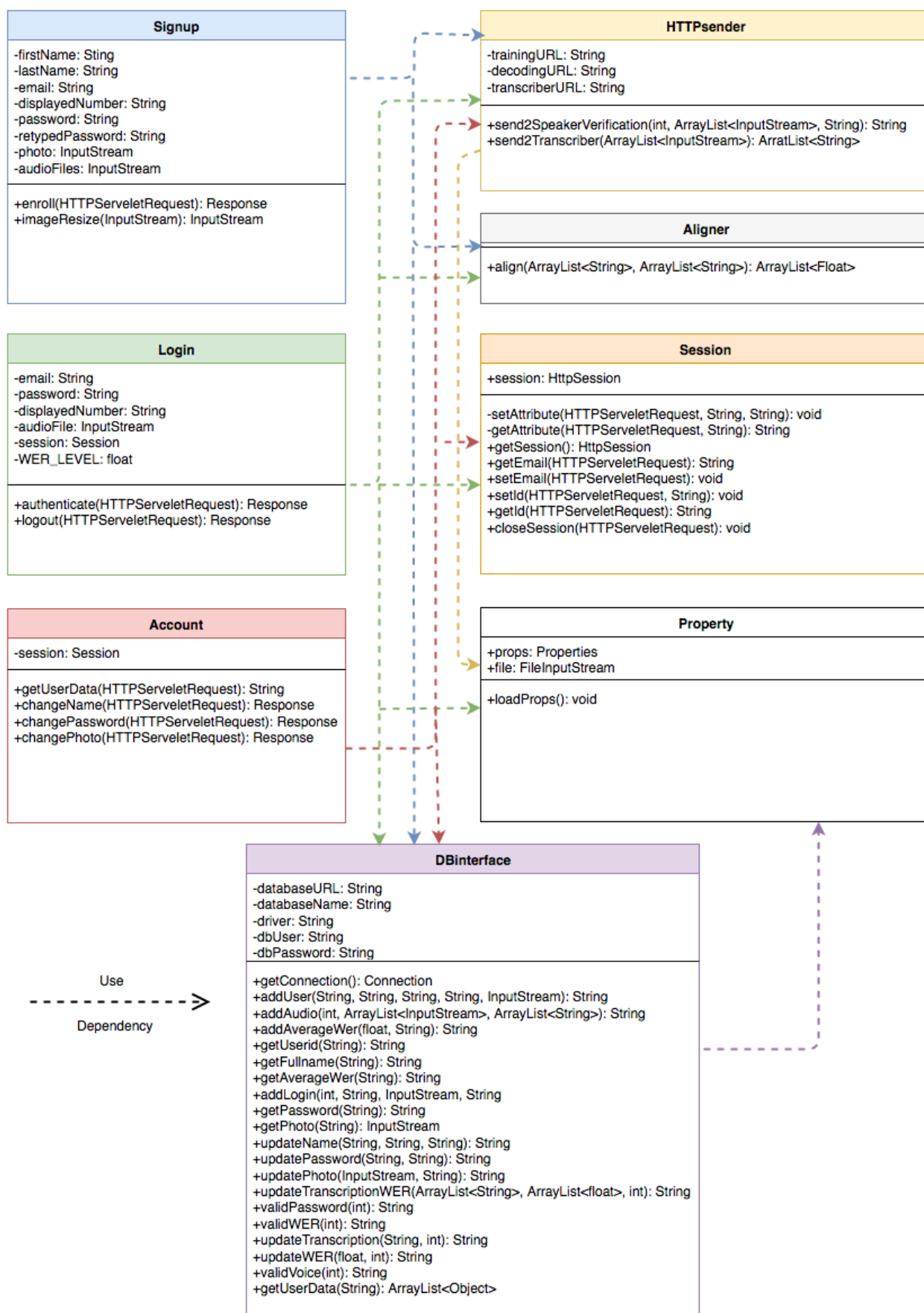


Figura 4.9 Diagramă UML a claselor

Clasa Dbinterface.java

Această clasă funcționează ca o interfață între serviciul REST propriu zis și baza de date, unde fiecărei proceduri stocate îi corespunde o funcție definită în clasă. Următoarea secvență prezintă funcția corespunzătoare extragerii identificadorului numeric al unui utilizator, cunoscându-se adresa sa de poștă electronică:

```
public String getUserId(String email){
    con = getConnection();
    try{
        CallableStatement statement = null;
        String query = "{call get_userid(?,?)}";
        statement = (CallableStatement) con.prepareCall(query);
        statement.setString(1, email);
        statement.registerOutParameter(2, java.sql.Types.INTEGER);
        int queryResult=statement.executeUpdate();
        userid = statement.getInt(2)+"";
        statement.close();
        con.close();
    }catch(SQLException e){
        return Finals.NONEXISTENT_EMAIL;
    }
    return userid;
}
```

După stabilirea conexiunii cu baza de date, se utilizează un obiect *CallableStatement*, căruia îi vor fi atribuiți parametri de intrare și ieșire, conform procedurii stocate apelate. În urma execuției, este extrasă valoarea parametrului de ieșire.

Clasa HTTPsender.java

Prin cele două funcții ale sale, clasa realizează comunicarea serviciului de autentificare cu celelalte două servicii, cel de recunoaștere de vorbire, respectiv recunoaștere de vorbitor. Către acestea sunt trimise prin protocolul HTTP clipurile audio primite de la aplicația client. În continuare sunt prezentate principalele secvențe de cod ale uneia dintre funcții:

```
public static ArrayList<String> send2Transcriber(ArrayList<InputStream> file){

    ClientConfig config = new DefaultClientConfig();
    Client client = Client.create(config);
    WebResource service =
client.resource(UriBuilder.fromUri(transcriberURL).build());
.....
FormDataMultiPart formDataMultiPart = new FormDataMultiPart();
    for (int i = 0; i < numberOfFiles; i++){
        formDataMultiPart.field("audioFile",
file.get(i),MediaType.APPLICATION_OCTET_STREAM_TYPE);
    }
.....
clientResp =
service.type(MediaType.MULTIPART_FORM_DATA_TYPE).post(ClientResponse.class,
formDataMultiPart);
.....
String response=clientResp.getEntity(String.class);
.....
}
```

Este creat și configurat un obiect *Client* capabil să execute o cerere HTTP către serviciul de recunoaștere de vorbire. Apoi este creat obiect de tip formular ce conține mai multe câmpuri, fiecare câmp fiind denumit *audioFile* și având atribuit un clip audio. În urma execuției cererii, răspunsul întors (transcrierea fișierelor) este stocat într-un obiect *String*.

Clasa Session.java

Cererile HTTP sunt de regulă complet independente. Totuși, uneori este nevoie ca anumite informații să fie reținute de la o cerere la alta. Acest lucru este realizabil prin intermediul sesiunilor. Atunci când un utilizator vizitează pentru prima dată o pagină web, se creează o sesiune caracterizată de un identificator numeric unic. Acesta este păstrat atât pe parte de server, cât și pe parte de client, într-un fișier *cookie*. O sesiune poate stoca atribute cu scopul de a le citi mai târziu. După ce un utilizator s-a autentificat, este de dorit ca el să fie recunoscut pe parcursul cererilor următoare. Astfel, se va salva într-o sesiune adresa sa de poștă electronică și identificatorul numeric corespunzător ce se regăsește în baza de date. Dacă ulterior va executa o acțiune care necesită în prealabil autentificare, se va consulta informația salvată în sesiune pentru a stabili identitatea utilizatorului. Clasa *Session.java* implementează metode de manipulare a atributelor unei sesiuni.

```
private void setAttribute(@Context HttpServletRequest request, String numeAtribut,
String valoareAtribut){
    session = request.getSession(true);
    session.setAttribute(numeAtribut, valoareAtribut);
}

private String getAttribute(@Context HttpServletRequest request, String numeAtribut){
    session = request.getSession(true);
    Object result = session.getAttribute(numeAtribut);
    if (result != null) {
        return result.toString();
    }else {
        return null;
    }
}

public void closeSession(@Context HttpServletRequest request){
    session = request.getSession(true);
    session.invalidate();
}
```

Clasa Aligner.java

Este o interfață pentru clasa *NISTAlign.java*, disponibilă în utilitarul CMU Sphinx. Pe baza secvenței de cifre afișată în aplicația client și a secvenței de cifre obținute în urma transcrierii unui clip audio, această clasă calculează rata de eroare la nivel de cuvânt.

Clasa Property.java

Este o clasă prin intermediul căreia sunt inițializate anumite constante din serviciul web cu valorile lor corespunzătoare aflate într-un fișier de configurare.

Clasa Signup.java

```
@Path("signup")
public class Signup {
    .....
    @POST
    @Path("enroll")
    @Consumes(MediaType.MULTIPART_FORM_DATA)
    public Response enroll(@Context HttpServletRequest request){

List<FileItem> items = new ServletFileUpload(new
DiskFileItemFactory()).parseRequest(request);

        for (FileItem item : items) {
            String fieldname = item.getFieldName();
            switch (fieldname) {
```

```

        case "audioFile":
            audioFiles.add(item.getInputStream());
            break;
        case "photo":
            photo = imageResize(item.getInputStream());
            break;
        case "displayedNumber":
            displayedNumber.add(item.getString());
            break;
        case "lastName":
            lastName = item.getString();
            break;
        case "firstName":
            firstName = item.getString();
            break;
        case "email":
            email = item.getString();
            break;
        case "password":
            password = item.getString();
            break;
        case "retypedPassword":
            retypedPassword = item.getString();
            break;
    }
}

```

```

        .....
        return Response.status(Status.OK).build();
    }

```

În urma unui HTTP POST către calea specificată în adnotări, va fi executată funcția *enroll*. Prin analiza conținutului cererii, va fi identificat un formular cu mai multe câmpuri. Acesta conține datele personale ale utilizatorului și clipurile audio înregistrate. Utilizând un obiect de tip *DBInterface*, datele vor fi introduse în baza de date, iar cu ajutorul clasei *HTTPsender* clipurile audio vor fi trimise la serviciile de verificare de vorbire și vorbitor. Dacă toate operațiile decurg cu succes, aplicația client va primi înapoi un mesaj ce confirmă acest lucru.

Clasa **Login.java** o structură similară cu aceasta și nu va mai fi prezentată.

Clasa Account.java

Dispune de o metodă de obținere a datelor utilizatorului și 3 metode de modificare a datelor acestuia.

```

@Path("/account")
public class Account {
    @GET
    @Path("/getUserData")
    @Produces("application/xml")
    public String getUserData(@Context HttpServletRequest request) throws{

```

```

        .....
        String email=session.getEmail(request);
        if(email == null){
            return "you are not logged!";
        }

```

```

        .....
        DBInterface db_manager = new DBInterface();
        userData=db_manager.getUserData(email);

```

```

        firstName=(String) userData.get(0);
        lastName=(String) userData.get(1);
        photo=(byte[]) userData.get(2);
        String photoBase64 = new sun.misc.BASE64Encoder().encode(photo);
        String xmlOutput("<user>"
            + "<firstName>" + firstName + "</firstName>"
            + "<lastName>" + lastName + "</lastName>"
            + "<photo>" + photoBase64 + "</photo>"
            + "<email>" + email + "</email>"
            + "</user>";
        return xmlOutput;
    }

```

Funcția *getUserData()* este apelată în urma unei cereri HTTP GET către calea specificată în adnotări și întoarce un răspuns în format XML. Deoarece extragerea datelor personale este o operație ce necesită ca utilizatorul să fie autentificat, acest lucru se va verifica prin citirea sesiunii. Aceasta ar trebui să conțină adresa de poștă electronică a utilizatorului. Apoi, prin utilizarea unui obiect *DBInterface*, datele personale ale utilizatorului vor fi extrase din baza de date și returnate aplicației client.

Funcțiile de modificare a datelor personale sunt apelate în urma unor cereri HTTP POST, fiind asemănătoare cu funcția *enroll* din clasa *Signup.java*.

4.8 CONCLUZII ȘI DEZVOLTĂRI ULTERIOARE

Aplicația oferă posibilitatea autentificării pe bază de parolă și recunoaștere de vorbire și amprentă vocală, reușind să îndeplinească sarcina pentru care a fost concepută. Precizia cu care unui utilizator îi este acceptată sau respinsă încercarea de autentificare este strict proporțională cu eficiența sistemelor de recunoaștere de vorbire și vorbitor. Totodată, aplicația poate fi considerată un punct de plecare în vederea realizării serviciilor de tip REST.

Pentru dezvoltările viitoare ale proiectului sunt vizate următoarele ținte:

- Implementarea unor soluții de eliminare a vulnerabilității descrisă în secțiunea 4.3 Scenarii de utilizare (Scenariul 4).
- Extinderea aplicației sub forma unui proiect pilot ce presupune provocarea cât mai multor utilizatori de a se înrola în aplicație. Crearea unui sistem competițional în care utilizatorii să fie încurajați să încerce autentificări neautorizate. În acest fel, aplicația va fi testată intens și vor putea fi observate alte posibile probleme. Cu această ocazie se vor obține baze de date de vorbire de mari dimensiuni, iar sistemul de recunoaștere de vorbire (recunoașterea secvențelor de cifre) ar putea fi îmbunătățit.
- Îmbunătățirea performanțelor aplicației din punct de vedere al timpului de răspuns. Au fost efectuate teste pe o mașină cu sistem de operare OS X, procesor Intel i5 2.4 GHz și 8 GB memorie. S-a observat că sistemul de recunoaștere de vorbire este mai lent, încetinind astfel aplicația mai ales în etapa de înrolare, ce presupune transcrierea a cel puțin 10 clipuri audio.

Deși implementarea serviciului de recunoaștere de vorbire nu este subiectul acestei lucrări, am încercat îmbunătățirea sa. În implementarea inițială, transcrierea clipurilor se face secvențial. Mai mult, resursele necesare (model acustic, fonetic, lingvistic) și obiectul ce realizează transcrierea (de tip *recognizer*, specific CMU Sphinx) sunt încărcate în momentul în care se primesc clipurile audio. Deoarece încărcarea lor are o durată ce nu poate fi neglijată, resursele necesare precum și obiectul *recognizer* au fost încărcate în memorie la momentul pornirii aplicației și ținute în așteptare până la apariția unui client. A fost creat un grup de fire de execuție, egal cu numărul de nuclee al procesorului (4). La înrolarea unui utilizator, fiecare dintre cele 10 clipuri audio este

procesat în paralel, pe un fir de execuție separat, folosind unul dintre obiectele *recognizer* preîncărcate. Au fost obținute următoarele valori experimentale:

- o crearea unui obiect *recognizer*: 0.6 s;
- o transcrierea unui singur clip audio: 0.6 (valoare medie);
- o crearea modelului (amprente vocale) a unui utilizator: 1.42 s;
- o verificarea identității (amprente vocale) a unui utilizator: 0.63 s;

De asemenea, au fost obținute și valori ce compara cele două implementări ale serviciului de recunoaștere de vorbire:

- o transcrierea secvențială a 10 clipuri audio, folosind un singur obiect *recognizer*, încărcat la momentul execuției: 5.84 s;
- o transcrierea paralelă cu 4 fire de execuție a 10 clipuri audio, folosind obiecte *recognizer* preîncărcate: 3.86 s.

Se observă că implementarea îmbunătățită conferă un câștig în ceea ce privește timpul de execuție. Acesta este însă destul de mic din cauza faptului că sarcina de transcriere are o durată mică și relativ comparabilă cu încetinirea apărută din cauza comutării între firele de execuție.

Așadar, sunt investigate în continuare noi metode de accelerare a timpului de execuție al aplicației.

CAPITOLUL 5

APLICAȚIA DE DETECȚIE A CUVINTELOR CHEIE

Aplicația înglobează într-un proiect software sistemul de recunoaștere automată a vorbirii "Molina", dezvoltat cu ajutorul utilitarului Kaldi și prezentat pe larg în secțiunea 2.2. Scopul aplicației este de a efectua 3 sarcini principale:

- Obținerea transcrierii text a unui clip audio din baza de date Molina;
- Identificarea posibilităților apariții ale cuvintelor cheie în clipul audio, dată fiind în prealabil o listă a cuvintelor cheie căutate;
- Verificarea răspunsului unui sistem interactiv cu răspuns vocal; așa cum s-a prezentat în secțiunea 2.2.1, baza de date Molina este formată din clipuri audio cu un format special, ce imită răspunsul oferit de un sistem interactiv cu răspuns vocal. Clipurile sunt compuse din părți statice și părți dinamice, considerându-se că aparțin aceluiași tip de frază dacă

diferă numai prin părțile lor dinamice, părțile statice fiind identice. Astfel, verificarea răspunsului unui sistem interactiv cu răspuns vocal constă în determinarea tipului de frază cărui aparține un clip audio din baza de date, dată fiind în prealabil o listă a tipurilor de fraze posibile.

5.1 DESCRIERE ȘI ARHITECTURĂ

Aplicația a fost dezvoltată în limbajul C sub forma a două module, client și server, capabile să comunice prin protocolul TCP.

Modulul client permite utilizatorului să citească un fișier audio de pe disc și să-l trimită serverului sub forma unui flux audio. Modulul server primește fluxul audio de la client și realizează transcrierea lui text, identifica cuvinte cheie și tipul frazelor transcrise având în vedere mai multe categorii de fraze ce pot constitui ieșirea unui sistem IVR. Toate acestea vor fi afișate în timp real în consola serverului. Imaginea de mai jos reprezintă arhitectura aplicației:

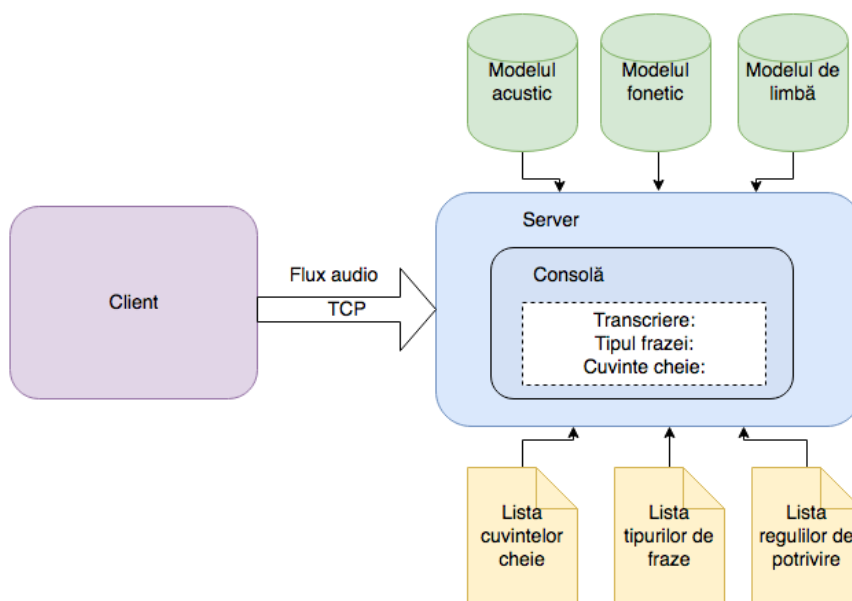


Figura 5.1 Arhitectura aplicației de detecție a cuvintelor cheie

Se pot observa cele două module (client și server). Serverul încarcă cele 3 modele (acustic, fonetic și lingvistic) corespunzătoare sistemului ASR Molina dar și 3 liste:

- o listă ce conține cuvintele cheie căutate;
- o listă ce conține tipurile de fraze ce pot constitui ieșirea unui sistem IVR.;
- o listă cu reguli care ajută la categorisirea frazelor în tipurile din lista precedentă.

Serverul astfel configurat va asculta pe un anumit port un flux de date audio și va scoate la ieșire transcrierea acestuia, tipul de fraze din care face parte precum și cuvintele cheie identificate.

5.2 FUNCȚIONALITATE DIN PERSPECTIVA UTILIZATORULUI

Mai întâi, pentru rularea modulului client este necesar ca Gstreamer 1.0 să fie instalat în sistemul de operare. Modulul dispune de un fișier de configurare în care se stabilește adresa IP a mașinii pe care se află modulul server, precum și portul pe care se va trimite fluxul audio.

Pentru rularea modulului server, suplimentar față de cazul precedent este necesar utilitarul Kaldi. Similar, acest modul dispune de un fișier de configurare în care se stabilesc căile către

resursele folosite, precum și adresa IP a mașinii pe care rulează, împreună cu portul pe care va aștepta conexiuni.

Aplicația rulează direct din linia de comandă, nefiind disponibilă o interfață grafică. Figura următoare ilustrează două terminale. În cel de sus rulează modulul server, în timp ce în cel de jos rulează modulul client.

```

kaldiuser@KaldiVM: ~/demoApp

TRANSCRIERE: member william robinson is not eligible for plan benefits
TIP FRAZA: member FIRST_NAME LAST_NAME is not eligible for plan benefits
Keywords: plan
Keywords: is not eligible

*****END STREAM*****

kaldiuser@KaldiVM: ~/demoApp
kaldiuser@KaldiVM:~/demoApp$ ./bin/client /home/kaldiuser/molinaDatabase/006/006_42_0040.wav
Running...
End of stream
Returned, stopping playback
Deleting pipeline
kaldiuser@KaldiVM:~/demoApp$

```

Figura 5.2 Rularea aplicației de detecție a cuvintelor cheie

Prima dată trebuie pornit modulul server. Acesta își va încărca resursele necesare și va aștepta o conexiune pe portul configurat. Modulul client se apelează având ca argument calea către un clip audio.

Se poate observa rezultatul afișat în urma rulării aplicației.

5.3 TEHNOLOGIA GSTREAMER

Gstreamer este o platformă de dezvoltare multimedia cu sursă deschisă care a fost implementată în mai multe limbaje de programare, iar în cazul de față fiind folosită biblioteca de funcții din Limbajul C. Aceasta funcționează pe majoritatea sistemelor de operare fiind folosită cu preponderență în sistemele Linux. Gstreamer este specializat în prelucrarea fluxurilor de date audio/video.

Un bloc Gstreamer are aspectul unui graf, ce poartă denumirea de *pipeline*, fiind format din elemente înlanțuite, unde datele curg dinspre elementul sursă (producător) către elementul final (consumator). De obicei, elementele implementează diverse filtre de semnal. Fiecare element conține un port de intrare în element, numit *sinkpad*, în timp ce portul prin care datele părăsesc elementul se numește *sourcepad*. De asemenea, fiecare bloc deține mai multe proprietăți ce pot fi configurate. Mai multe detalii despre un anumit element Gstreamer pot fi obținute prin comandă `gst-inspect-1.0 <nume-element>` tastată în terminal. Imaginea următoare ilustrează cele două blocuri *pipeline* ale aplicației, cel corespunzător modulului client și cel corespunzător modulului server.

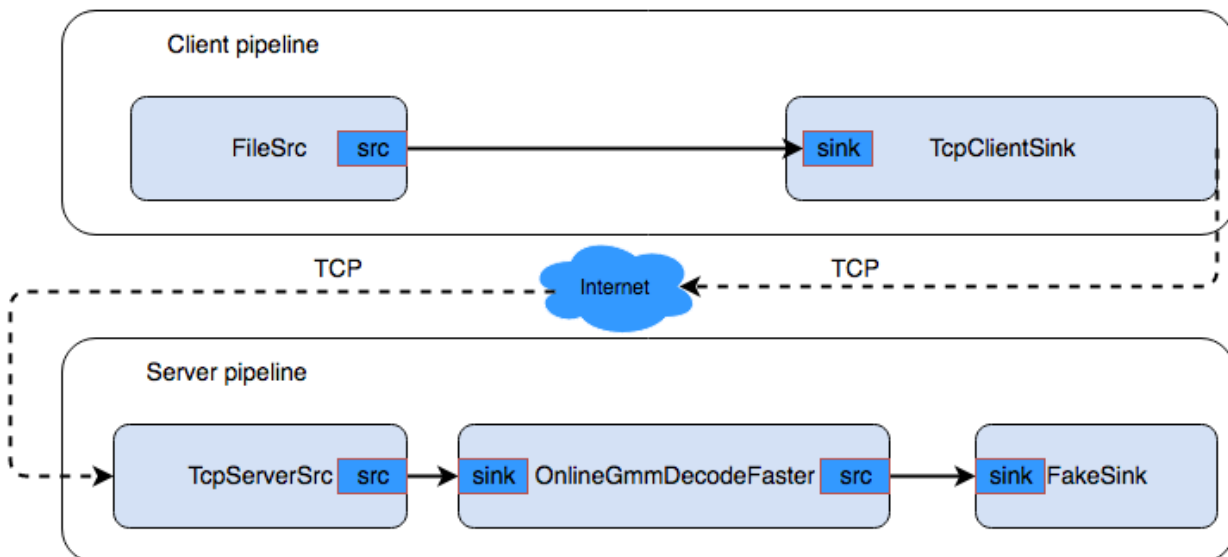


Figura 5.3 Pipeline-ul GStreamer

FileSrc este un element sursă, acesta citește un fișier de pe disc și îl poate trimite mai departe prin portul *src*. Cele mai importante proprietăți ale sale sunt: *location* (calea către fișierul ce va fi citit), *blocksize* (dimensiunea în octeți a unui bloc de date prelucrat), *name* (elementului îi poate fi atribuit un nume).

TcpClientSink este un element ce implementează un client capabil să trimită date peste rețea prin protocolul TCP. Acesta are un singur port, *sink*, prin care primește date de la elementele precedente din *pipeline* și le trimite mai departe spre o anumită adresă IP și un port specific. Cele mai importante proprietăți ale sale sunt *host* (adresa IP către care trimite datele) și *port* (portul pe care va trimite datele).

TcpServerSrc este un element ce implementează un server care poate primi date din rețea prin protocolul TCP. Prin portul *src* va trimite mai departe datele recepționate către alte elemente din *pipeline*. Similar cu elementul client, acestuia îi pot fi configurate proprietățile *host* și *port*.

OnlineGmmDecodeFaster este un element ce implementează un modul de transcriere în timp real, specific utilitarului Kaldi. Acesta dispune de două porturi, unul de intrare prin care primește fluxul audio (în formatul S16LE, 1 canal, 16 KHz) și altul de ieșire prin care furnizează transcrierea text. Proprietățile sale trebuie configurate stabilind calea către resursele unui sistem de recunoaștere automată a vorbirii dezvoltat cu Kaldi. Dintre proprietățile sale pot fi amintite: *model*, *fst*, *lda-mat*, *word-syms*, etc. Specific elementelor Gstreamer, **OnlineGmmDecodeFaster** are capacitatea de a emite semnale ce vor declanșa execuția unor funcții. În cazul de față, în momentul în care este obținută transcrierea text a unui flux audio, semnalul *hyp-word* este emis, fapt ce va declanșa apelul funcției cu următorul antet:

`void word_output (GstElement *asr, gchar *word, gpointer dată)`, unde *asr* este o referință la elementul **OnlineGmmDecodeFaster**, iar *word* reprezintă transcrierea text generată.

Deoarece în Gstreamer orice date produse trebuie consumate, elementul **FakeSink** are rolul de fals consumator, transcrierea text primită fiind redirectionată către linia de comandă.

5.4 DESCRIEREA ALGORITMULUI

Secțiunea curentă va descrie pașii executați de către aplicație pentru a îndeplini cele 3 sarcini propuse: transcrierea clipurilor, identificarea cuvintelor cheie, identificarea tipului de frază. Următoarea diagramă UML de stare ilustrează pașii ce vor fi prezentați:

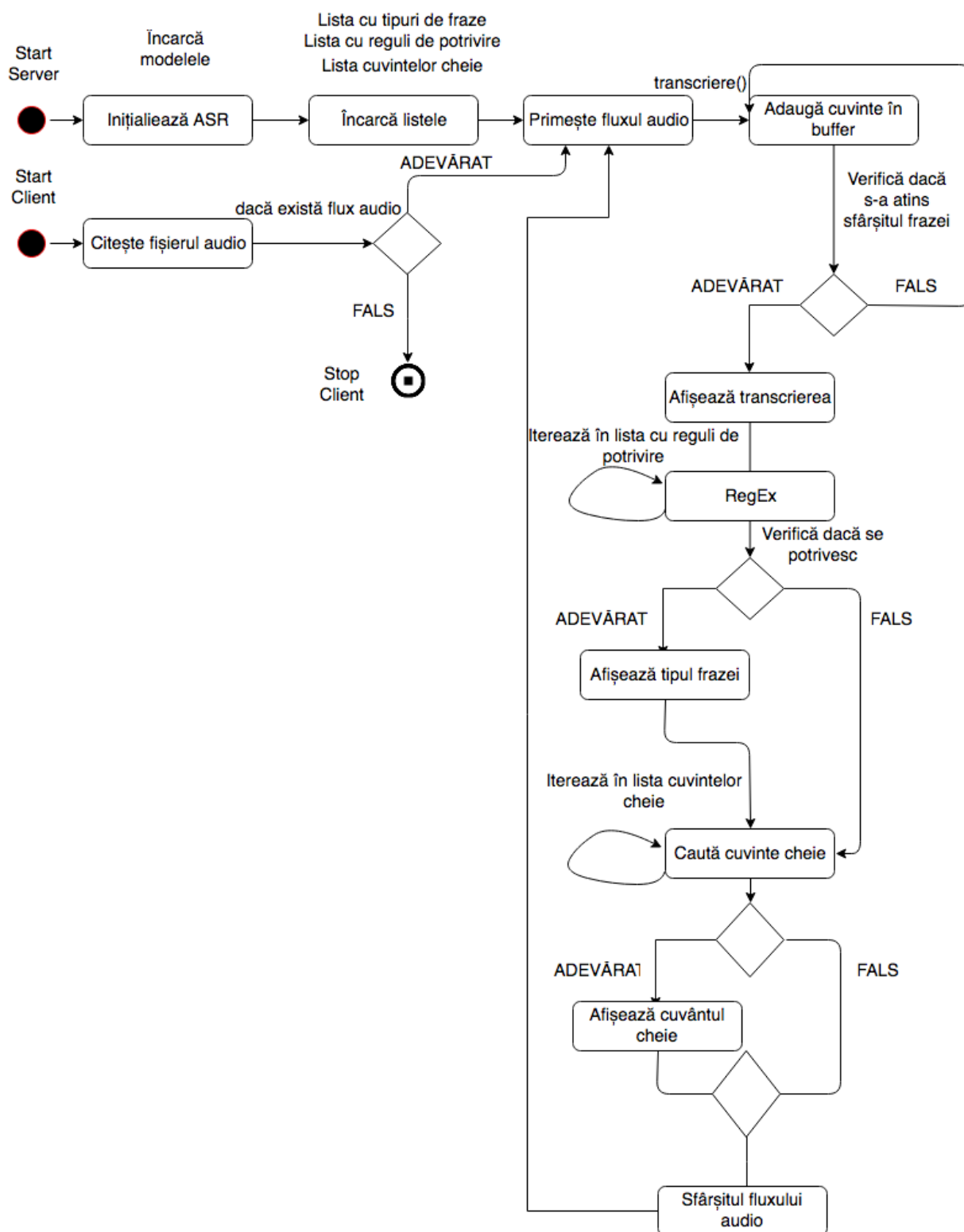


Figura 5.4 Diagramă UML de stare a aplicației de detecție a cuvintelor cheie

Modulul server

1. Citește fișierul de configurare.
2. Încarcă resursele ASR (modelul acustic, graful de decodare, lista cuvintelor, lista fonemelor, matricea cu caracteristici LDA-MLLT).

3. Încarcă listele (lista tipurilor de fraze, lista regulilor de potrivire, lista cuvintelor cheie).
4. Așteaptă conexiunea clientului. Pentru fiecare client conectat:
 - 4.1. Preia fluxul audio. Începe transcrierea fluxului. Pentru fiecare cuvânt transcris:
 - 4.1.1. Adaugă cuvântul într-un șir în care se va forma fraza.
 - 4.1.2. Verifică dacă sistemul de transcriere a returnat simbolul de delimitare a frazelor (<#s>).
 - 4.2. Când delimitatorul de frază este returnat, șirul va conține transcrierea integrală a clipului, aceasta fiind afișată.
 - 4.3. Se încearcă potrivirea șirului ce conține transcrierea integrală cu fiecare regulă de potrivire.
 - 4.4. Dacă o potrivire a fost realizată cu succes, tipul frazei este afișat.
 - 4.5. Fiecare cuvânt cheie este căutat în transcrierea text.
 - 4.6. Cuvintele cheie identificate sunt afișate la consolă.
 - 4.7. Șirul cu transcrierea text este resetat.
 - 4.8. Algoritmul se întoarce la pasul 4 (așteaptă o nouă conexiune din partea clientului).

Modulul client

1. Citește un fișier audio de la disc.
2. Dacă fișierul nu este gol, clientul îl trimite la server.
3. Atunci când tot șirul de octeți ai fișierului au fost trimiși, aplicația client se oprește.

5.5 VERIFICAREA UNUI SISTEM INTERACTIV CU RĂSPUNS VOCAL

La începutul capitolului s-a reamintit în ce constă verificarea unui sistem interactiv cu răspuns vocal. Această sarcină a fost îndeplinită folosind o bibliotecă C ce implementează **expresii regulate**. O expresie regulată este un șir de caractere text ce permite crearea unui model cu ajutorul căruia se vor face potriviri, se va manipula sau se va indentifica o secvență text.

Următorul exemplu constituie un **tip de frază** specific bazei de date Molina:

member FIRST_NAME LAST_NAME is not eligible for plan benefits plan name PLAN_NAME

Cuvintele scrise cu majuscule sunt părțile dinamice ale frazei, în timp ce restul cuvintelor scrise cu litere mici sunt părțile statice.

Folosind expresiile regulate, pentru acest tip de frază a fost creată următoarea **regulă de potrivire**:

member .+? is not eligible for plan benefits plan name .+? <#s>

Operatorul punct (".") va realiza potrivirea cu orice caracter, mai puțin caracterul de linie nouă. Operatorul "+" potrivește cu una sau mai multe apariții ale caracterului precedent. Operatorul "?" potrivește cu cel mult o apariție a caracterului precedent. Cei trei operatori luați împreună, ".+?", vor potrivi orice caracter de un număr nelimitat de ori, dar în același timp de cât mai puține ori posibil. Restul elementelor din expresia regulată (cuvântul "member" sau grupul de cuvinte "is not eligible for plan benefits") trebuie să se regăsească în mod identic pentru a fi posibilă potrivirea. Grupul de caractere <#s> este generat de către sistemul de transcriere a semnalului vocal în text și reprezintă o parte statică o frazei, marcând sfârșitul acesteia.

Considerând următoarele fraze ca fiind ieșirea unui sistem interactiv cu răspuns vocal, se va încerca potrivirea lor cu expresia regulată definită anterior:

1. *member John Smith is not eligible for plan benefits plan name United Healthcare Plan <#s>*
2. *member Lauren Thompson is eligible for plan benefits plan name United Healthcare Plus <#s>*

3. *member Nicole Brown is not eligible* <#s>

Prima frază se va potrivi cu expresia regulată deoarece părțile statice sunt identice, în timp ce părțile dinamice reprezintă o expansiune de caractere luată de oricâte ori. Cea de-a doua frază nu va realiza potrivirea deoarece lipsește cuvântul "not" din cea de-a doua parte statică a frazei. Nici ultima frază nu va realiza potrivirea deoarece lipsește partea statică "benefits for plan name" și ultima parte dinamică, "PLAN_NAME".

5.6 IMPLEMENTARE

Ambele module au fost implementate în limbajul C sub sistemul de operare Linux, folosind librăria GNU C (glib), specifică sistemelor Unix, ce asigură înaltă performanță și portabilitate.

Modulul server

Pentru a funcționa corespunzător, acesta are nevoie de următoarele:

- Lista cuvintelor cheie căutate (câte unul pe fiecare rând);
- Lista regulilor de potrivire a frazelor (câte una pe fiecare rând, similar cu cea explicată în secțiunea 5.5);
- Lista tipurilor de fraze (câte una pe fiecare rând, similar cu cea oferită ca exemplu în secțiunea 5.5);
- Resursele sistemului de recunoaștere: modelul acustic, graficul de decodare, lista cuvintelor, lista fonemelor, matricea cu caracteristici LDA-MLLT;
- Un fișier de configurare, format din mai multe secțiuni: secțiunea cailor spre liste și spre resursele sistemului de recunoaștere, secțiunea informațiilor referitoare la conexiune (IP, port), secțiune ce păstrează numărul cuvintelor din listă de cuvinte cheie, respectiv numărul tipurilor de fraze și al regulilor de potrivire.

Codul modulului este structurat într-un singur fișier ce conține funcția principală, împreună cu definițiile altor câteva funcții apelate de către aceasta.

Funcția *read_config_file()* este funcția care citește fișierul de configurare. Elementele din acesta sunt identificate în funcție de numele lor și numele secțiunii. Fie exemplu următor:

```
HOST_IP = g_key_file_get_string(gKeyFile,"connection","HOST_IP",&error);
HOST_PORT = g_key_file_get_integer(gKeyFile,"connection","HOST_PORT",&error);
```

Celor două variabile declarate global le vor fi atribuite valorile din fișierul de configurare ce corespund secțiunii *connection*, mai exact valorile elementelor *HOST_IP* și *HOST_PORT*.

Funcția *read_resources()* încarcă în memorie cele 3 liste necesare modulului. Conținutul listelor este citit și plasat în câte o structură de tip matrice.

Funcția *bus_call()* reprezintă un sistem de gestiune a mesajelor de eroare ce pot apărea într-un *pipeline* Gstreamer.

Funcția *search_for_keywords()* ia fiecare cuvânt din lista cuvintelor cheie și îl caută în transcrierea rezultată.

Funcția *word_output()* este o funcție *callback*, fiind executată atunci când se declanșează în funcția principală semnalul *hyp-word*, specific apariției unei noi transcrieri. Cât timp nu au fost generate caracterele ce marchează sfârșitul frazei (<#s>), fiecare nou cuvânt transcris este adăugat într-o listă în care se va forma fraza. Când fraza a fost transcrisă integral, se va încerca determinarea tipului ei prin apelul funcției *test_matching()*.

Funcția `int test_matching(gchar *buffer)` primește ca argument transcrierea integrală a frazei. Prin folosirea bibliotecii `regex.h` se va încerca potrivirea cu fiecare regulă de potrivire din fișier, în acest fel fiind determinat tipul frazei. Instrucțiunea `regex = g_regex_new (pattern, 0, 0, &err);` va compila regulă de potrivire ("pattern"), în timp ce instrucțiunea `g_regex_match (regex, buffer, 0, &matchInfo);` va testa potrivirea cu transcrierea frazei regăsită în variabilă `buffer`.

Funcția `main()` este funcția principală a programului. Aceasta începe prin declararea unor obiecte specifice librăriei Gstreamer. Astfel, cele 3 elemente Gstreamer ale serverului sunt declarate, împreună cu `pipeline-ul`, un container în care acesta va rula și elementul `bus` de gestiune a mesajelor de eroare.

```
gint main (gint  argc, gchar *argv[])
{
    GMainLoop *loop;
    GstElement *pipeline, *src, *asr, *fakesink;
    GstBus *bus;
```

Se va inițializa lista ce va conține transcrierea frazelor generată de sistemul ASR.

```
transcriptionBuffer= g_string_new(NULL);
```

Se vor încărca cele 3 liste necesare: cuvintele cheie, regulile de transcriere și tipurile de fraze.

```
read_config_file();
phraseTypes=read_resources(PHRASE_TYPES_PATH, PHRASE_TYPES_NR);
matchRules=read_resources(MATCH_RULES_PATH, PHRASE_TYPES_NR);
keywords=read_resources(KEYWORDS_PATH, KEYWORDS_NR);
```

Containerul Gstreamer este inițializat.

```
gst_init (&argc, &argv);
loop = g_main_loop_new (NULL, FALSE);
```

Este creat un obiect de tip pipeline și este inițializat obiectul bus.

```
pipeline = gst_pipeline_new ("pipeline");
bus = gst_pipeline_get_bus (GST_PIPELINE (pipeline));
gst_bus_add_watch (bus, bus_call, loop);
gst_object_unref (bus);
```

Sunt create cele 3 elemente Gstreamer descrise în secțiunea 5.3.

```
src = gst_element_factory_make ("tcpserversrc", "source");
asr=gst_element_factory_make("onlinemmdecodefaster", "asr");
fakesink=gst_element_factory_make("fakesink","fakesink");
```

Sunt configurate proprietățile elementelor.

```

g_object_set (G_OBJECT (src), "host", HOST_IP, NULL);
g_object_set (G_OBJECT (src), "port", HOST_PORT , NULL);

g_object_set (G_OBJECT (asr), "fst", FST_PATH, NULL);
g_object_set (G_OBJECT (asr), "lda-mat", LDA_PATH, NULL);
g_object_set (G_OBJECT (asr), "model", AM_PATH, NULL);
g_object_set (G_OBJECT (asr), "word-syms", WORDS_PATH, NULL);
g_object_set (G_OBJECT (asr), "silence-phones", SILENCE_PHONES, NULL);
g_object_set (G_OBJECT (asr), "max-active", MAX_ACTIVE, NULL);
g_object_set (G_OBJECT (asr), "beam", BEAM, NULL);
g_object_set (G_OBJECT (asr), "acoustic-scale", ACOUSTIC_SCALE, NULL);
g_object_set (G_OBJECT (asr), "inter-utt-sil", INTER_UTT_SIL, NULL);
g_object_set (G_OBJECT (asr), "silent", FALSE, NULL);
g_object_set (G_OBJECT (fakesink), "dump", FALSE, NULL);

```

Elementele sunt adăugate în pipeline și sunt create legăturile dintre acestea.

```

gst_bin_add_many (GST_BIN (pipeline), src, asr, fakesink, NULL);
gst_element_link_many (src, asr, fakesink, NULL);

```

Semnalul generat de elementul ce realizează transcrierea (asr) este conectat la funcția ce va fi apelată la declanșarea sa.

```

g_signal_connect (asr, "hyp-word", G_CALLBACK (word_output), NULL);

```

Pipeline-ul este pornit într-o buclă infinită și permite prelucrarea a oricâte fluxuri audio primite secvențial.

```

while(1)
{
    gst_element_set_state (pipeline, GST_STATE_PLAYING);
    g_main_loop_run (loop);
    gst_element_set_state (pipeline, GST_STATE_NULL);
}
gst_object_unref (GST_OBJECT (pipeline));
return 0;
}

```

Modulul client a fost creat într-un mod similar cu modulul server. Elementele Gstreamer ale modulului client au fost descrise în secțiunea 5.3.

5.7 CONCLUZII ȘI DEZVOLTĂRI ULTERIOARE

Aplicația demonstrativă oferă funcționalitățile pentru care a fost concepută. Pentru fluxurile audio ce imită răspunsul unui sistem IVR sunt realizate transcrierile, sunt identificate cuvintele cheie și este determinat tipul de frază.

În momentul actual, aplicația folosește pentru transcrierea fluxurilor audio modele de tip HMM-GMM antrenate în Kaldi. Pentru dezvoltările ulterioare ale proiectului se urmărește integrarea în aplicație a modelelor acustice antrenate cu rețele neuronale. Astfel de modele au fost deja obținute cu succes, iar acuratețea transcrierilor este superioară față de modelele bazate pe GMM.

Tot în cadrul dezvoltărilor viitoare se va avea în vedere crearea unei interfețe standard a serverului în raport cu aplicațiile client. Cunoscând acesta interfață, acestea vor fi capabile să se conecteze la server pentru a obține atât transcrieri în timp real ale fluxurilor ce provin de la microfon sau alte surse, cât și a fișiere audio ce vor fi încărcate pe server. Totodată, foarte importantă este paralelizarea aplicației, fiind posibilă manipularea cererilor venite simultan de la mai mulți clienți.

Pentru sarcina de identificare a cuvintelor cheie este important de menționat momentul de timp al apariției acestora în raport cu momentul de început al fluxului audio. Vor fi aduse modificări ale elementului *Gstreamer OnlineGmmDecodeFaster*, în vederea obținerii acestei facilități.

CONCLUZII FINALE

Această lucrare a avut ca scop descrierea implementării unor aplicații software ce includ sisteme de recunoaștere a vorbirii și detecție a cuvintelor cheie. S-a dorit evaluarea acurateții sistemelor și a performanțelor acestora.

Capitolul 1 a reprezentat o introducere teoretică în domeniul tehnologiei vorbirii. Au prezentate detalii despre producerea și percepția semnalului vocal, extragerea parametrilor acustici, precum și resursele necesare unui sistem de recunoaștere automată a vorbirii. Au fost sumarizate informații despre utilitățile folosite și metricile de evaluare. De asemenea, au fost prezentate principalele metode de detecție a cuvintelor cheie, precum și metricile de evaluare ale unui astfel de sistem.

Capitolul 2 s-a axat pe construcția unor sisteme de recunoaștere automată a vorbirii.

În prima secțiune a capitolului, s-a descris crearea unui sistem de recunoaștere a cifrelor în limba română. Au fost create modelele necesare și s-au efectuat mai multe teste și optimizări. S-a constatat faptul că un model acustic dependent de context este mult mai performant decât unul independent de context. De asemenea, sistemul independent de vorbitor este capabil să realizeze transcrieri corecte chiar și în cazul vorbitorilor necunoscuți. Prin optimizări succesive s-a constatat că modelul optim este cel cu 100 senone și 128 densități gaussiene. S-au evaluat ratele de eroare la nivel de cuvânt pentru fiecare vorbitor în parte și s-au constatat sursele ce pot provoca erori. Prin replicarea experimentelor folosind ambele utilitare de recunoaștere automată a vorbirii, a fost realizată o comparație directă între acestea. S-a dovedit faptul că utilitarul Kaldi oferă o acuratețe mai bună față de CMU Sphinx.

În a doua secțiune a capitolului s-a prezentat crearea unui sistem de recunoaștere automată a vorbirii cu vocabular redus în limba engleză. Au fost create modelele necesare și s-au efectuat teste de acuratețe, dar și teste privind consumul resurselor. În cadrul celor din urmă, s-a constatat că un model de limba de tip gramatică este mai eficient decât un model probabilistic de tip n-gram.

Capitolul 3 a prezentat crearea unui sistem de detecție a cuvintelor cheie folosind un modul din utilitarul Kaldi. S-a efectuat o analiză asupra consumului de resurse cauzat în special de către modelul de limbă. S-a constatat că un model de limbă bazat pe reguli, ce conține numai cuvintele cheie căutate este mult mai eficient decât un model n-gram. Apoi, baza de date de testare a fost etichetată, fiind marcate aparițiile cuvintelor cheie. S-a evaluat sistemul de detecție obținut. Din păcate, nu au fost duse la capăt toate scenariile de testare propuse. Acest lucru nu a fost posibil din cauza problemelor apărute la utilizarea aplicațiilor de evaluare.

Capitolul 4 s-a concentrat pe implementarea software a serviciului de autentificare prin voce. A fost descrisă funcționalitatea aplicației și arhitectura. Au fost oferite detalii de implementare a aplicației client, dar și a serviciului de autentificare și a bazei de date. Aplicația finală și-a atins scopul propus, fiind posibilă înrolarea și autentificarea utilizatorilor pe bază de recunoaștere de vorbire și vorbitor. S-au constatat însă unele limitări cauzate de serviciul de transcriere de cifre, fapt ce determină creșterea timpului de așteptare la crearea contului. Au fost propuse și analizate noi metode de implementare ale serviciului de transcriere a cifrelor. În continuare, acest aspect se caută a fi îmbunătățit. Nu în ultimul rând, a fost descoperită o eventuală problemă de securitate la autentificare. Aceasta a fost descrisă pe larg în secțiunea corespunzătoare modurilor de utilizare ale aplicației. Au fost propuse câteva metode de rezolvare a problemei.

Capitolul 5 a descris implementarea unei aplicații software de detecție a cuvintelor cheie. Aceasta a fost realizată cu succes, îndeplinind și funcționalități de transcriere text a fluxurilor audio primite, precum și de identificare a tipului de frază corespunzător. Totuși, aplicația este una demonstrativă, ce nu oferă foarte multe opțiuni utilizatorului. Modulul server suporta un singur client la un anumit moment de timp. Modulul client este și el unul limitat, permite numai citirea unui fișier audio de pe disc și trimiterea lui la server, nefiind posibilă transcrierea unui flux captat de la microfon. Totuși, aplicația este funcțională în limitele descrise și poate fi un bun punct de plecare în cazul unor viitoare dezvoltări.

CONTRIBUȚII PERSONALE

Într-o mare măsură, lucrarea de față a reușit să îndeplinească obiectivele propuse. Contribuțiile personale ale autorului ce au condus la atingerea obiectivelor pot fi sumarizate după cum urmează:

1. Sistemele de recunoaștere automată a vorbirii

1.1 *Sistemul Rodigits: recunoașterea cifrelor în limba română*

- 1.1.1 Crearea modelelor (acustic, fonetic, lingvistic);
- 1.2.2 Crearea sistemului dependent de vorbitor;
- 1.2.3 Crearea sistemului independent de vorbitor;
- 1.2.4 Evaluarea acurateții și optimizarea sistemului;
- 1.2.5 Analiza rezultatelor, depistarea posibilelor surse de eroare;
- 1.2.6 Replicarea proiectului în vederea realizării comparației celor două utilitare.

1.2 *Sistemul Molina: recunoașterea vorbirii în limba engleză cu vocabular redus*

- 1.2.1 Crearea bazei de date audio de evaluare;
- 1.2.2 Crearea modelelor (acustic, fonetic, lingvistic);
- 1.2.3 Evaluarea acurateții sistemului;
- 1.2.4 Evaluarea consumului de resurse la momentul decodării.

2. Sistemul de detecție a cuvintelor cheie

- 2.1 Analiza limitării generată de consumul resurselor în cazul modelului de limba 3-gram
- 2.2 Testarea unui model de limbă alternativ
- 2.3 Etichetarea manuală a cuvintelor cheie din baza de date de evaluare
- 2.4 Obținerea rezultatelor în urma detecției
- 2.5 Evaluarea sistemului de detecție a cuvintelor cheie

3. Serviciul web de autentificare prin voce

- 3.1 Implementarea aplicației client
- 3.2 Implementarea bazei de date
- 3.3 Implementarea serviciului de autentificare
- 3.4 Analiza limitării cauzată de timpul de execuție al serviciului de recunoaștere de cifre

4. Aplicația de detecție a cuvintelor cheie

- 4.1 Implementarea modulului server
- 4.2 Implementarea modulului client

BIBLIOGRAFIE

- [1] Horia Cucu, "Proiect de cercetare-dezvoltare în Tehnologia Vorbirii", Indrumar de proiect, Universitatea "Politehnica" București.
- [2] Keith Allan, "Linguistic Meaning", Routledge, 2014.
- [3] Anca Popescu, Carmen Pătrașcu, Interfețe Om-Mașină, Laborator 3, http://lpsv.pub.ro/images/IOM/CURS_-_IOM/IOM_3_AP.pdf
- [4] The mel frequency scale and coefficients, http://kom.aau.dk/group/04gr742/pdf/MFCC_worksheet.pdf
- [5] Jordan Critten, Parker Evans, Speaker recognition, ECE 576 FINAL PROJECT, https://courses.cit.cornell.edu/ece576/FinalProjects/f2008/pae26_jsc59/pae26_jsc59/
- [6] Lindasalwa Muda, Mumtaj Begam and I. Elamvazuthi, "Voice Recognition Algorithms using Mel Frequency Cepstral Coefficient (MFCC) and Dynamic Time Warping (DTW) Techniques", JOURNAL OF COMPUTING, VOLUME 2, ISSUE 3, MARCH 2010, ISSN 2151-9617
- [7] Roberto Togneri, Daniel Pullella, „An overview of speaker identification: Accuracy and Robustness Issues”, IEEE Circuits and systems magazine, 2011
- [8] Anca Popescu, Carmen Pătrașcu, Interfețe Om-Mașină, Laborator 3, http://lpsv.pub.ro/images/IOM/Lab_02_Analiza_de_timp_scurt_a_semnalului_vocal.pdf
- [9] Andi Buzo, "Automatic speech recognition over mobile telecommunication channels", PhD Thesis, Universitatea "Politehnica" București, Oct 2011.
- [10] Huang, X., Acero, A., Wuen-Hon, H., Spoken Language Processing – A Guide to Theory, Algorithm, and System Development, Prentice Hall, 2001.
- [11], Mixture Model, http://www.en.wikipedia.org/Mixture_model
- [12], CMU Sphinx architecture, <http://www.di.uniba.it/~loglisci/FELIPE/>
- [13], Daniel Povey, "The Kaldi Speech Recognition Toolkit", Microsoft Research, 2011
- [14], [14, Josh James, Data never sleeps 2.0, <https://www.domo.com/blog/2014/04/data-never-sleeps-2-0/>
- [15], KWS 15 Keyword Search Evaluation Plan
<https://www.nist.gov/sites/default/files/documents/itl/iad/mig/KWS15-evalplan-v05.pdf>
- [16] Horia Cucu, "Proiect de cercetare-dezvoltare în Tehnologia Vorbirii", îndrumar de proiect, Universitatea "Politehnica" București.
- [17] Molina Healthcare IVR system,
http://www.molinahealthcare.com/providers/wi/medicaid/forms/PDF/forms_WI_10_IVR_reference_guide.pdf
- [18], Anthony Rousseau, Paul Deléglise, Yannick Estève, "TED-LIUM: an Automatic Speech Recognition dedicated corpus", Laboratoire Informatique de l'Université du Maine (LIUM), 2012
- [19], Ted Conference, <https://www.ted.com/>
- [20] Alina Bănică, "Serviciu web de autentificare prin voce. Sisteme de recunoaștere de cifre conectate și verificare de vorbitor", Proiect de Licență, Universitatea "Politehnica" București, 2015
- [22] Boris Smus, "Web Audio API", O'Reilly, 2013
- [23] RecorderJS, <https://github.com/mattdiamond/Recorderjs>