

Motto:

**“ Each domain has its Holy Grail, and
automatic speech recognition is ours “**

J. Flanagan

ASRSV – curs 9

Recunoasterea Semnalului Vocal

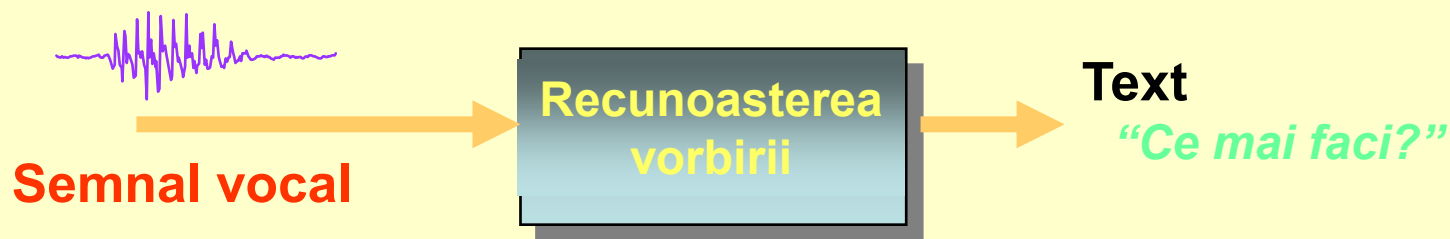
http://home.dei.polimi.it/matteucc/Clustering/tutorial_html/index.html

<https://www.ibm.com/watson/services/speech-to-text/>

<https://medium.com/descript/a-brief-history-of-asr-automatic-speech-recognition-b8f338d4c0e5>

Cuprins

1. Comunicarea orală om-mășină
2. Recunoașterea vorbirii
3. Metrică în spațiul acustic
4. Metode de recunoașterea vorbirii
5. Metoda alinierii dinamice (DTW)

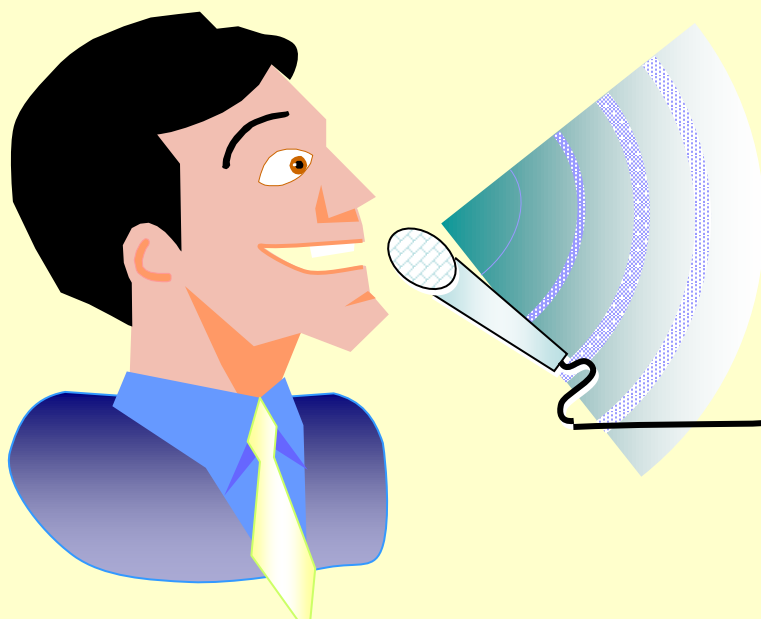


1.COMUNICAREA ORALA OM-MASINA

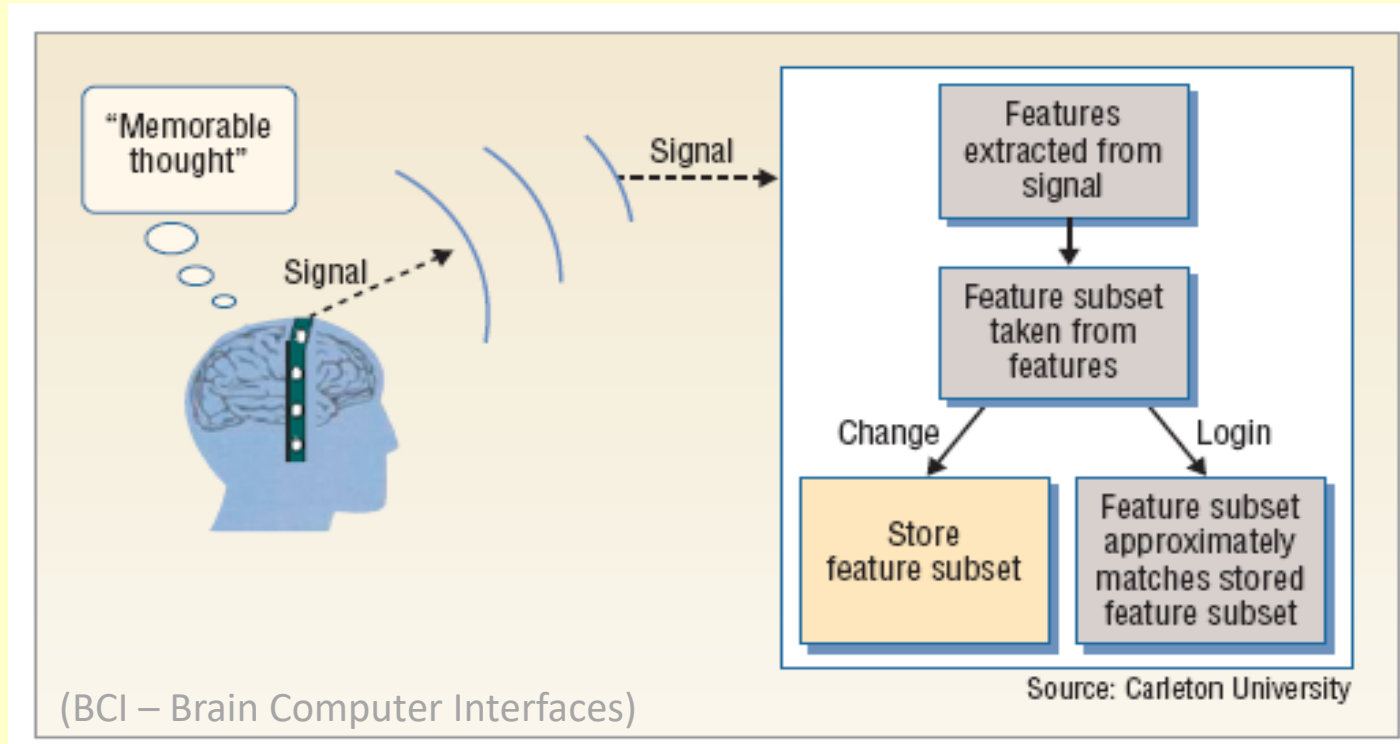
- Tehnologia care permite calculatoarelor să identifice și să proceseze vorbirea umană, transformând-o în text sau comenzi executabile (speech to text) si sa transmita mesaje vocale ca raspuns (text to speech)

Impact și Utilizare

- Accesibilitate: Pentru persoanele cu dizabilități
- Productivitate: Dictare, transcrieri automate
- Interfață naturală: Dispozitive inteligente, mașini
- Eficiență: Suport pentru clienți, analiză de interacțiuni



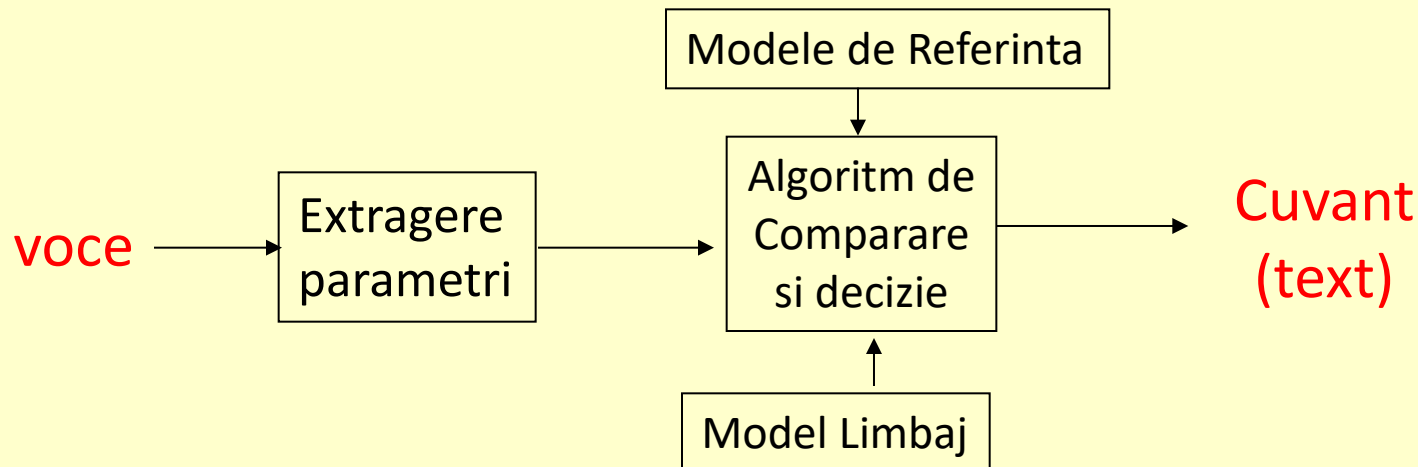
Carleton University's proposed BCI-based biometric system



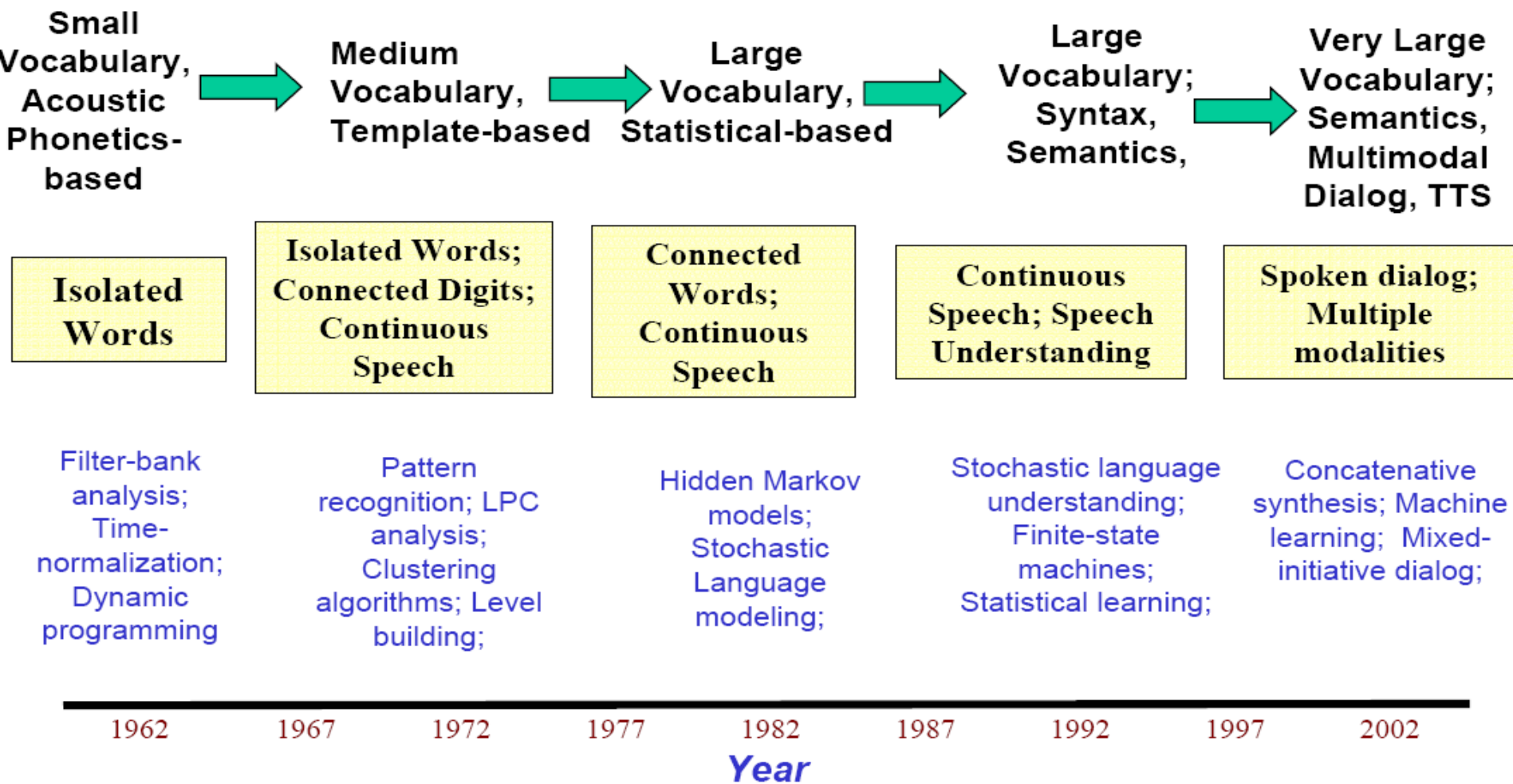
- Subjects use specific thoughts as passwords (called **pass-thoughts**). When someone tries to access a protected computer system or building, they think of their pass-thought. A headpiece with electrodes records the brain signals. The system extracts the signal's features for computer processing, which includes identification of the feature subset that best and most consistently represents the pass-thought.
- The biometric system then compares the subset to those recorded for authorized users.

2. Recunoasterea vorbirii

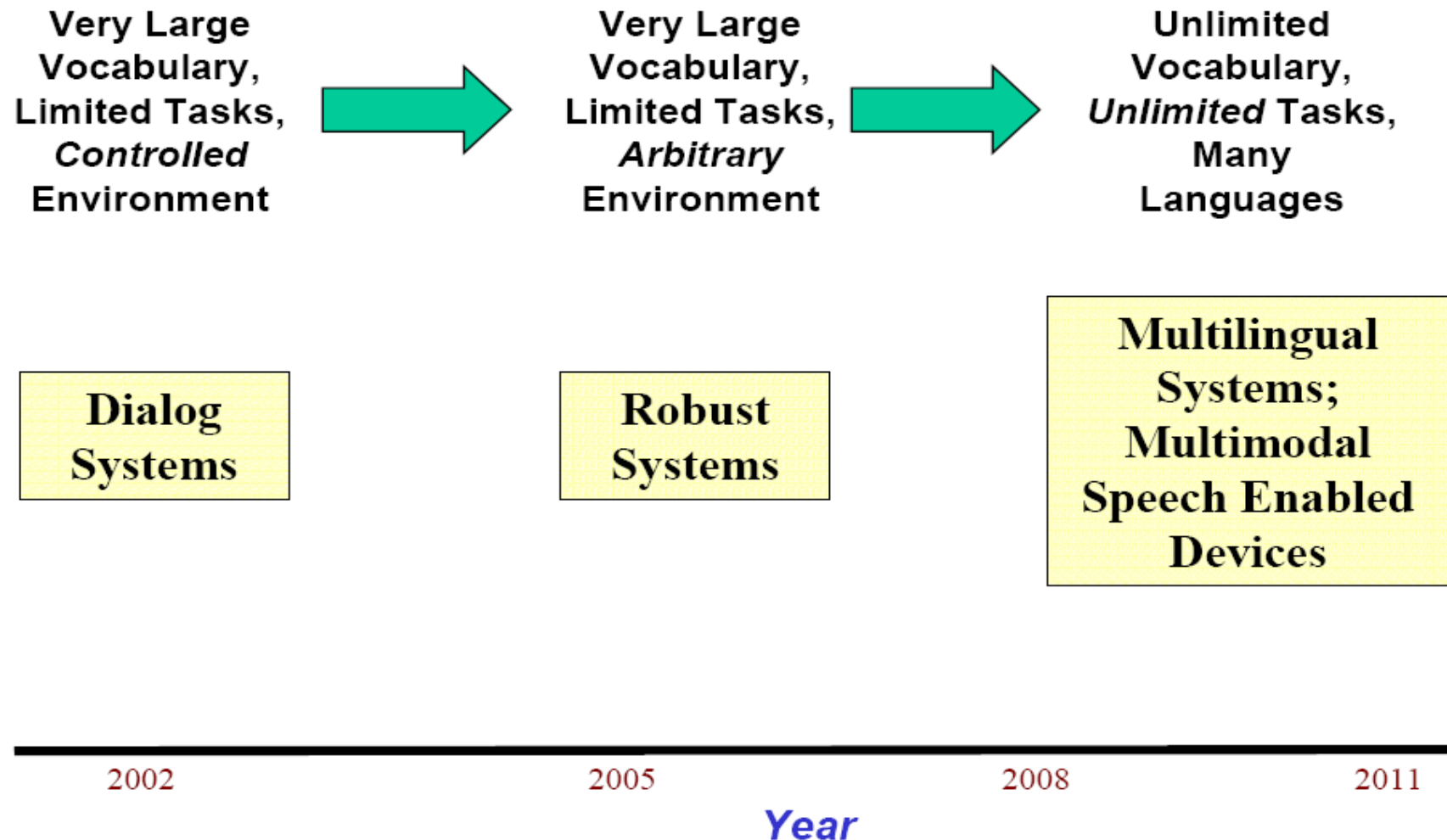
- Recunoasterea vorbirii (RV) este *baza interactiunii om-calculator* folosind vocea (speech to text)
- Un sistem ASR este o piatră de temelie a inteligenței artificiale
- Recunoasterea vorbirii poate fi in diferite contexte
 - Dependenta/ independenta de vorbitor
 - Cuvintele rostite izolat sau continuu
 - Vocabular redus sau mare
 - In mediu linistit sau zgomotos



Milestones in Speech and Multimodal Technology Research



Future of Speech Recognition Technologies



Voice reply to customer



"What number did you want to call?"



Speech



Automatic Speech Recognition

Words spoken

"I dialed a wrong number"

Spoken Language Understanding

Meaning

"Billing credit"

Dialog Management

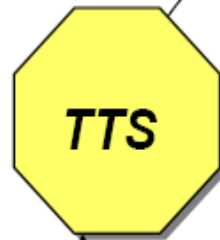
Action

Spoken Language Generation

Words: What's next?

"Determine correct number"

Text-to-Speech Synthesis



1. Clasificarea sistemelor de recunoastere automata a SV

Factorii de complexitate care determina dificultatea recunoasterii pot fi clasificati astfel:

- Dependentă/independentă de vorbitor
- Cuvintele rostite izolat sau continuu
- Extensia vocabularului
- Limbajul
- Mediul ambiant

- *Dependența sau independenta de vorbitor*

Recunoasterea semnalului vocal pentru vocea unui singur vorbitor nu este o problema chiar usoara datorita variabilitatii intralocutor inerente procesului de vorbire. S-a constatat ca variabilitatea este legata de :

- variatii de debit verbal
- emotie
- stress, oboseala, raceala etc.

Recunoasterea independenta de vorbitor (multilocutor) este o problema dificila, in masura in care la **variabilitatea intralocutor** se adauga cea **interlocutor** legata de caracteristicile anatomice, a habitului lingvistic sau accentelor regionale etc.

- *Cuvintele sunt rostite izolat sau continuu*

Daca vorbitorul (vorbitorii) marcheaza cu o pauza dupa fiecare cuvant pronuntat, complexitatea problemei se reduce, pentru ca granitele intre cuvinte sunt bine definite (spre deosebire de vorbirea continua) ceea ce limiteaza mult combinatiile pentru a accede la un lexic, de exemplu.

- cuvintele pot fi considerate ca entitati globale nu ca o succesiune de entitati elementare.

- **Dimensiunea vocabularului** - Este mai dificil de a lucra pe un vocabular foarte extins (mii, zeci de mii de cuvinte) decat pe un vocabular restrans (cateva zeci de cuvinte). Dimensiunea vocabularului este un parametru insuficient, un grup de cuvinte foarte diferite unele de altele fiind mult mai usor de tratat decat cuvinte *apropiate din punct de vedere fonetic* (ex. moale, foale, oale, boale, toale).
- **Limbajul** - Daca la tratarea semnalului vocal se tine cont si de sintaxa limbajului folosit de locutor (si eventual si de semantica si pragmatica in cadrul unei aplicatii) va fi mult mai usor pentru un limbaj rigid, foarte constrins decat pentru unul suplu, natural.
 - Exista metode cantitative de studiu al acestui tip de complexitate in legatura cu dimensiunea vocabularului folosind notiuni ca: *factor de ramificatie mediu, perplexitate, vocabular echivalent* etc.
- **Mediul ambiant** - Independent de cele enumerate anterior, mediul acustic si conditiile de captare a sunetului constituie un factor important: prezenta zgomotului (chiar stationar) degradeaza in general foarte mult performantele sistemelor de recunoastere. Mai mult un zgomot puternic atrage dupa sine o variabilitate sporita a locutorului (efect Lombard).

MEDIU	Zgomot de reverberatie corelat cu vorbirea Zgomot de reflexie necorelat Zgomot aditiv (stationar/nestationar)
VORBITOR	Atributele vorbitorului:- dialect, sex, vârsta Modul de vorbire : - zgomot de respiratie si buze - stress - efect LOMBARD - debit verbal - nivel - frecventa fundamentala - cooperativitate
ECHIPAMENT DE INTRARE	Microfon (transmitator) Distanța de microfon Filtre Sistemul de transmisie (distorsiuni, zgomot, ecou) Echipament de înregistrare

Sistemele de recunoaștere se pot clasifica în ordinea dificultății crescătoare astfel :

Sistemele de recunoaștere se pot clasifica în ordinea dificultății crescătoare astfel :

- **Recunoașterea cuvintelor izolate aparținând unui vocabular limitat (<100 cuvinte) și sistemul adaptat la un vorbitor dat (sistem monolocator)**
- **Recunoașterea cuvintelor izolate aparținând unui vocabular extins (mii de cuvinte) și sistemul adaptat la un vorbitor dat (sistem monolocator)**
- **Recunoașterea cuvintelor izolate sistemul fiind independent de vorbitor (sistem multilocutor)**
- **Recunoașterea cuvintelor în lanț dintr-un vocabular limitat (ex. cele 10 cifre) , dar pronunțate în orice ordine fără pauză între ele.**
- **Recunoașterea de fraze scurte bazate pe un vocabular limitat, monolocator.**
- **Recunoașterea de fraze scurte bazate pe un vocabular limitat , multilocutor.**
- **Recunoașterea vorbirii continue mono/multilocutor**
- **Înțelegerea vorbirii este o etapă suplimentară care cere sistemului o acțiune adecvată conținutului mesajului vocal**

O problemă foarte apropiată de recunoașterea vorbirii este *recunoașterea locutorului*.

Sisteme de recunoașterea vorbirii

Recunoașterea și înțelegerea

- Sunt două obiective diferite care se pot atașa unui sistem. Recunoașterea în sens strict conduce la o aplicație de tip dictare automată (mașina de scris fonetică), în timp ce *înțelegerea caută să ajungă la o semnificație* a enunțului vorbit (ca de ex. în cazul dialogului om-mașină). Această distincție nu se remarcă decât în cazul sistemelor destul de evolute. Pentru un limbaj de comandă simplu, prin cuvinte izolate sau intrări de date recunoașterea = înțelegere.

Sisteme auto-organizate și bazate pe cunoștințe

- Un sistem poate folosi metode matematice pentru a extrage date reprezentative (esantioane de învățare) structuri și parametri pe care îi va folosi pentru efectuarea recunoașterii. Este altfel spus *“auto-organizat”*, dar aceasta nu înseamnă că este total ignorant la cunoștințele umane (care au stat la baza definirii unei metode matematice pe un tip specific de date), ci le folosește implicit. (DTW, HMM)
- O altă abordare *utilizează explicit în sistemele de cunoștințe experți umani* care au studiat vorbirea la mai multe nivele: acustic, fonetic, lingvistic - *sisteme bazate pe cunoștințe*. Astfel de sisteme sunt concepute să realizeze *prelucrări de tipul inteligentă artificială* cooperând cu sistemele clasice de tratare a semnalului și de recunoaștere a formelor.

Recunoașterea globală și recunoașterea analitică

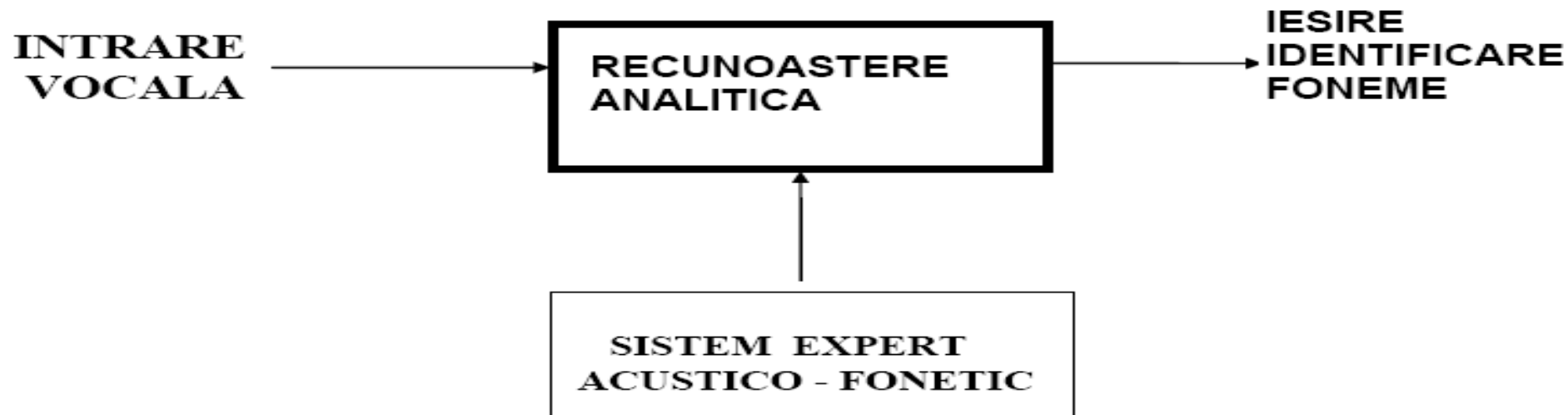
- Recunoașterea globală consta in compararea entitatilor de tip cuvint (de identificat) intre ele fara a le descompune in alte subentitati, in acest caz avem recunoasterea globala.



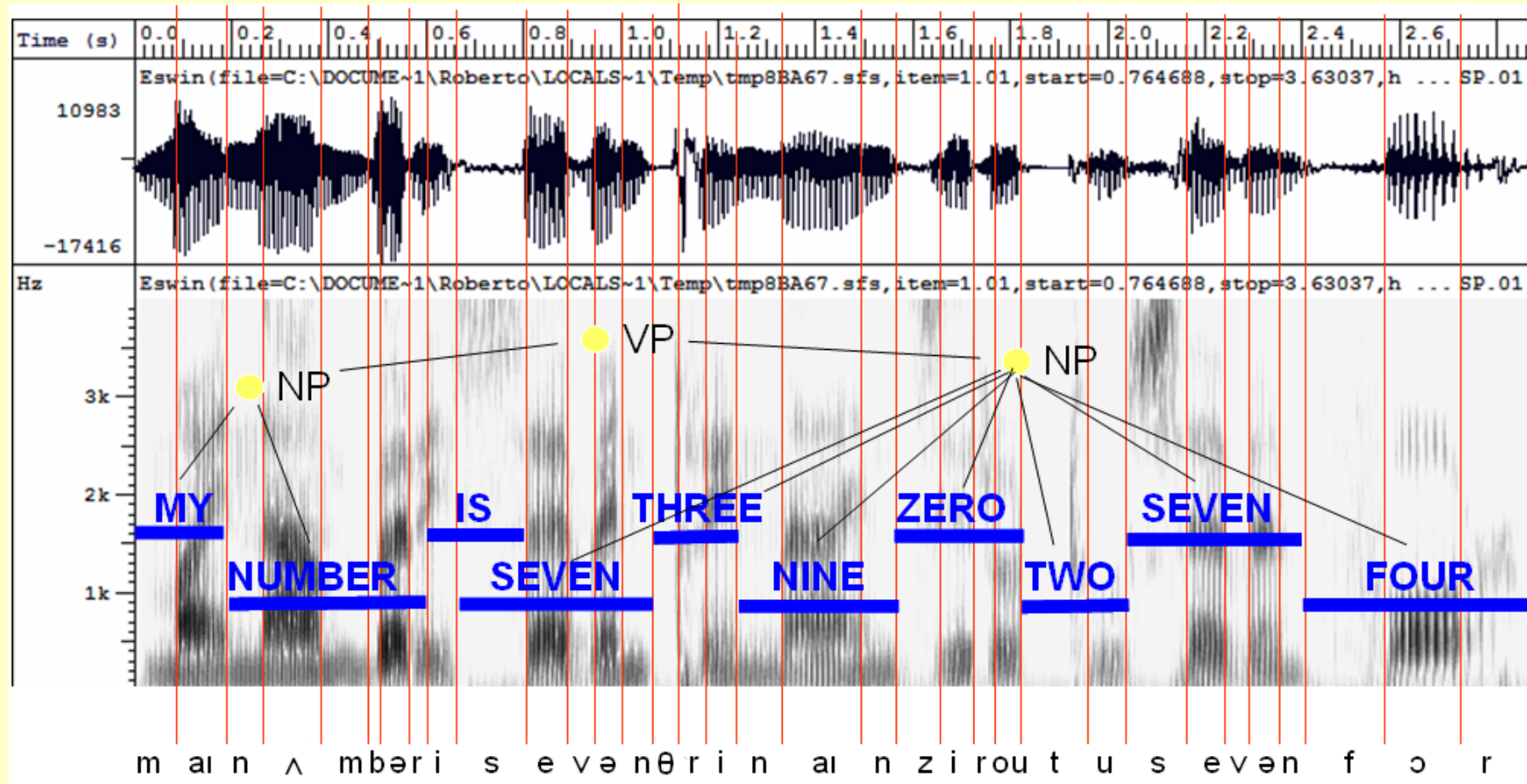
Sistem de recunoaștere globală.

- **Recunoșterea analitică** consta din metode care folosesc într-un *stadiu intermediar entitati elementare* (foneme, difoni, demisilabe) care permit descrierea entitatilor de nivel superior

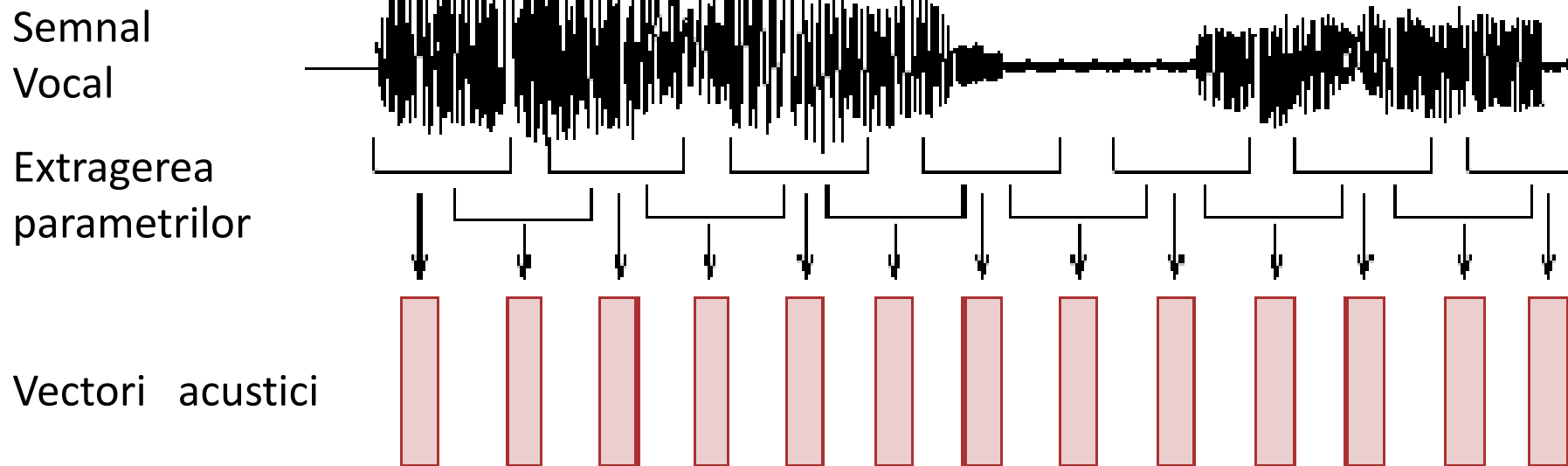
- anumite entitati elementare ale recunoasterii analitice nu au nevoie de o definitie fonetica, ele putand fi definite “ad-hoc” (definite de exemplu prin cuantizarea vectoriala).



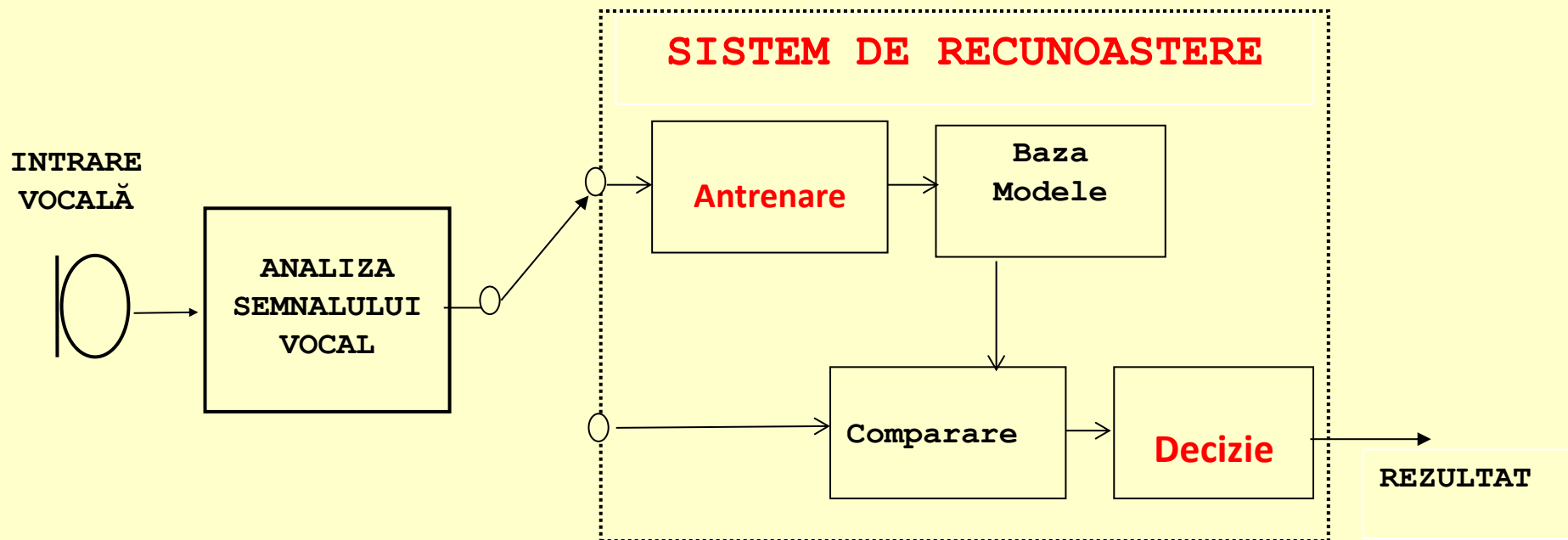
Sistem de recunoaștere analitică.



- Identificarea fonemelor individuale
- Identificarea cuvintelor
- Identificarea structurii și/sau semnificației propozițiilor
- Interpretarea caracteristicilor prozodice (tonul, intensitatea, durata)



- *Nivelul de baza pentru recunoasterea vorbirii este nivelul acustic*
- parametrii care caracterizeaza un spectru vocal se numeste *vector acustic/spectral*
- se fixeaza o metrica in spatiul marginit de un anumit tip de parametri
- distantele intre vectorii spectrali sunt mai mult sau mai putin eficiente in functie de capacitatea lor de a caracteriza spatiul acustic si de volumul de calcul care il implica estimarea lor.
- ptr. reducerea numarului de distante de calculate, adesea se imparte spatiul vectorilor acustici in clase, fiecare clasa fiind reprezentata de un vector particular numit centroid (cuantizarea vectoriala).



- **antrenarea** – se achiziționează cunoștințele și le organizează în clase sau modele

- **recunoasterea**

• Rostire necunoscută: $X=\{x1,x2,...xI,...,xN\}$, unde $xI=[xI1,xI2,..., xId]^T$

ROSTIRI CUNOSCUTE (Baza modele)

$$Y1=\{y11,y12,...,y1J\} \quad (1)$$

$$Y2=\{y21,y22,...,y2J\} \quad (2)$$

...

$$YK=\{yk1,yk2,...,ykJ\} \quad (M)$$

• Calculez $D(X,Yk)$ ptr. $k=1,...M$

• Recunoscut: $X= Y_{best}$ cu $D(X,Y_{best}) \leq D(X,Y_k)$ ptr. $k=1,...M$

• OK - ptr. recunoasterea cuvintelor izolate dependente de vorbitor

- Sunt mai multe *abordări clasice* ale clasificării/deciziei în recunoaștere, astfel:
 - clasificarea automată
 - discriminarea funcțională
 - metodele statistice bayesiene
 - metoda proximității maxime a celor k-vecini (k-NN)
 - metode stochastice (HMM)
 - metode conexioniste (rețele neuronale)
 - metode structurale.
- *Inteligența artificială* permite învățarea automată a relației complexe între semnalele audio și structurile lingvistice fără intervenție umană detaliată.

Principalele modele AI moderne includ:

- DNN (Deep Neural Networks) - utilizate inițial pentru clasificarea fonemelor.
- CNN (Convolutional Neural Networks) – captează modele spațiale în spectrograma audio.
- RNN / LSTM / GRU – modelează dependențele temporale lungi din vorbire.
- Transformer / Attention Models (ex: Whisper, wav2vec 2.0)** – procesează întregul semnal în paralel și oferă performanțe superioare cu latență redusă.

3. Metrică în spațiul acustic

Distanța sau măsura deosebirii

Într-un spațiu metric n-dimensional ***distanța dintre doi vectori a și b*** trebuie să satisfacă următoarele condiții :

$$d(a,b) > 0, d(a,b)=0 \text{ dacă } a=b \quad (1)$$

$$d(a,b)=d(b,a) \quad - \text{ simetrie} \quad (2)$$

$$d(a,b) < d(a,c)+d(c,b) \quad - \text{ inegalitatea triunghiului} \quad (3)$$

Dacă simetria nu este respectată se folosește:

$$d_S(a,b) = 1/2 [d(a,b)+d(b,a)] \quad (4)$$

Definirea unei distanțe între 2 spectre trebuie :

- să aibă semnificație în planul acustic
- să se preteze formalismului matematic
- să definească un spațiu de parametri judicios aleși.

- Reprezentarea parametrică a caracteristicilor unui segment de vorbire poate fi privită ca un vector N-dimensional sau ca un punct în spațiul N-dimensional al caracteristicilor, unde N este numărul caracteristicilor reprezentate.

Pentru un vector cu n componente norma Holder este:

$$d_p(\mathbf{x}, \mathbf{y}) = \left[\sum_{k=1}^N |\mathbf{x}_k - \mathbf{y}_k|^p \right]^{1/p} = \|\mathbf{x} - \mathbf{y}\|_p$$

Ea definește o distanță și are diferite semnificații pentru valorile lui p, astfel :

p=1 - suma absolută a diferentelor coordonatelor (City block, Manhattan, L_1 norm)

$$d_1(\mathbf{x}, \mathbf{y}) = \sum_{i=1}^N |\mathbf{x}_i - \mathbf{y}_i|$$

Un exemplu tipic al acesteia este distanța Hamming, care este numărul de biți care diferă între 2 vectori binari)

p=2 distanța medie pătratică (euclidiană)

$$d_2(\mathbf{x}, \mathbf{y}) = \sqrt{\sum_{i=1}^N (\mathbf{x}_i - \mathbf{y}_i)^2}$$

p= ∞ distanța maximă între componente (distanța maximului)

$$d_{\infty}(\mathbf{x}, \mathbf{y}) = \max |\mathbf{x}_i - \mathbf{y}_i|_{i=1, N}$$

Daca B este o matrice definită pozitiv de ordinul k, se poate defini o distantă ponderată astfel:

$$d_B(\mathbf{x}, \mathbf{y}) = (\mathbf{x} - \mathbf{y})' \mathbf{B} (\mathbf{x} - \mathbf{y})$$

- Deoarece această matrice B depinde de unul din vectori, această definiție nu satisface proprietatea de simetrie (2).
- Nu toate distanțele folosite în recunoașterea automată a vorbirii sunt metrice, de exemplu *distanța Mahalanobis* sau distanța covarianței ponderate, care pentru două modele x și y este dată de relația :

$$d(\mathbf{x}, \mathbf{y}) = (\mathbf{x} - \mathbf{y})^T \mathbf{W}^{-1} (\mathbf{x} - \mathbf{y})$$

unde $\mathbf{W}^{-1}$ este o matrice definită pozitiv care ponderează diferit caracteristicile individuale ale modelului depinzând de utilitatea lor în identificarea segmentelor de vorbire în spațiul caracteristicilor.

Distanța euclidiană rezulta ptr. $\mathbf{W} = \mathbf{I}$, în timp ce distanța Mahalanobis face ca $\mathbf{W}$ să fie matricea de autocovarianță corespunzătoare vectorului de referință.

- avantajele teoretice ale distanței Mahalanobis din păcate nu sunt folosite pentru că este dificil de estimat W , dintr-un set redus de date de antrenament
- se folosește de obicei distanța euclidiană sau LPC care necesită numai N înmulțiri pentru un vector caracteristic n -dimensional față N^2 de pentru distanța Mahalanobis.
- se mai folosesc distanțele Cebîșev sau “city-block “ care reduc precizia sistemului de recunoaștere în favoarea unui volum de calcul mai redus.

$$D_c(T, R) = \sum_{n=1}^N |T_n - R_n|$$

unde T_n, R_n reprezintă al n -lea parametru din modelul de test și din cel de referință.

O altă distanță înrudită, cu volum de calcul redus, este următoarea :

$$D_s(T, R) = \frac{\max_n |T_n - R_n| - \min_n |T_n - R_n|}{\sum_{n=1}^N T_n}$$

Distanțe LPC

Două distanțe eficiente și asemănătoare sunt folosite cu coeficienții LPC:

- distanța Itakura-Saito (IS)
- LLR (log likelihood ratio).

Ele prezintă avantajul folosirii modelării matematice LPC pentru ponderarea efectelor diferiților coeficienți fără a fi nevoie de ortogonalitate. Forma cea mai folosită a d_{IS} între două modele reprezentate prin vectorul LPC de test T și unul de referință a_i ($i=1,...,L$) este următoarea :

$$dis_i = dis(a_i, a_T) = \frac{\sigma_i^2 a_i^T v_T a_i}{\sigma_T^2 a_T^T v_T a_T} + \log \left(\frac{\sigma_T^2}{\sigma_i^2} \right) - 1$$

unde v_T este matricea de autocorelație a rostirii de test, σ_T, σ_i - parametrii de câștig pentru modelele de test și de referință.

Distanța LLR este o versiune normalizată a distanței Itakura-Saito, LLR presupune că *amplitudinea este irelevantă* pentru recunoaștere deoarece cuvântul poate fi rostit mai tare sau mai slab. :

$$d_{LLR_i} = d_{LLR}(a_i, a_T) = \log \frac{a_i^T v_T a_i}{a_T^T v_T a_T}$$

Distanțe cepstrale

- Cepstrul unui semnal este definit ca transformata Fourier a log spectrului semnalului.
- Pentru un spectru de putere $X(\omega)$ care este simetric față de $\omega=0$ și este periodic pentru o secvență de eșantioane, seria Fourier a $\log|X(\omega)|$ se poate exprima astfel :

$$\log X(\omega) = \sum_{n=-\infty}^{\infty} c_n e^{-jn\omega}$$

Dacă semnalul vocal este modelat cu un spectru numai poli, de fază minimă :

$$X(\omega) \rightarrow \frac{\sigma^2}{|A(e^{j\omega})|^2}$$

- cepstrul rezultat are mai multe proprietăți interesante. Pentru un filtru stabil numai poli $\log A(z)$ se află în interiorul cercului unitate și se poate reprezenta prin dezvoltare Laurent :

$$\log[\sigma/A(z)] = \log \sigma + \sum_{m=1}^{\infty} c_n z^{-n}$$

Diferențiind cei doi membri ai ecuației în raport cu $1/z$, obținem:

$$c_n = -a_n - \frac{1}{n} \sum_{k=1}^{n-1} k c_k a_{n-k} \quad , n > 0$$

$a_0=1$ și $a_k=0$ pentru $k > p$.

$$d_c^2(L) = \sum_{n=1}^L (C_n - C_n')^2$$

Sisteme și metode utilizate în recunoaștere

• *Abordările traditionale* de recunoaștere a vorbirii și vorbitorului în raport cu metodele și tehnicile folosite se pot grupa în mai multe familii :

- sisteme bazate pe spectrograme
- sisteme bazate pe metodele programării dinamice (DTW)
- sisteme care folosesc cuantizarea vectorială (VQ)
- sisteme bazate pe modele Markov ascunse (HMM)
- sisteme utilizând rețelele neuronale NN
- sisteme bazate pe metoda TESPAN FANN
- sisteme folosind metode algebrice/statistice
- combinații ale lor.

• *Abordări IA*

1. Metode Tradiționale (Bazate pe Componente)

Înainte de dominația modelelor end-to-end, sistemele ASR erau alcătuite din mai multe componente separate, antrenate individual:

- **Modelul Acustic:** Utilizează algoritmi precum **Modelele Markov Ascunse (HMM)** sau, mai târziu, **Rețelele Neurale (DNN/RNN)** pentru a converti caracteristicile acustice ale sunetului (ex. coeficienți MFCC) în **unități fonetice** (foneme).
- **Modelul de Pronunție (Lexicon):** Conține o listă de cuvinte și modul în care acestea sunt pronunțate (secvențe de foneme).
- **Modelul de Limbă:** Utilizează **n-grame** sau alte metode statistice/neurale pentru a prezice probabilitatea ca o anumită secvență de cuvinte să apară într-o limbă. Ajută la rezolvarea ambiguității ("*casa ta e mare*" vs. "*casata Ema re*").
- Aceste componente erau integrate prin algoritmi de decodare, cum ar fi **algoritmul Viterbi**, pentru a găsi cea mai probabilă secvență de cuvinte.

2. Metodele End-to-End

Metodele **end-to-end** reprezintă o schimbare de paradigmă. Ele antrenează un **singur model neuronal** pentru a mapa direct semnalul audio de intrare la secvența de text de ieșire, eliminând nevoia de componente separate (model acustic, lexicon și model de limbă distincte).

Avantajele principale includ:

- **Simplitate:** arhitectură unică, mai ușor de implementat și optimizat.
- **Performanță Superioară:** Modelul optimizează direct eroarea finală (de cuvânt), iar componentele învață să lucreze împreună mai eficient.

Trei abordări majore end-to-end sunt:

A. Connectionist Temporal Classification (CTC)

B. Sequence-to-Sequence (Seq2Seq) cu mecanism de atenție (arhitectura encoder-decoder)

C. Modele bazate pe Transformer

- Se bazează exclusiv pe mecanismul de atenție (self-attention) pentru a procesa datele audio și a genera textul, permițând o *paralelizare* mult mai mare și, implicit, o antrenare mai rapidă și pe seturi de date mai mari.
- Exemple notabile sunt modelele de mari dimensiuni precum **Wav2Vec 2.0** sau arhitecturi folosite în **OpenAI Whisper**.

3. Modele preantrenate (self-supervised) - tendințe actuale

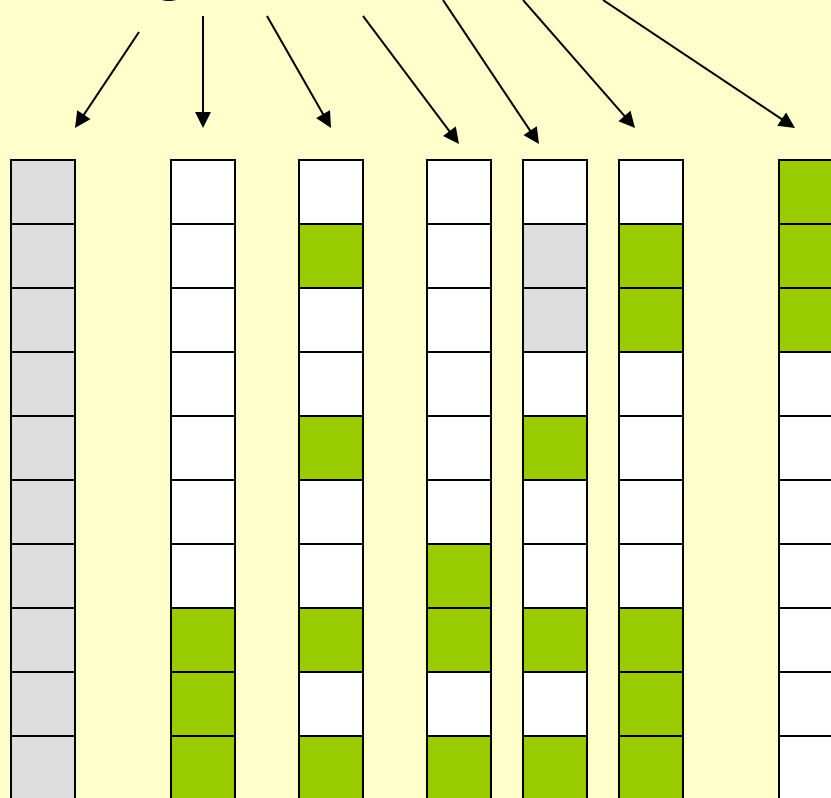
- Cea mai recentă evoluție în ASR se bazează pe **învățarea auto-supervizată** (*Self-Supervised Learning* - SSL):

- **Pre-antrenare:** Modelele sunt antrenate pe un volum masiv de date audio *netranscrise* (fără etichete) pentru a învăța reprezentări bogate și generale ale vorbirii (ex. prezicerea unor segmente de vorbire mascate).
- **Ajustare Fină (Fine-Tuning):** Modelul pre-antrenat este apoi ajustat (antrenat în continuare) pe un set de date *transcrise* (cu etichete) mult mai mic, folosind o metodă end-to-end (precum CTC). Această abordare (ex. *Wav2Vec 2.0*, *HuBERT*) a îmbunătățit dramatic performanța ASR, mai ales pentru limbile cu resurse de date transcrise limitate.

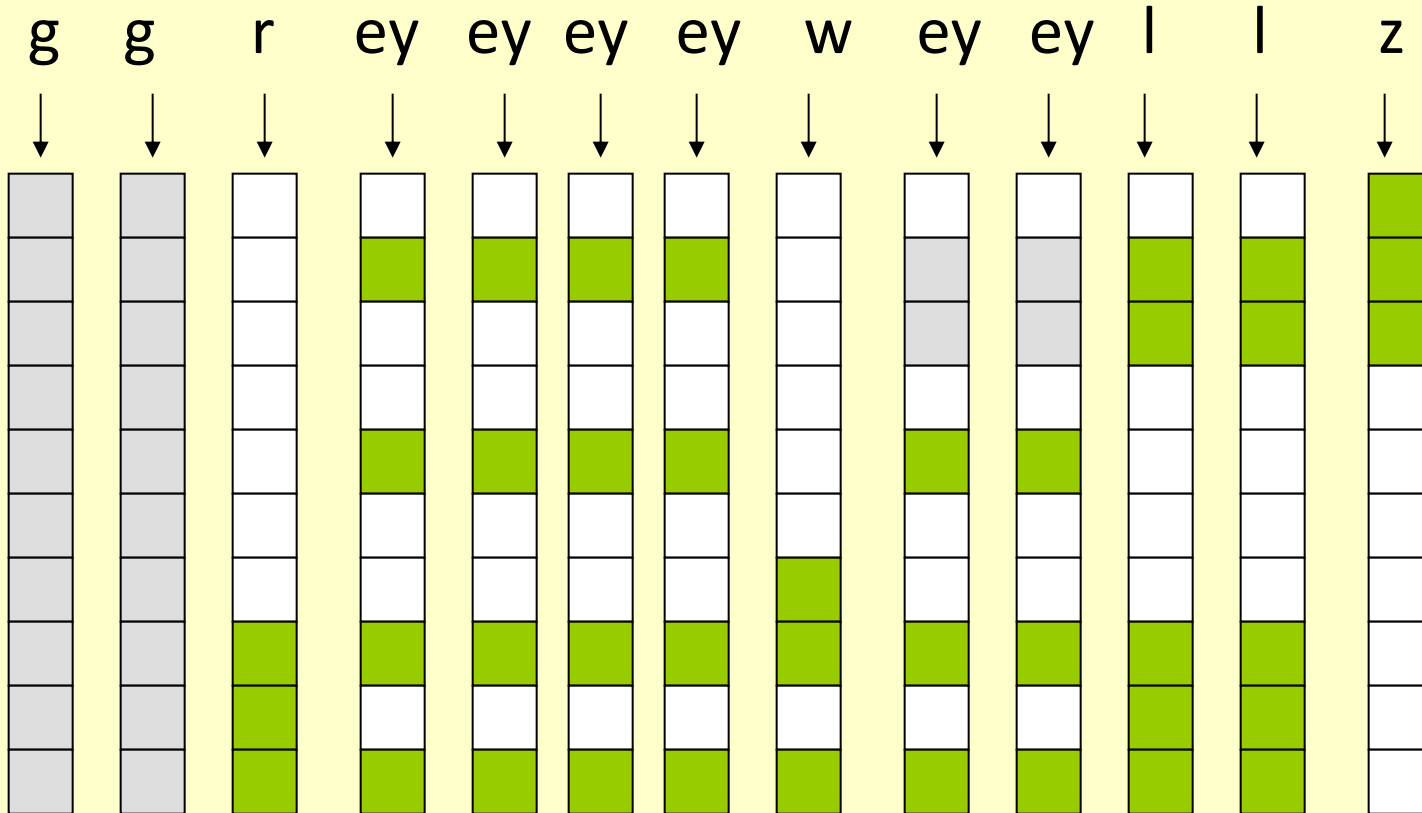
4. Metode de recunoasterea vorbirii

Cuvinte: grey whales

Foneme: g r e^y w e^y l z

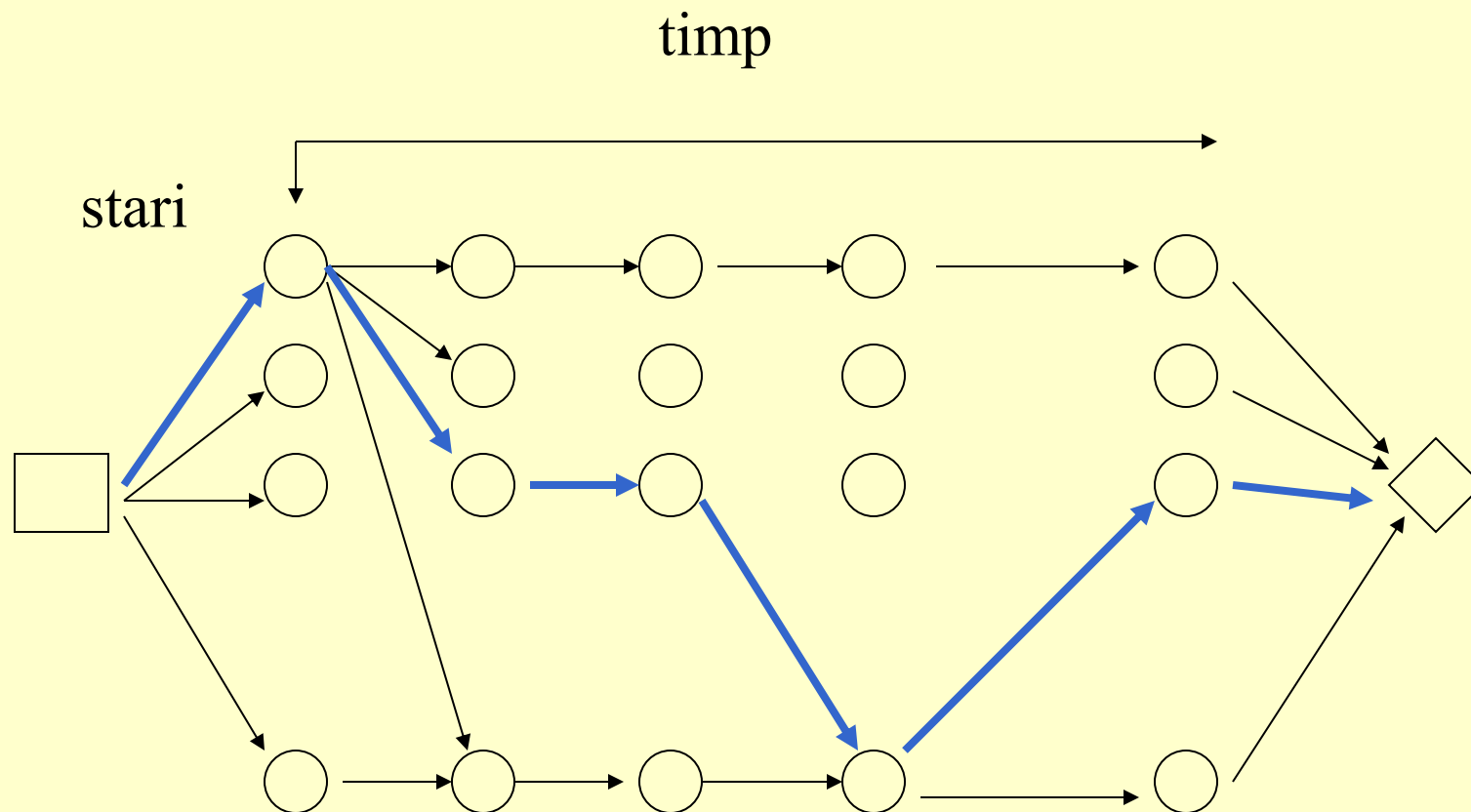


-Fiecare fonema are diferite caracteristici (de ex. Distributia de putere).



Cum se pot compara cuvintele daca exista diferente temporale sau altele?

Solutia: Programarea Dinamica



- Se *cauta o cale* care corespunde fonemelor cele mai probabile care au generat secventa de observatii "vectorii acustici" (extrasi din fiecare cadru de vorbire)

5. METODA ALINIERII DINAMICE (DTW)

- **Alinierea liniară și neliniară** - La calcularea distanțelor dintre două rostiri se presupune că cele două modele au aceeași lungime. Variațiile în lungime ale cuvintelor se datorează în mod special modului de rostire al vocalelor: rapid sau lent

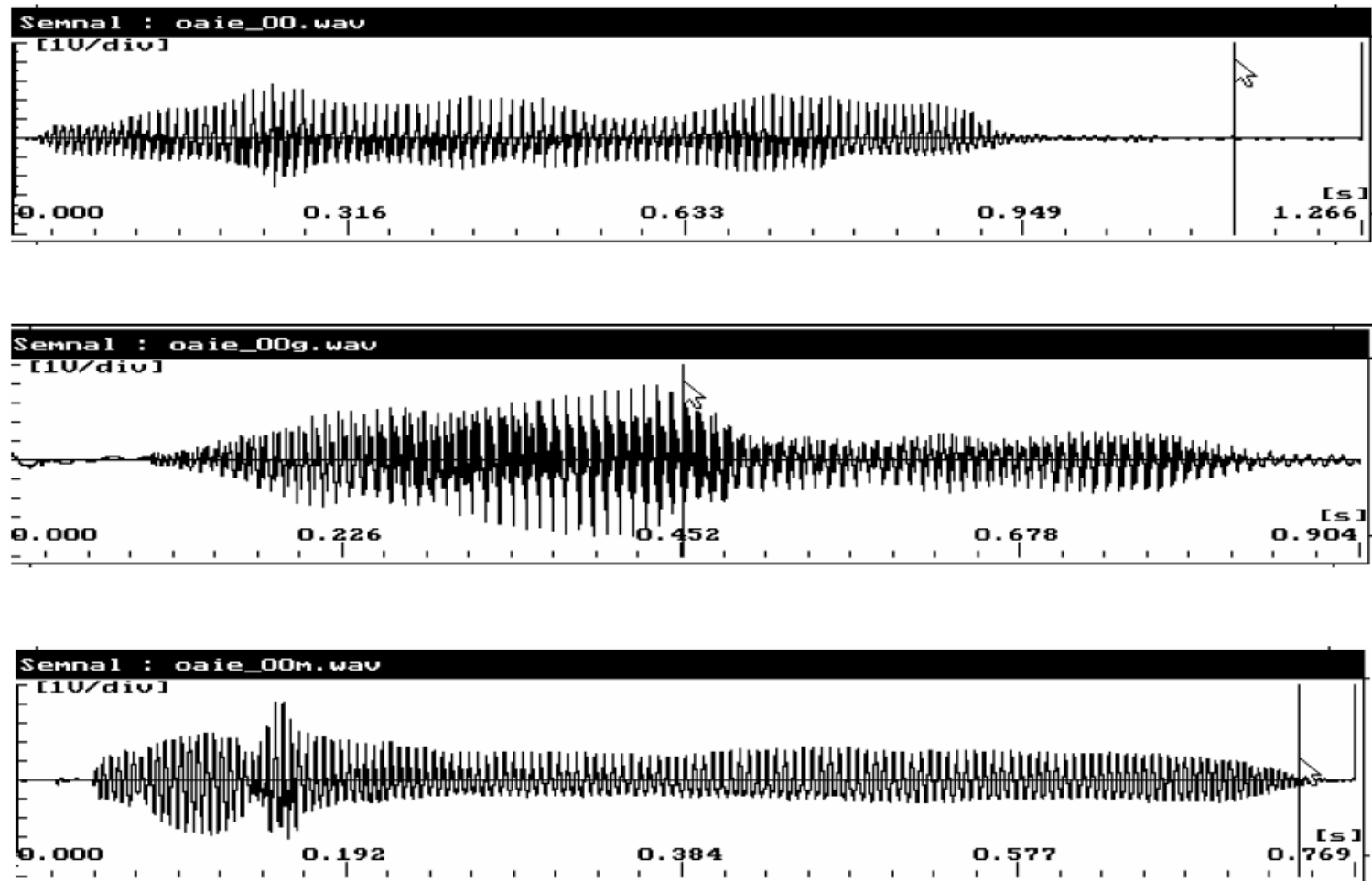


Fig. 2. Trei rostiri ale cuvântului "oaie" pronunțate de trei vorbitori

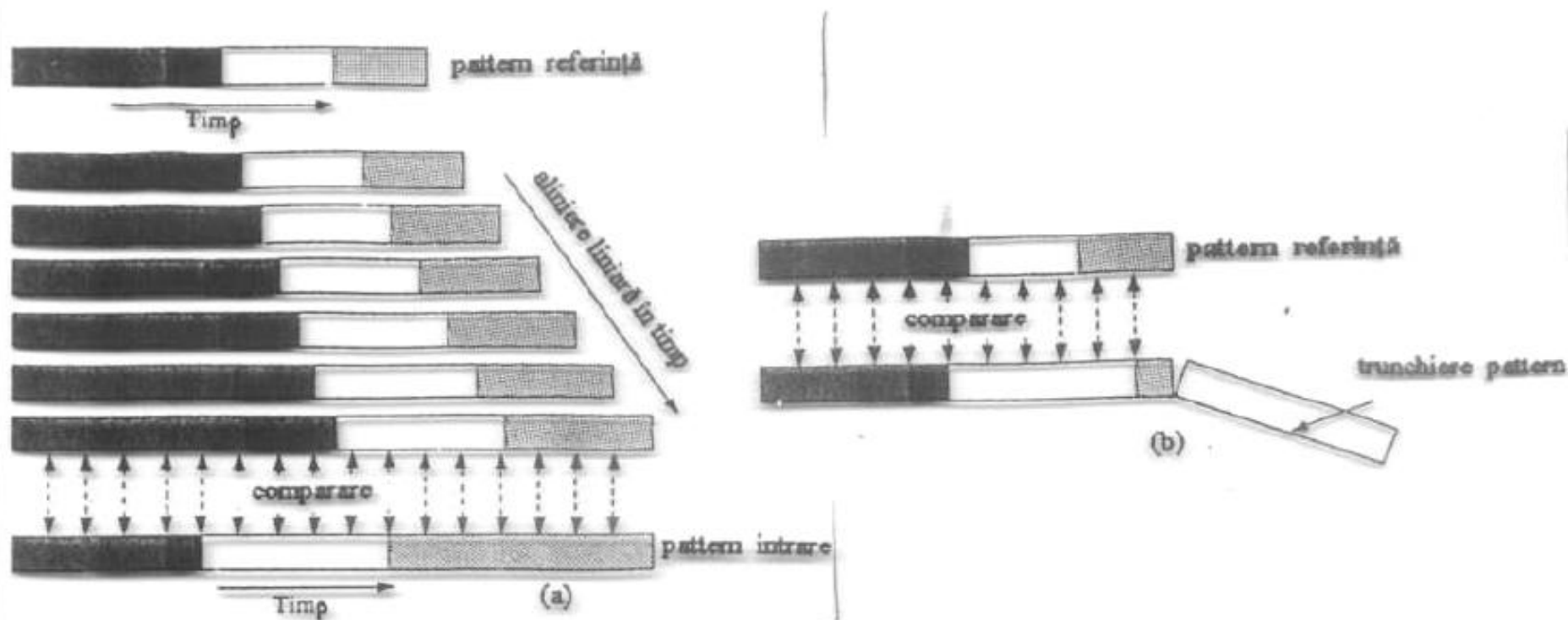


Fig. 3. Alinierea liniară în timp prin comprimare/dilatate proporțională (a) sau prin trunchiere (b)

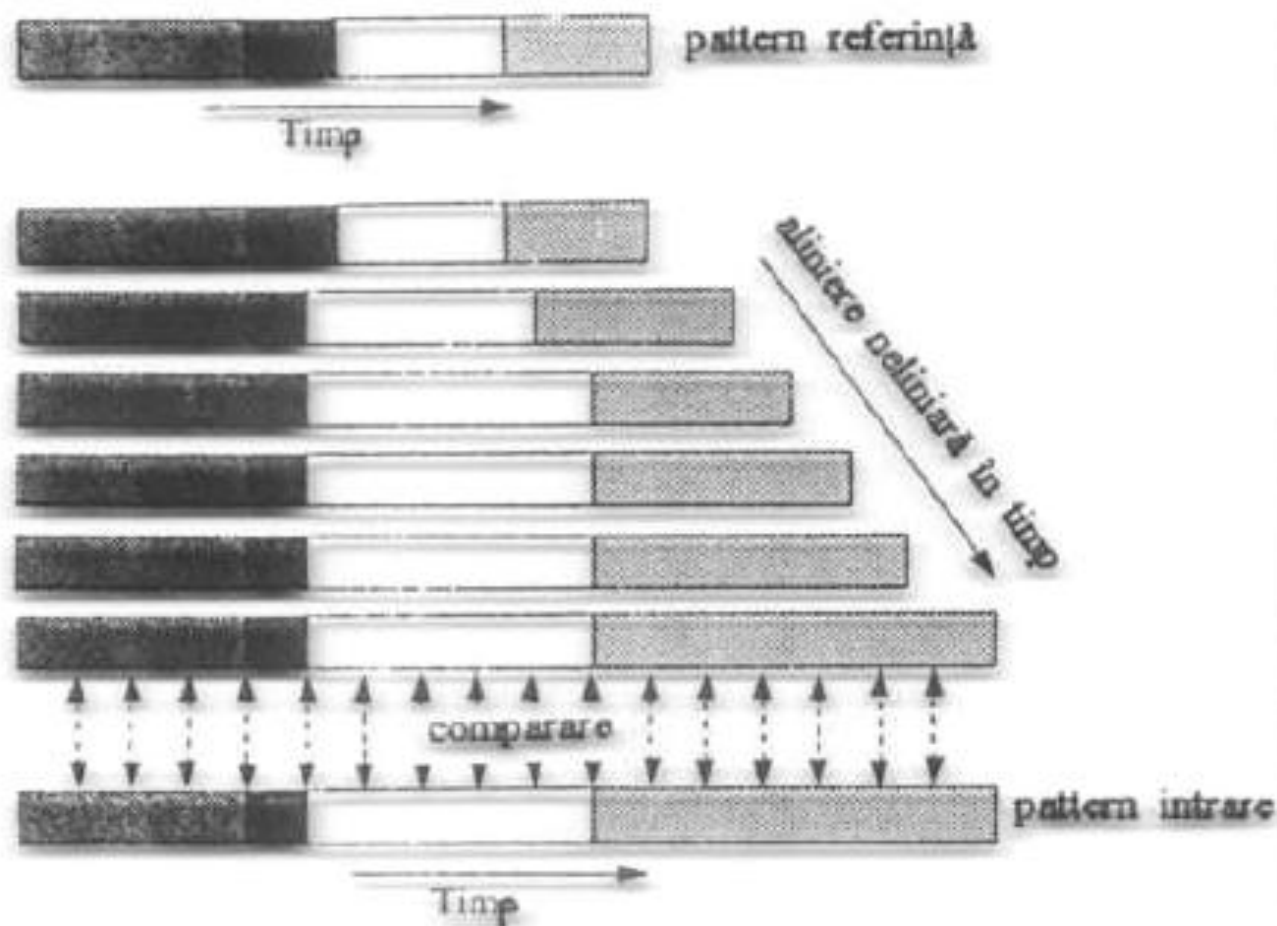


Fig. 4. Alinierea neliniară în timp

DTW este un algoritm care *calculează similaritatea* dintre două secvențe de date care pot varia în timp sau viteză. În recunoașterea vorbirii, acesta compară semnalele de vorbire (de exemplu, un cuvânt pronunțat de un utilizator cu un model de referință).

De ce este necesar?

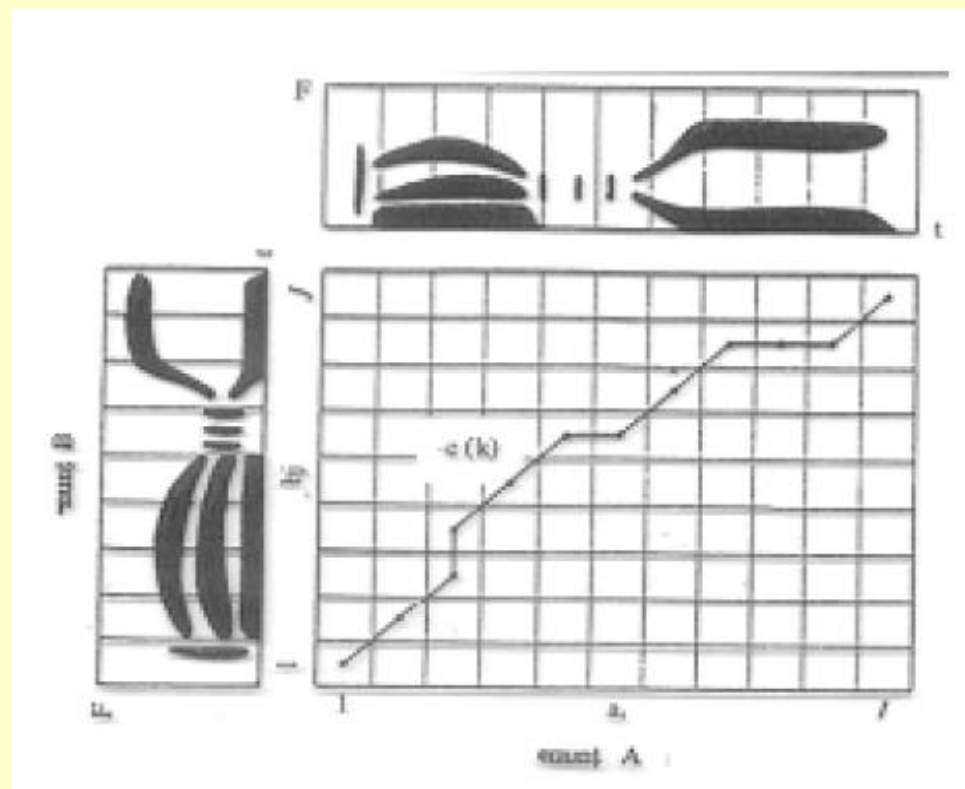
- Oamenii vorbesc cu **viteze diferite** (ex: "Hei" poate fi rostit rapid sau lent).
- Metodele simple (ex: distanța Euclidiană) eșuează dacă semnalele nu sunt sincronizate temporal.
- DTW **aliniază** nelinear cele două semnale în timp, pentru a găsi corespondența optimă.

Metoda DTW sau "potrivirea elastică" a modelelor este o tehnică bazată pe programarea dinamică, des folosită la rezolvarea problemelor de cost/drum minim - Această metodă DTW, *combină alinierea și calculul distanței* printr-o procedură de programare dinamică care, pentru varianta standard, presupune următoarele :

- *variația globală în viteza de vorbire pentru un vorbitor care rostește același cuvânt în diferite ocazii poate fi prelucrată prin normalizare liniară în timp*
- *Variațiile locale în viteza de vorbire, pentru care alinierea liniară este inadecvată, sunt mici și pot fi tratate folosind ponderi de distanță cunoscute ca și constrângeri locale de continuitate*
- *fiecare cadru din rostirea de test contribuie în mod egal la recunoaștere*
- *se folosește o singură măsură de distanță aplicată uniform pe toate cadrele*

Algoritmul de aliniere dinamică

- Algoritmul DTW comprimă sau dilată axa timp, neliniar pentru a potrivi poziția aceluiași fonem în modelul de intrare (test) și cel de referință
- Acest proces poate fi realizat folosind tehnica programării dinamice propusă de Bellman
- Algoritmul DTW - alinierea dinamică temporală a vorbirii - Slutsker, Vintsyuk, Velichko și Zaguruyko în Rusia și în paralel de Sakoe și Chiba în Japonia (1968).
- Această tehnică a avut impact în recunoașterea vorbirii, fiind în prezent o metodă importantă și des aplicată



Prezentarea algoritmului

Se presupun două secvențe în timp, de parametrii reprezentați de vectorii :

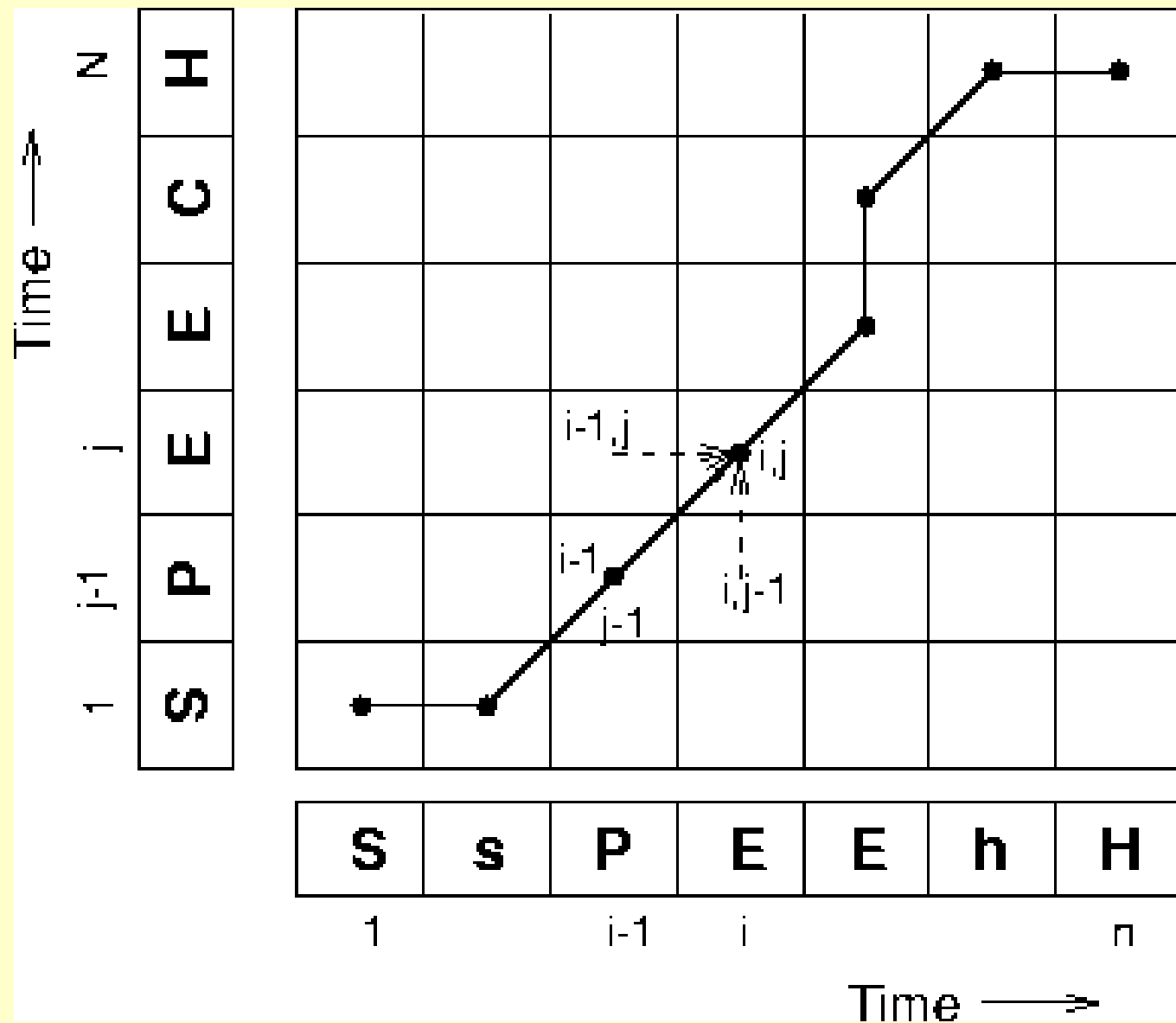
$$A= a_1, a_2,.....a_i \quad B= b_1,b_2,.....b_j \tag{1}$$

care trebuiesc comparați. Considerând planul împărțit de A și B ca în figura .1., funcția de aliniere în timp indică corespondența între axele timp ale secvențelor A și B și poate fi reprezentată printr-o secvență de puncte (latice) în plan $c =(i,j)$.

$$F = c_1, c_2,C_k.....,C_K \quad , \quad C_K=(i_k, j_k) \tag{2}$$

Dacă distanța spectrală între doi vectori caracteristici a_i și b_j se reprezintă prin: $d(c) = d(i,j)$, atunci suma distanțelor până la sfârșitul secvenței de-a lungul lui F se poate reprezenta prin: (3)

$$D(F) = \frac{\sum_{k=1}^K d(c_k)W_k}{\sum_{k=1}^K W_k}$$



Cu cât valoarea este mai mică, cu atât potrivirea (asemănarea) dintre A și B este mai bună. W_k reprezintă ponderi corelate cu F care trebuie să respecte următoarele [Fur89]:

- condiția de monotonie și continuitate:

$$i_{k-1} \leq i_k, \quad j_{k-1} \leq j_k \quad \text{și} \quad 0 \leq i_k - i_{k-1} \leq 1, \quad 0 \leq j_k - j_{k-1} \leq 1 \quad (4)$$

- condiția la limite (capete):

$$i_1 = j_1 = 1, \quad i_k = I, \quad j_k = J \quad (5)$$

- condiția de limitare a ferestrei:

$$|i_k - j_k| \leq r, \quad r = \text{constant}. \quad (6)$$

care se aplică pentru a preveni contracții/extensii exagerate.

Ecuția se simplifică, definind W_k astfel încât numitorul relației 3 să fie constant, deci independent de drum (F). De exemplu pentru:

$$W_k = (i_k - i_{k-1}) + (j_k - j_{k-1}), \quad i_0 = j_0 = 0 \quad (7)$$

W_k devine distanța ‘city-block’ (3.9.13) și

$$\sum_{k=1}^K W_k = I + J \quad (8)$$

iar relația (3) devine :

$$D(F) = \frac{1}{I + J} \sum_{k=1}^K d(c_k) \cdot W_k \quad (9)$$

Obiectivul minimizării funcției F poate fi realizat fără a examina toate posibilitățile funcției F, ținând cont de caracterul aditiv:

$$\begin{aligned}
 g(c_k) = g(i,j) &= \min_{c_1, \dots, c_{k-1}} \sum_{m=1}^K d(c_m)w_m = \min_{c_1, \dots, c_{k-1}} \left[\sum_{m=1}^{K-1} d(c_m)w_m + d(c_k)w_k \right] = \\
 &= \min_{c_{k-1}} \left[\min_{c_1, \dots, c_{k-2}} \left[\sum_{m=1}^{K-1} d(c_m)w_m \right] + d(c_k)w_k \right] = \\
 &= \min_{c_{k-1}} [g(c_{k-1}) + d(c_k)w_k]
 \end{aligned} \tag{10}$$

Folosind condiționările (.4), (.5) și (.6) pentru F și formularea de mai sus pentru W_k , relația (.10) se poate rescrie:

$$g(i,j) = \min \begin{pmatrix} g(i, j-1) + d(i, j) \\ g(i-1, j-1) + 2d(i, j) \\ g(i-1, j) + d(i, j) \end{pmatrix} \tag{11}$$

Fie condițiile inițiale:

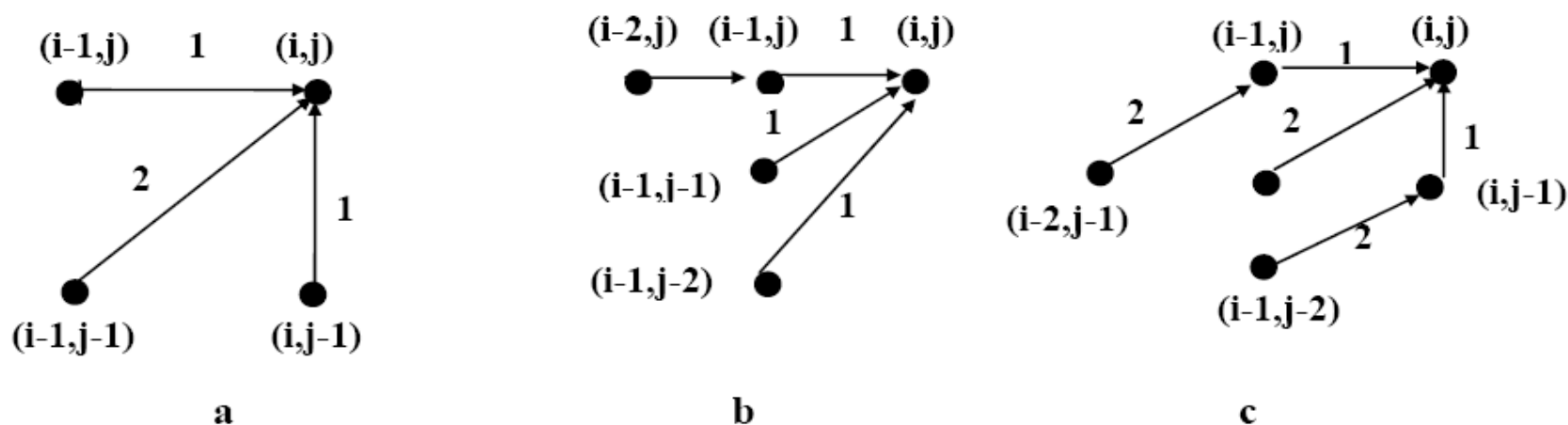
$$g(1,1) = 2 * d(1,1) \text{ și } j=1 \tag{12}$$

Plecând de la această distanță locală se calculează distanța între cele două secvențe A și B cu ajutorul relației (.11), variind I între limitele ferestrei definite. Calculele se repetă incrementând j până j=J, iar distanța globală se obține ca:

$$D(F) = D(A,B) = g(I,J)/(I+J) \tag{13}$$

Variante ale comparației prin programare dinamică

Pentru evaluarea funcției de aliniere și calculul distanței $D(F)$ (.3) trebuie definiți funcția de ponderare W_k și coeficientul de normalizare $\sum W_k$. Sunt mai multe tipuri de funcții de ponderare propuse (Sakoe, Myers) care rezultă din experimentele de recunoaștere și care reprezintă constrângeri locale ce se impun (figura .2) .



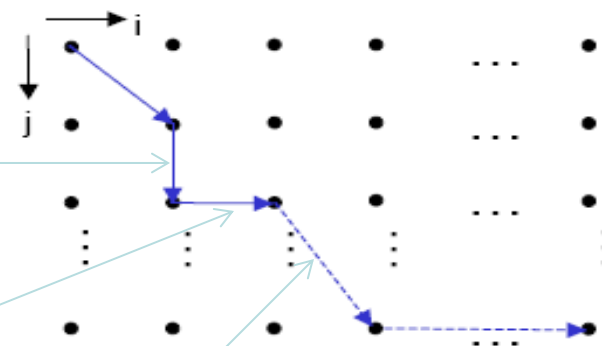
Constrângeri locale

Weighted Directed Path: **Arcs**

Insertion : $\{G[i, j], G[i, j + 1]\} \rightarrow w(0,1)$

Deletion : $\{G[i, j], G[i + 1, j]\} \rightarrow w(1,0)$

Change or Substitution : $\{G[i, j], G[i + 1, j + 1]\} \rightarrow w(1,1)$



```

# define Mx 100000
g[0][0]=0
for (j=1; j ≤ J; j++) g[0][j]=Mx ;
for (l=1 ; l ≤ L; l++)
{
    g[l][0]=Mx ;
    for (j=1; j ≤ J; j++)
    {
        d= dist (i,j) ;
        g[l][j]=min ( g[l-1][j]+d ,
                      g[l-1][j-1]+d+d,
                      g[l][j-1]+d ) ;
    }
}

D=g[L][J]/L+J;

```

unde $\text{dist}(i,j)$ este o funcție care returnează distanța spectrală între evenimentele: i a rostirii A și j din rostirea B.

Volumul de calcul cerut de algoritmul DTW

- algoritmul trebuie să calculeze o **distanță locală** $d(i,j)$ și una **cumulată (parțială)** $g(i,j)$.
- pentru a compara două cuvinte izolate de lungime I și J se tratează toate punctele (i,j) ale rețelei de puncte $[1,I] \times [1,J]$, la recunoașterea unui cuvânt izolat este nevoie să se analizeze $M \times Q \times L$ puncte, unde
 - M reprezintă dimensiunea vocabularului,
 - Q este numărul mediu de referințe/cuvânt ,
 - L este lungimea medie a cuvintelor.

În general, numărul real de puncte evaluate este mult mai mic din două motive:

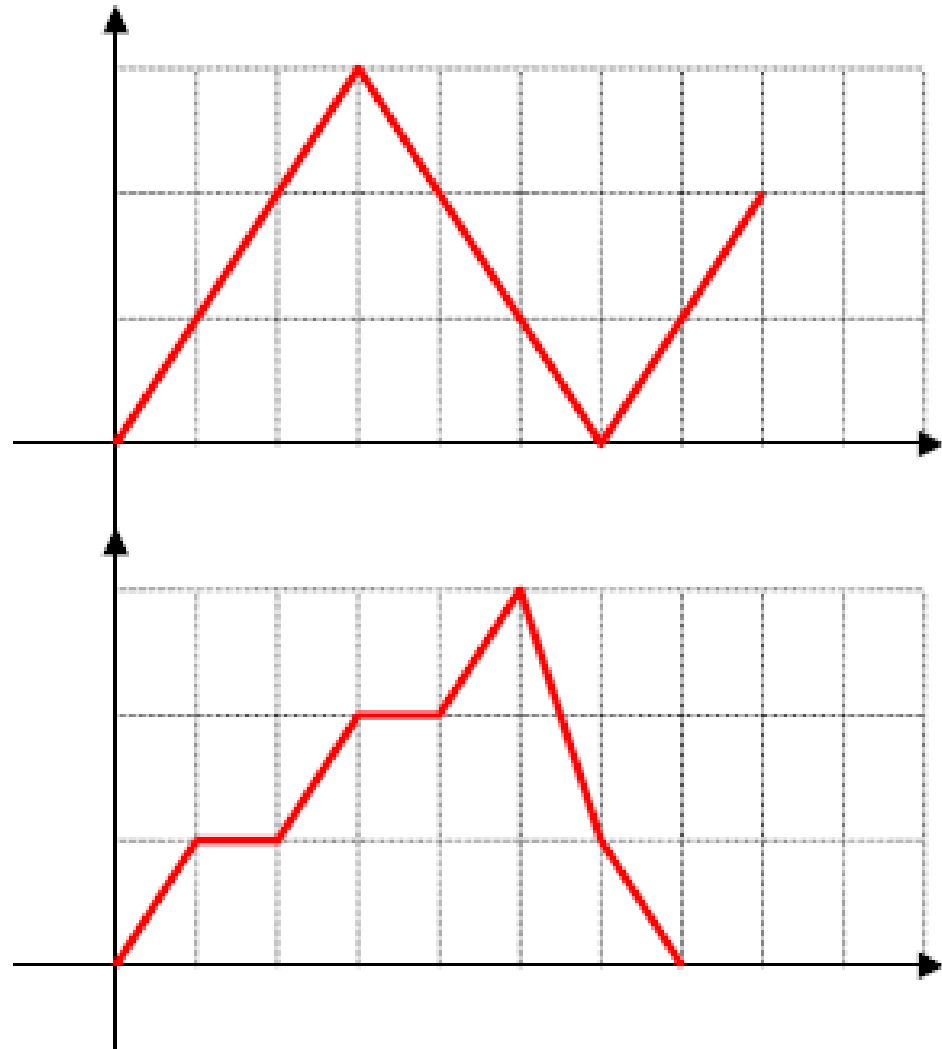
- din domeniul $[1,I] \times [1,J]$ interesează doar o zonă în jurul diagonalei, suprafața analizată reducându-se cu un factor de 0,4 ...0,1
- compararea unui cuvânt de recunoscut cu o referință poate fi abandonată dacă: indicele de distanță $D(T,R)$ nu se poate situa sub un prag dat.

Tema. Cautati variante modificate ale algoritmului DTW.

Exemplu. Găsiți calea de aliniere și distanța între următoarele două secvențe.
(City-block)

$$x[n] = \{0, 1, 2, 3, 2, 1, 0, 1, 2\}$$

$$y[n] = \{0, 1, 1, 2, 2, 3, 1, 0\}$$



$$\underline{d(i, j) = |x[i] - y[j]|}$$

		$x[n]$									
		0	1	2	3	2	1	0	1	2	
$y[n]$	0	0	1	2	3	2	1	0	1	2	
	1	1	0	1	2	1	0	1	0	1	
	1	1	0	1	2	1	0	1	0	1	
	2	2	1	0	1	0	1	2	1	0	
	2	2	1	0	1	0	1	2	1	0	
	3	3	2	1	0	1	2	3	2	1	
	1	1	0	1	2	1	0	1	0	1	
	0	0	1	2	3	2	1	0	1	2	

	0	1	2	3	2	1	0	1	2
0	0	1	3	6	8	9	9	10	12
1	1	0	1	3	4	4	5	5	6
1	2	0	1	3	4	4	5	5	6
2	4	1	0	1	1	2	4	5	5
2	6	2	0	1	1	2	4	5	5
3	9	4	1	0	1	3	5	6	6
1	10	4	2	2	1	1	2	2	3
0	10	5	4	5	3	2	1	2	4

1. **Extragere caracteristici:** Semnalele de vorbire sunt convertite în secvențe de vectori de caracteristici (ex: coeficienți MFCC).
2. **Calcul matrice distanțe:** Se calculează o matrice de distanțe între fiecare pereche de puncte din cele două secvențe.
3. **Găsire cale optimă:** Se găsește calea cu costul minim de-a lungul matricei, care "deformează" temporal una dintre secvențe pentru a se potrivi cu cealaltă.
4. **Scor similaritate:** Costul total al căii reprezintă gradul de similitudine.

Exemplu:

- **Secvența A:** "H e e e y" **Secvența B:** "H e y"
- DTW va alinia "e e e" din A cu un singur "e" din B, compensând diferența de durată.

Avantaje:

- Robust la variații de viteză în vorbire.
- Simplu și eficient pentru vocabular limitat.
- Folosit în sisteme "template-based" (fiecare cuvânt este stocat ca un model).

Dezavantaje:

- Computațional costisitor pentru baze de date mari.
- Sensibil la zgomot și diferențe de intonație.
- Înlocuit parțial de modelele HMM și rețelele neuronale în sistemele moderne.

Aplicații:

- Recunoaștere vorbire izolată (cuvinte sau comenzi scurte).
- Verificare vorbitor.
- Aliniere temporală în bioinformatică sau analiză gesturi.