

ASRSV – curs 10

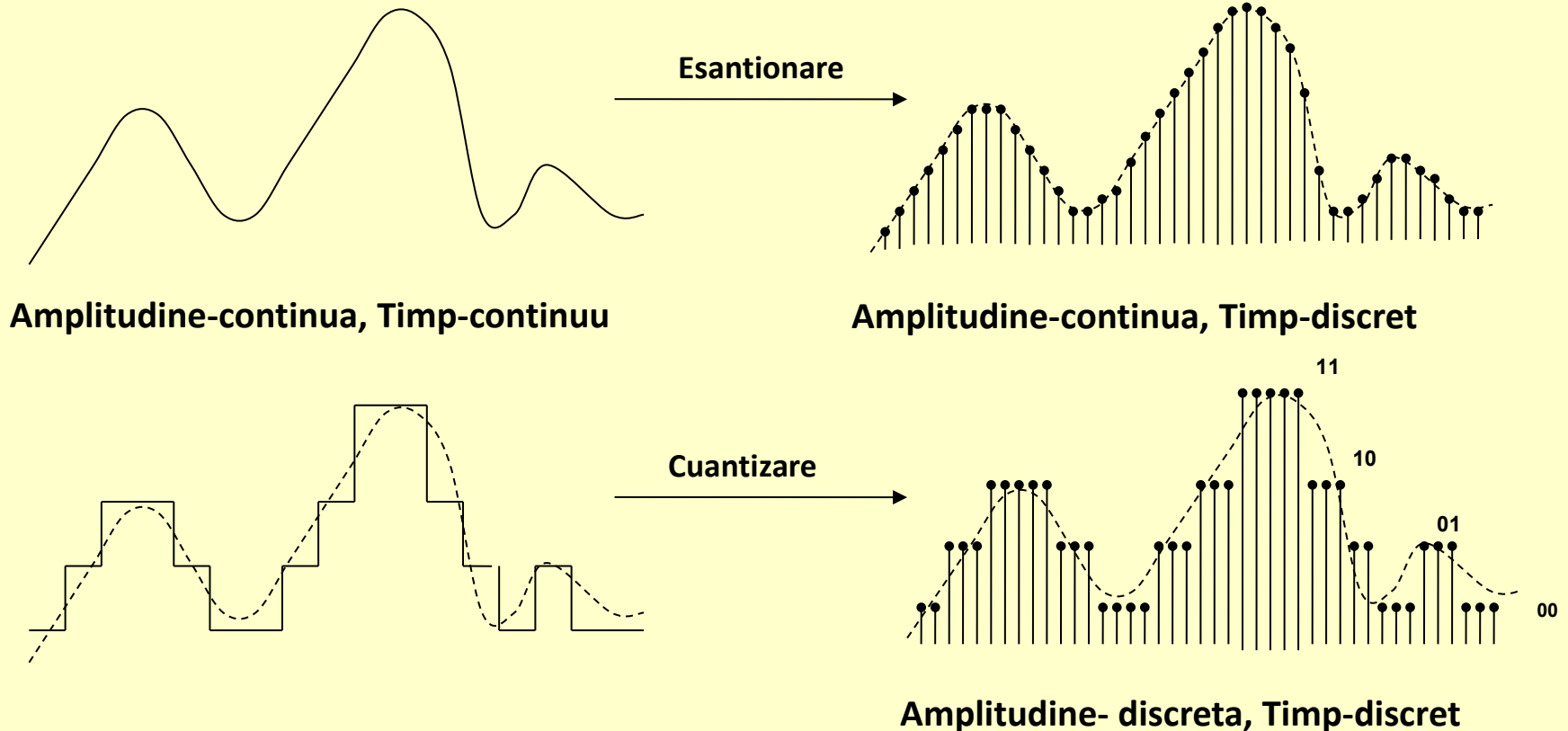
Metode si tehnici de recunoastere a vorbirii (2)

- **Sisteme și metode utilizate în recunoaștere**
- Abordările de recunoaștere a vorbirii/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 (MMA) - HMM
 - sisteme bazate pe mixturi Gaussiene GMM
 - sisteme utilizând rețelele neuronale NN
 - sisteme bazate pe metoda TESPAN FANN
 - sisteme folosind metode algebrice/statistice
 - combinații ale lor.
 - Abordări IA

Cuantizarea vectorială (VQ)

Introducere

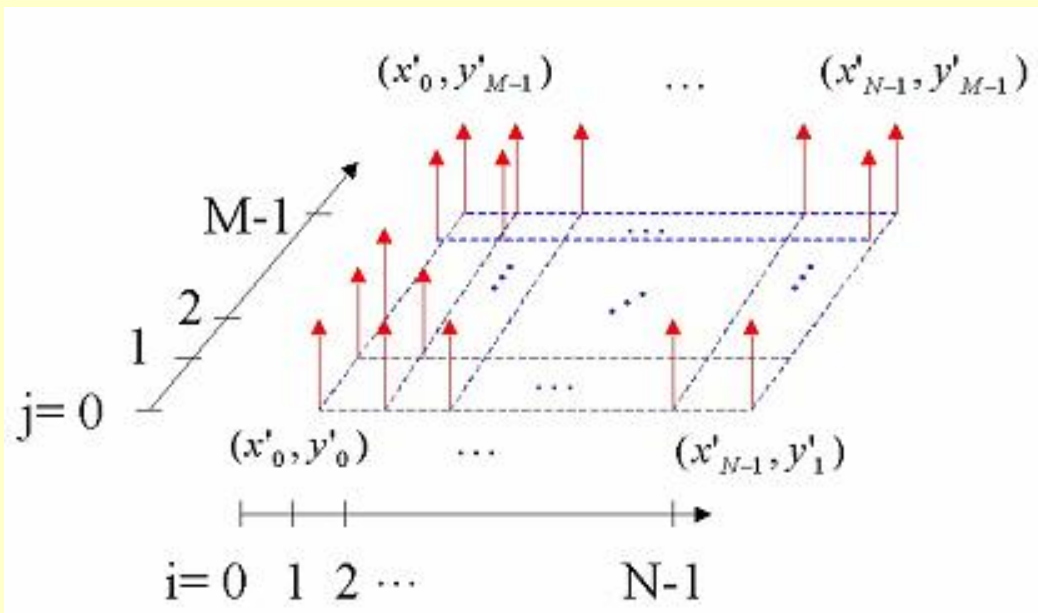
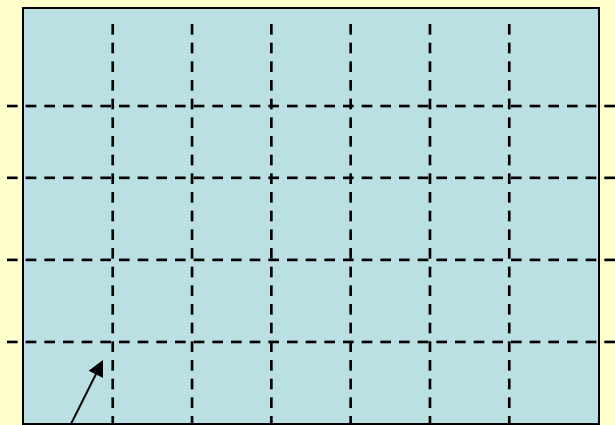
- Majoritatea semnalelor in natura sunt analogice



Exemplu pe imagine.

Esantionarea

- Se esantioneaza valorile imaginii in nodurile unei grile regulate pe planul imaginii.
- Un pixel (element de imagine) la coordonatele (i, j) este valoarea intensitatii imaginii in punctul grilei indexat cu număr întreg de coordonate (i, j) .



Exemple de esantionare



256x256



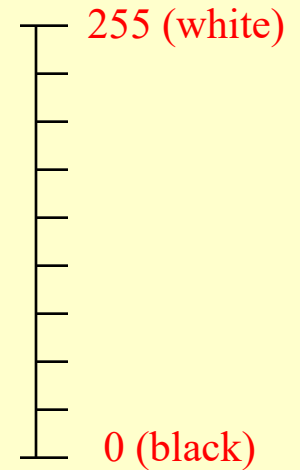
64x64



16x16

Cuantizarea scalara (SQ)

- Este un proces de transformare a unei imagini reale esantionate luand numai un numar finit de valori distincte.
- Fiecare valoare esantionata a imaginii pe o scara de 256 nivele de gri este reprezentata pe 8 biti.



Exemple de cuantizare



8 bits / pixel



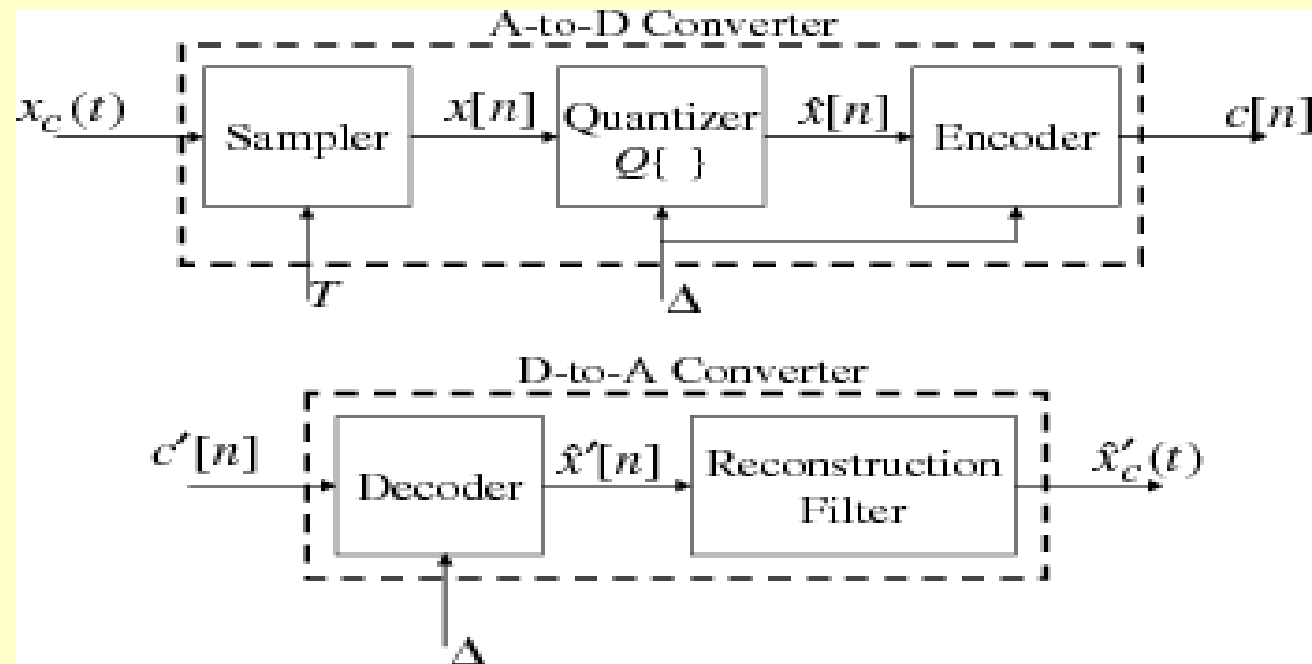
4 bits / pixel



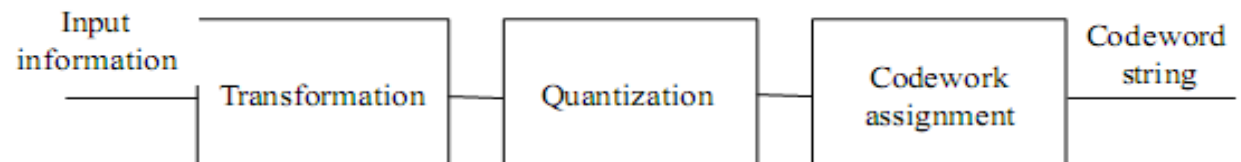
2 bits / pixel

Cuantizarea in sistemele de prelucrarea semnalelor

- Conversia A-D:



- Codarea



(a) source encoder



(b) source decoder

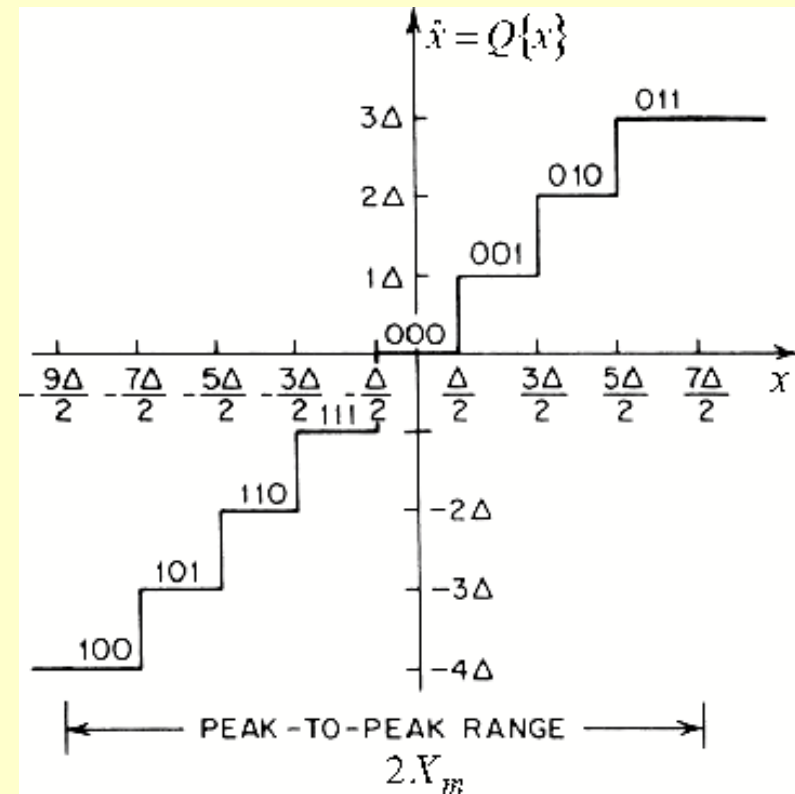
- **Problemele cuantizarii:**
 - Cuantizarea este un proces ireversibil
 - Cuantizarea este o sursa de pierdere a informatiei (erori)
- Cuantizarea este un nivel critic in compresia semnalelor, cu efecte asupra:
 - Distorsiunii semnalelor reconstruite
 - Debitului binar al codorului
- Obiectivul – cuantizor optimal sau de inalta calitate:
 - Distorsiuni medii cat mai mici (D).
 - debit cat mai redus (R).
- Cuantizarea: a) Uniforma
 - b) Neuniforma
 - c) Adaptiva
 - d) Predictiva

- **a) Cuantizarea uniforma**

- Toate regiunile de decizie a cuantizarii sunt de dimensiune egala :

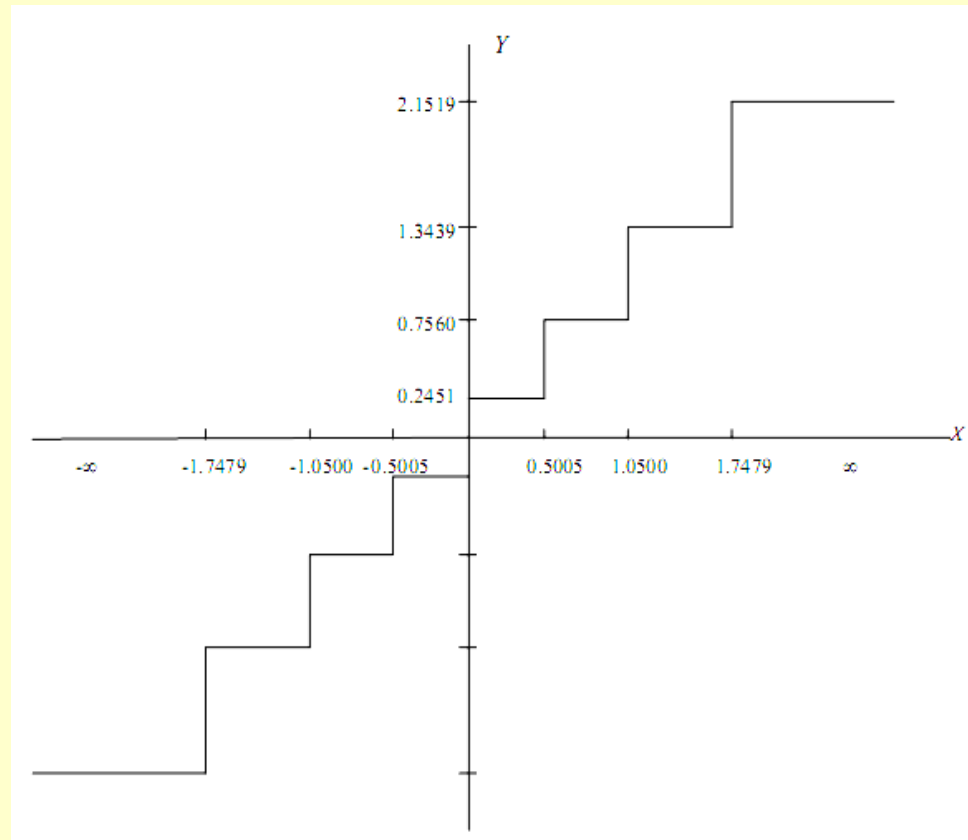
$$\Delta = 2X_m / 2^B = 2X_m / N$$

- Nivelele de cuantizare sunt la jumatatea nivelelor de decizie
- La N regiuni de cuantizare, se folosesc $B = \log_2(N)$ biti pentru a reprezenta fiecare valoare cuantizata.



b) Cuantizare Neuniforma

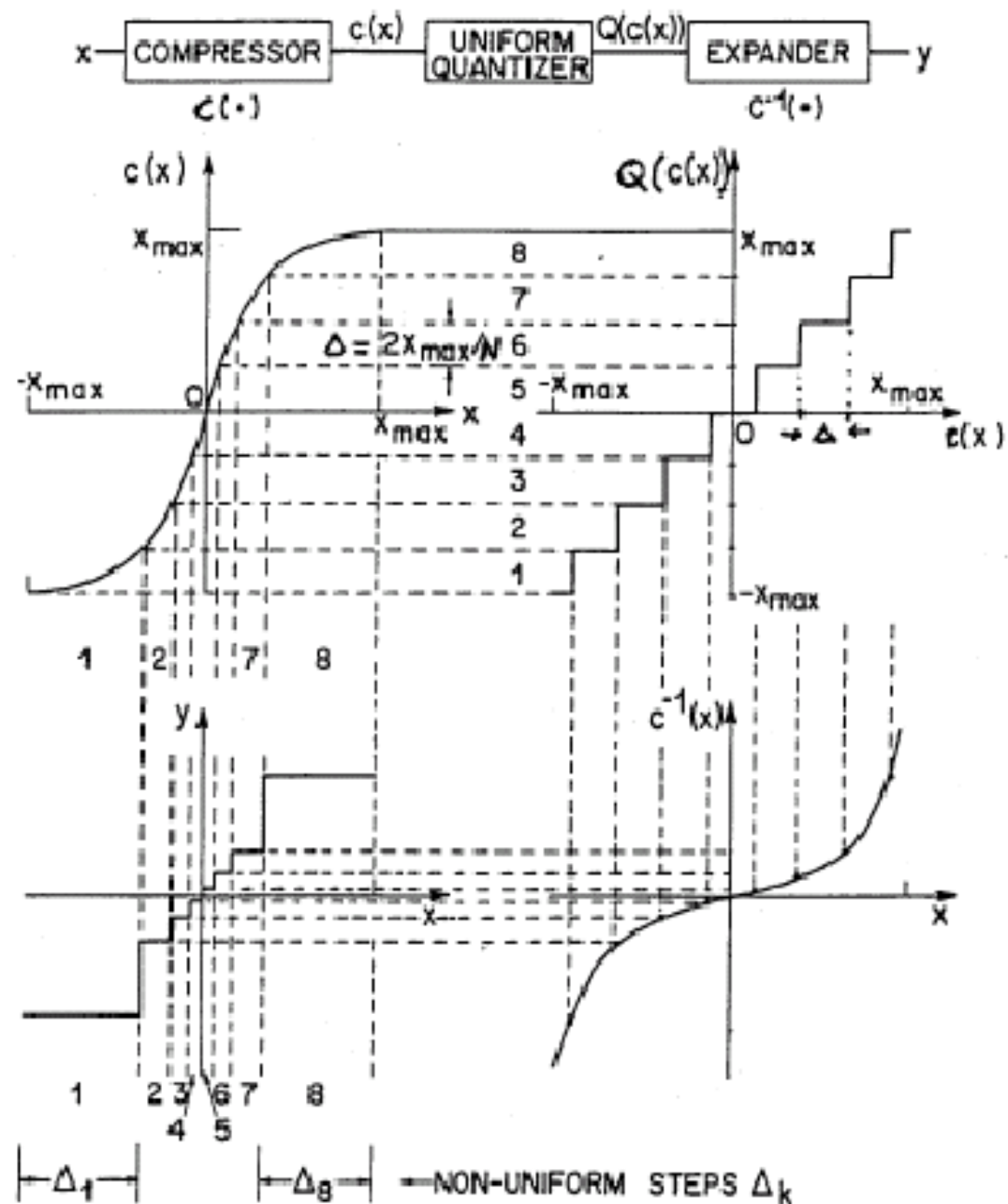
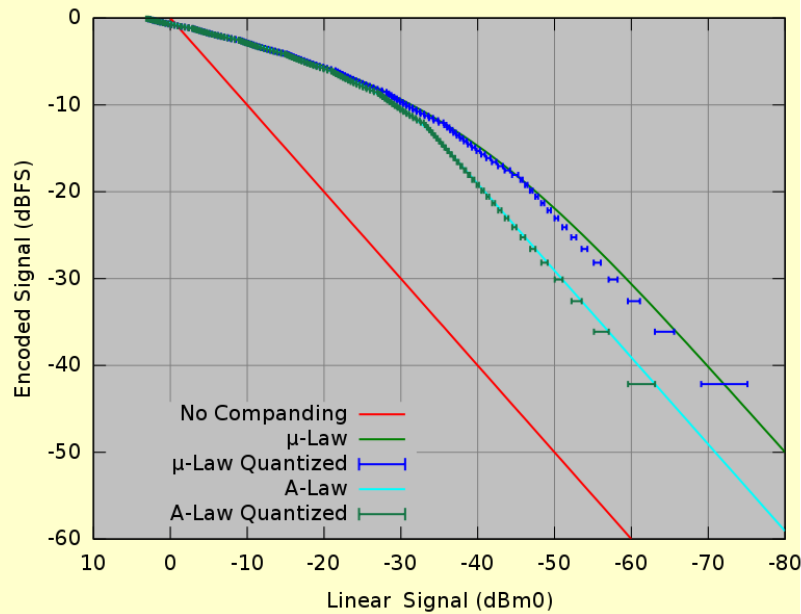
- Semnalele cu Fct Densit Probab. de tip (Gaussian, Laplacian etc.).
- Semnalele de intrare de amplitudine mica trebuie cuantizate mai precis
- Semnalele de intrare de amplitudine mare se cuantizeaza mai grosier
- Algoritmul Lloyd-Max: algoritm numeric iterativ ptr. cuantizorul optimal



Exemplu – G.711

- G.711 - standard ptr. compandare audio, un codor al formei de unda a vocii, cunoscut si ca PCM logaritmic.
- Frecventa de esantionare 8 kHz
- PCM uniform foloseste 12-16 bits/sample
- G.711 foloseste cuantizare neuniforma cu 8 bits/sample; compander (a-law/ μ -law) + cuantizor uniform
- debit binar 64 kbit/s (8 kHz x 8 bits/sample).
- Intarzierea tipica algoritmica este 0.125 ms

Exemplu – G.711



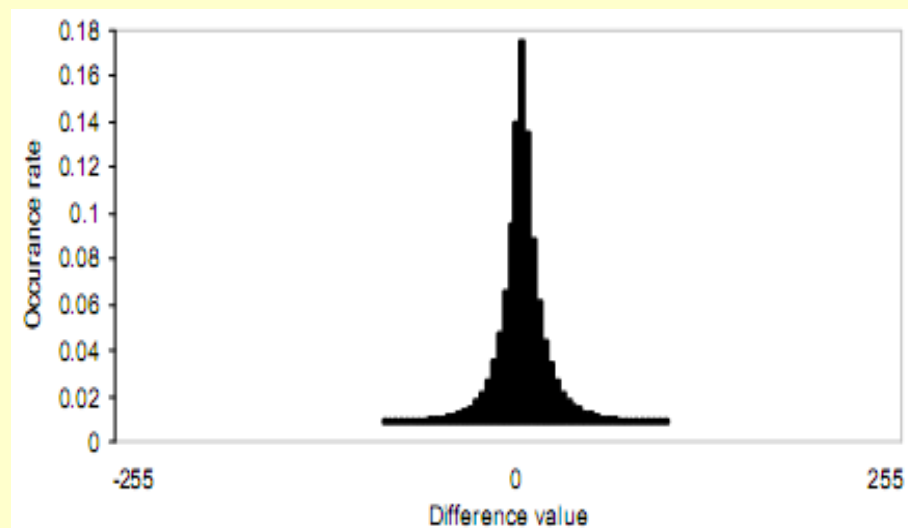
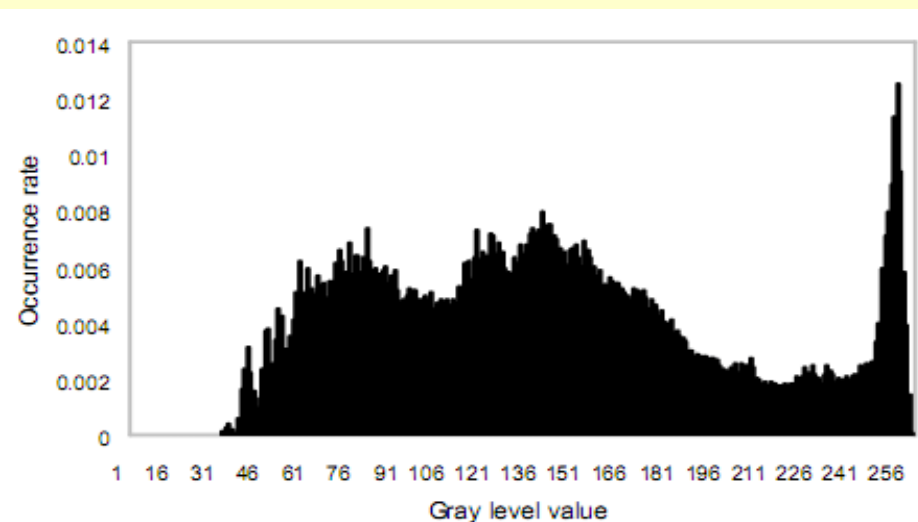
(From Jayant & Noll)

c) Cuantizarea Adaptiva

- Cuantizarea Uniforma : regiunile de cuantizare fixe
- Cuantizarea neuniforma : regiunile de cuantizare variabile, dar constante in timp
- Cuantizarea Adaptiva :
 - Regiunile de cuantizare variaza in timp
 - Cuantizorul se adapteaza la modificarile statisticii semnalului de intrare
 - Furnizează distorsiune mai mică - performanță mai bună
 - Costul: intarziere prin procesare suplimentara si cerinte suplimentare de stocare

d) Cuantizarea Predictiva

- în locul codării directe a semnalului, se codifică diferența dintre semnalul însasi și cel prezis prin tehnica de codare diferențială;
- este cunoscut și sub numele de codor de predicție
- Cuantizorul necesită gamă de cuantizare mai redusă
- Ex: histograma imagine în comparație cu o histogramă a imaginii diferența



Cuantizare Vectoriala

- Este o operație care generalizează cuantizarea scalară și constă în reprezentarea unui vector X , ale cărui M componente sunt valori reale și continue, printr-un vector $(x \in \mathbb{R}^M)$, aparținând unei mulțimi finite:

$y = \{y_i \in \mathbb{R}^M, i = 1, 2, \dots, K\}$ de vectori care alcătuiesc **dicționarul**.

$$y = q(x), \quad y \in Y$$

y_i - cuvânt de cod (centroid, codeword) ,

k - mărimea dicționarului (numărul de nivele de cuantizare)

- Scopul cuantizării vectoriale poate fi **reducerea dimensiunii bazei de date**

N parametri

2	3	5	4	4
1	2	3	12	3
2	4	5	7	8
...	...	...	...	...
...	...	...	...	...
5	7	3	19	4
6	7	9	0	2

P

V
e
c
t
o
r
i



N parametri

2	3	5	4	4
1	2	3	12	3
...	...	...	...	...
...	...	...	...	...
5	7	3	19	4

Q
V
e
c
t
o
r
i

Q < P

- Alegerea unui dictionar constă în partitionarea spatiului M - dimensional al vectorului x în k regiuni $\{C_i\}$, $i = \overline{1, k}$, k clase (clustere), fiecărei clase fiindu-i asociat câte un cuvânt de cod y_i .

Cuantizorul asignează cuvântul de cod y_i dacă x se află în regiunea C_i :

$$q(x) = y_i$$

- La realizarea cuantizării $x \rightarrow y$ apare o măsură de distorsiune (distanță) $d(x, y)$ care se definește între x și y și care de obicei se măsoară prin distanța euclidiană:

$$d(x, y) = \sum_{i=1}^M (x_i - y_i)^2$$

Distorsiunea totală medie se calculează prin relația:

$$D = \lim_{M \rightarrow \infty} \frac{1}{M} \sum_{n=1}^M d(x(n), y(n))$$

Un cuantizor este *optim dacă distorsiunea totală este minimizată* pe toate cele k nivele de cuantizare.

•Sunt două condiții pentru optim :

1)cuantizorul să fie realizat pe baza distorsiunii minime sau a regulii de selectie a celui mai apropiat vecin:

$$q(x)=y_i \text{ ddacă } d(x,y_i) \leq d(x,y_j) , \forall j \neq i , j=1,...k$$

2) fiecare cuvânt de cod y_i se alege astfel încât să se minimizeze distorsiunea medie în regiunea C_i ; centroidul pentru fiecare regiune depinde de modul de definire a măsurii de distorsiune:

$$y_i = \frac{1}{n_i} \sum_{x_i \in c_i} x_i , \quad i=1,...k$$

•Scopul urmărit la un sistem de codare este:

- **minimizarea distorsiunii medii** pentru un debit dat, sau
- **reducerea debitului** pentru o distorsiune acceptată;
- În cazul recunoașterii vorbirii, obiectivul urmărit este **reducerea volumului de calcul** printr-o compresie de date.

- Problema majora a VQ este **complexitatea computationala** ridicata.
- Solutia** : algoritm eficient de cautare corecta a centroidului pentru vectorul de intrare.
- Rezolutia Cuantizorului:

$$R = \frac{\log_2 N}{k} [bits/sample]; N = 2^{kR}$$

$$K \uparrow \Rightarrow R \downarrow \Rightarrow D \uparrow \text{ si}$$

$$K \uparrow \Rightarrow R \uparrow \Rightarrow D \downarrow$$

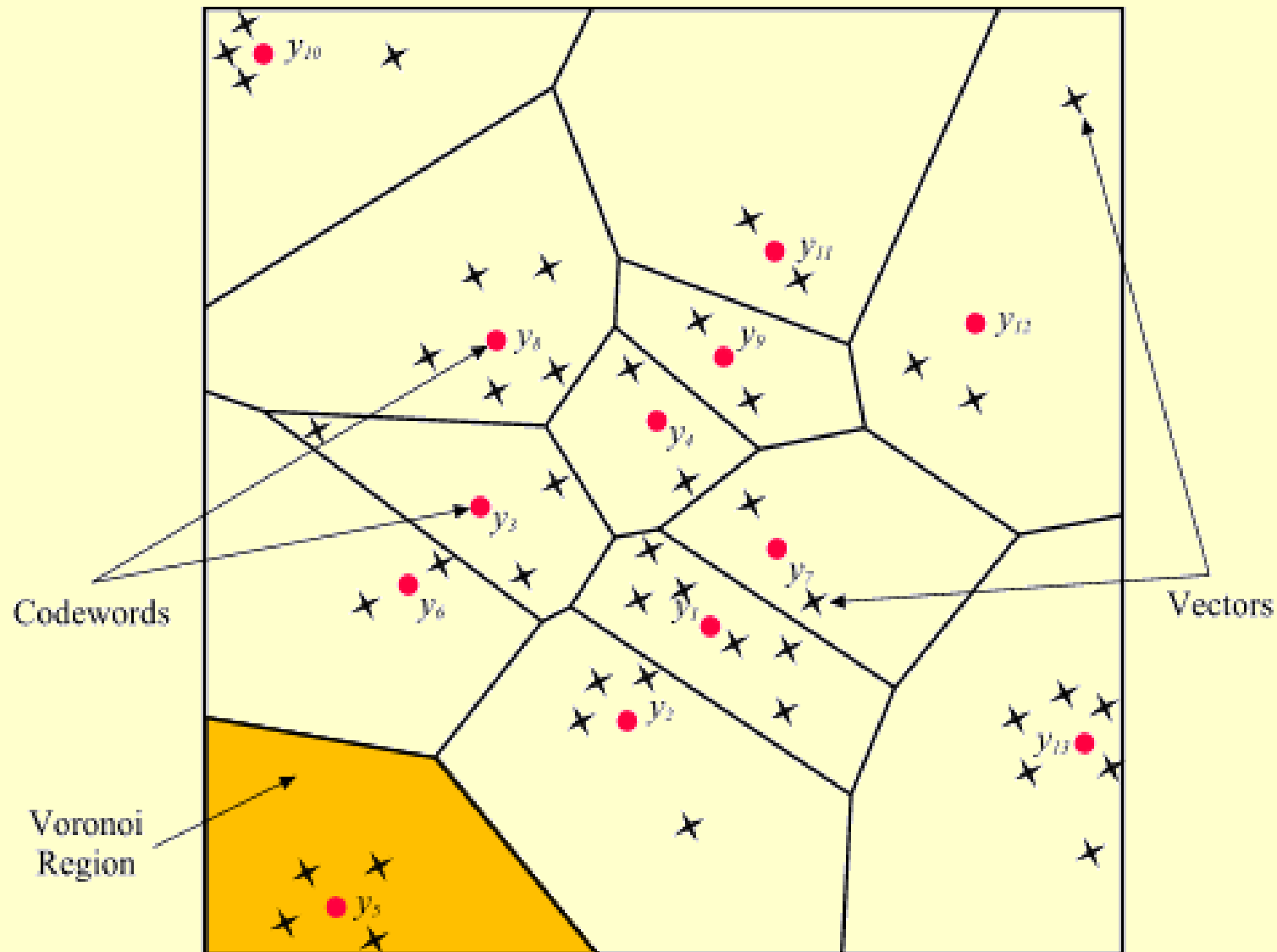
unde K- dimensiunea vectorilor, N- dimensiunea dictionarului, D- distorsiunea

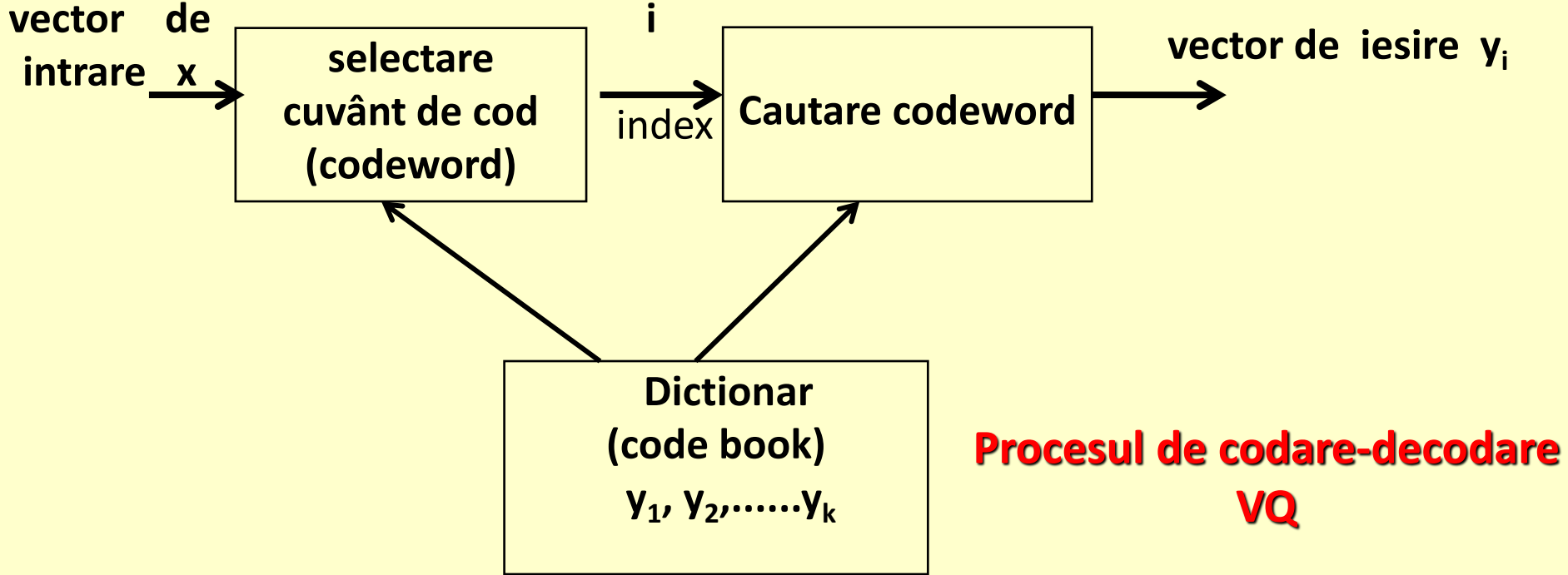
- Asociat fiecarui cuvânt de cod, y_i , este regiunea celor mai apropiati vecini numita regiune Voronoi, definita de :

$$V_i = \left\{ x \in R^k : \|x - y_i\| \leq \|x - y_j\|, \forall j \neq i \right\}$$

- Setul de regiuni Voronoi partitioneaza intreg spatiul R^k

Ex. de regiuni Voronoi bidimensionale ($k=2$) ptr. o sursa cu FDP neuniforma:
 Vectorii de intrare sunt marcati cu x, iar cuvintele de cod (codewords) cu
 cercuri rosii; regiunile Voronoi sunt separate de linii de delimitare.





Imagine

100	100	110	110
80	80	90	90
110	110		
100	100		

→ **Cod vector**

$X=(100,100,80,80)$

index	Codebook			
1	0	0	10	20
2	10	10	10	10
:			:	
k	100	100	90	90
:			:	
Nc	255	255	200	200

⇒ **Transmis k**

Problema Determinarii Dictionarului de VQ

- Pentru o secventa de antrenament compusa din $T = M > 10 \cdot N$ de vectori sursa, de dimensiune k (M suficient de mare pentru a captura proprietatile statistice ale sursei) :

$$T = \{x_i\}_{i=1}^M$$

- dat fiind numarul de vectori de cod: $N = 2^{kR}$
- Sa se afle un dictionar $Y = \{y_i\}_{i=1}^N$ si o partitie, $V = \{v_i\}_{i=1}^N$
- Pentru care sa rezulte distorsiunea medie minima:

$$D_{\text{avg}} = \frac{1}{Mk} \sum_{m=1}^M \|x_m - Q(x_m)\|^2; Q(x_m) = y_n \text{ daca } x_m \in v_n$$

Algoritmi pentru obținerea dicționarului

- Algoritmul "cu prag"
- Algoritmul LLOYD (k-means)
- Algoritmul LBG (Linde- Buzo-Gray)

Algoritmul "cu prag"

- este un algoritm simplu și rapid care nu dă un număr fixat de clase .
- fie distanța "d" pe care o folosim și $\{x_1, x_2, \dots, x_k\}$ mulțimea vectorilor de antrenare
- Primul punct, x_1 , formează singur o clasă i el va fi reprezentantul clasei, $C_1 = x_1$
- Pe parcursul funcționării clasarea punctului x_i va avea următorul efect:
- dacă distanțele $d(x_i, c_j)$ între x_i și clasele deja formate c_j (între x_i și reprezentantul c_i , centroidul său) sunt toate superioare unui prag dat "p" se creează o nouă clasă al cărei singur element $c_i = x_i$, dacă x_i se atribuie clasei celei mai apropiate se reactualizează eventual reprezentantul său.
- Algoritmul se oprește când $i = N$.

Begin

$\{x_1\} \rightarrow C_1$, $x_1 \rightarrow y_1$, $k=1$

do $i=2, N$

$j \leftarrow \text{Argmin}_{1 \leq k \leq K} d(x_i, y_k)$

if $d(x_i, C_j) > p$ atunci $k=k+1$, $\{x_i\} \rightarrow C_k$ $y_k \leftarrow x_i$,

else x_i se adaugă C_j , recalculez. y_j

end if

end do

End

Dezavantaje:

- sensibil la ordinea de sosire a datelor, distanța folosită și pragul joacă un rol primordial
- numărul de clase nu este controlabil cu precizie, fiind dependent de prag

Avantaje:

- simplu
- permite tratarea unui volum mare de date în timp de calcul rezonabil
- este posibilă întreruperea antrenării și obținerea unui dicționar și apoi refolosirea și mărirea lui în cazul în care apar noi date

Reprezentantul unei clase poate fi actualizat sau nu, pentru aceasta existând două strategii:

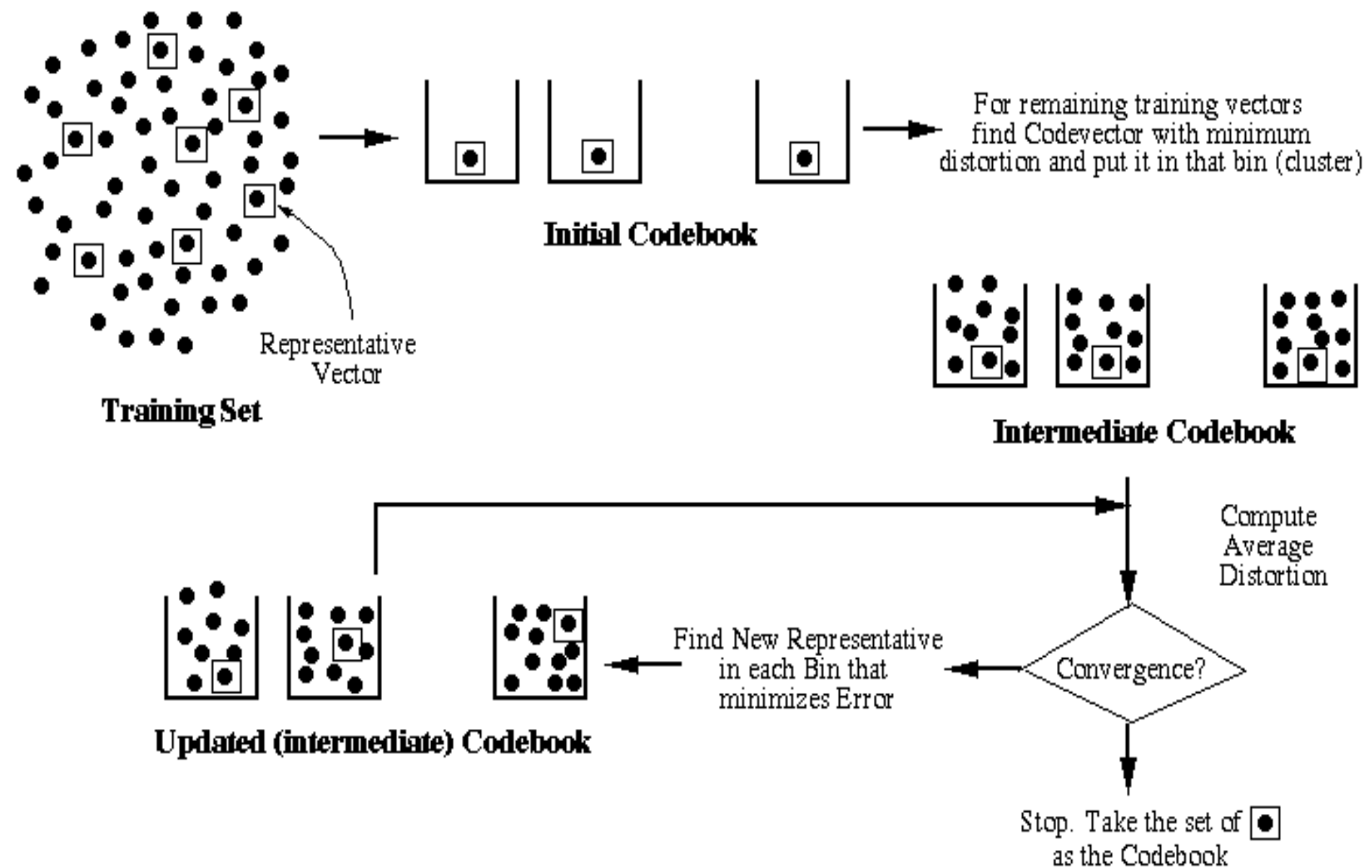
- a) reprezentantul unei clase este un element care a servit la crearea ei (această tehnică garantează distorsiunea minimă, dar poate duce la clase dispersate)
- b) reprezentantul se actualizează cu fiecare nou vector adăugat la clasă - poate fi de exemplu centrul de gravitație al clasei - dar există riscul creării de clase alungite care se pot limita printr-un prag relativ redus.

Algoritmul LLOYD (k-means)

- Acest tip de algoritm a fost studiat și generalizat pentru *clasificare automată și recunoaștere a formelor* sub numele "*k-means*" și sub o formă mult mai generalizată ca metodă numită "nori dinamici".
- Plecând de la un dicționar dat a priori, se construiesc o serie de dictionare având același număr de centroizi și diminuând distorsiunea globală la fiecare etapă, iar algoritmul se oprește când se consideră că am atins o convergență dorită.
- Nimic nu asigură că optimul găsit este global, un alt dicționar de plecare putând duce la o convergență diferită.

PROPRIETATI:

- Dicționarul se modifică numai după prezentarea întregului set de date
- MSE scade la fiecare iteratie
- Riscul de a se bloca într-un minim local este ridicat
- Dicționarul final depinde de cel initial



a) Initializarea : $m=0$, m - indicele de buclare

$\{x(n)\}$, $n=1,\dots,N$ - set antrenare .Se alege dictionarul initial de dimensiunea k :

$$\{Y_1^{(0)}, Y_2^{(0)}, \dots, Y_k^{(0)}\}$$

b) Clasificarea

- se clasifică setul de vectori de antrenare $\{x(n)\}$, $1 \leq n \leq N$ în clase $\{C_i(m)\}$ pe baza regulii proximității maxime (a vecinului cel mai apropiat)

$$x \in C_i(m) \text{ ddacă } d[x, Y_i(m)] \leq d[x, Y_j(m)] \quad \forall j \neq i \quad (1)$$

unde d reprezintă o distantă adecvată

c) Determinarea centroizilor

$m=m+1$ și se actualizează centroizii în fiecare cluster (clasă) prin recalcularea lor din vectorii de antrenare : $Y_i(m) = \text{centroiod}(C_i(m))$, $i=1,\dots,k$

$$Y_i^{(m)} = \frac{1}{\text{card}\{C_i^{(m)}\}} \sum_{x \in C_i^{(m)}} x \quad (2)$$

d) **Calculul erorii globale:**

$$D(m) = \frac{1}{N} \sum_{j=1}^M \sum_{x \in c_j(m)} d(x, Y_j) \quad (3)$$

e) **Verificarea condiției de sfârșit :**

$$\text{Dacă } \frac{D(m) - D(m-1)}{D(m)} \leq \varepsilon \quad (4) \text{ sau } m=m_0 \text{ atunci este gata,}$$

dacă nu, atunci se revine la b).

Pentru determinarea centroidului se pot folosi si alte relatii în locul lui (2).

Fiecare clasă poate fi reprezentată de axa de inerție Δc_i :

$$Y_i(m) = \inf_{y \in \Delta c_i} d(x, y) \quad (5)$$

sau o clasă se poate reprezenta prin p puncte $(y_1, y_2, \dots, y_p)$ care o minimizează:
(6)

$$y_i(m) = \frac{1}{p} \sum_{p \in [1, p]} d(x, y_p)$$

Cum se initializeaza algoritmul Lloyd?

1. Aleator in spatiul vectorilor de intrare
2. Primii Q vectori de date x^i - Se folosesc primii 2^R vectori (ca vectori centroid) in secventa de antrenare ca si dictionar initial.
3. Q vectori x^i selectati aleator din secventa de antrenament si distantati intre ei. (uneori aceasta abordare este numita generarea aleatoare a codurilor)
4. Cresterea setului initial – aplicam algoritmul cu prag pina ajungem la dimensiunea dictionarului
5. Vecinii pereche cei mai apropiati: primul dictionar e format din toate punctele x^i ; – cei mai apropiati 2 centroizi se contopesc (centru de greutate); se repeta procedura pana ajugem la dimensiunea dorita

Algoritmul LBG (Linde- Buzo-Gray)

- Algoritmul Lloyd nu generează decât un dicționar optimal local, alegerea centroizilor inițiali fiind importantă.
- Algoritmul LBG propus ne plasează în metrica euclidiană și se creează un dicționar de $N=2^p$.
- Pașii algoritmului sunt următorii:

1. Inițializarea

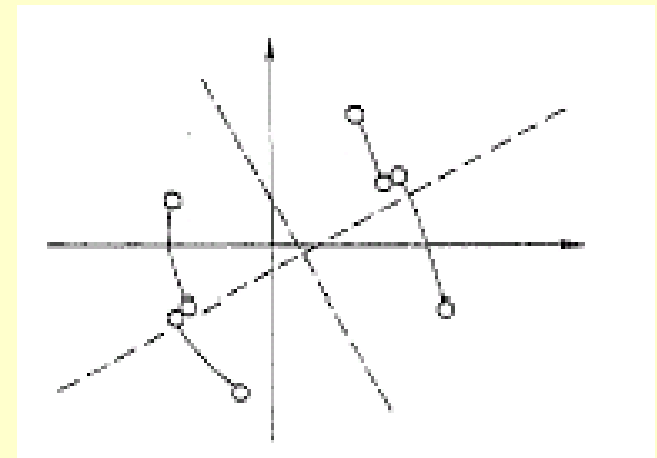
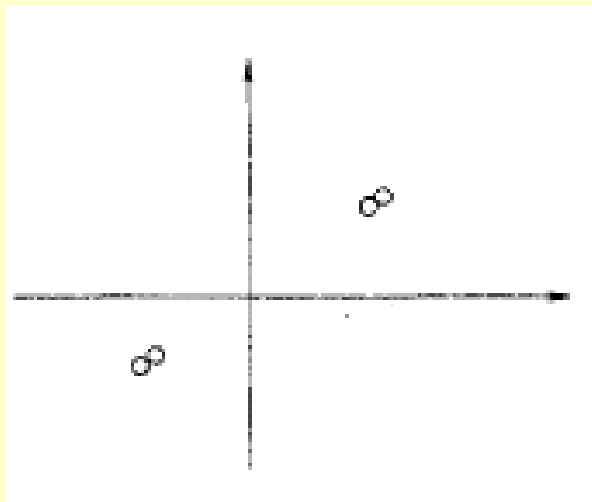
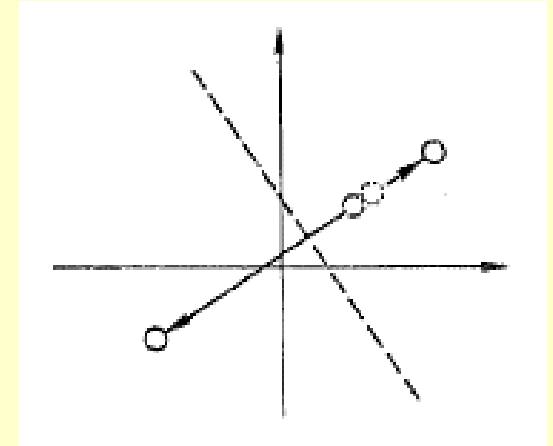
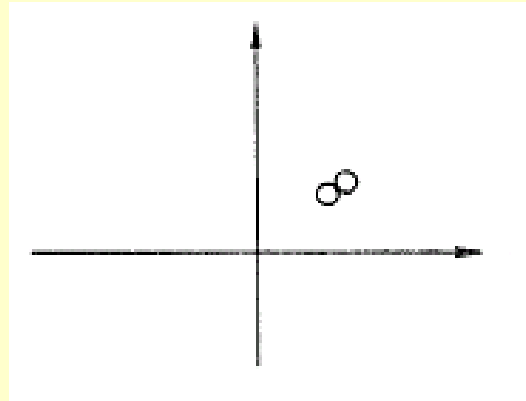
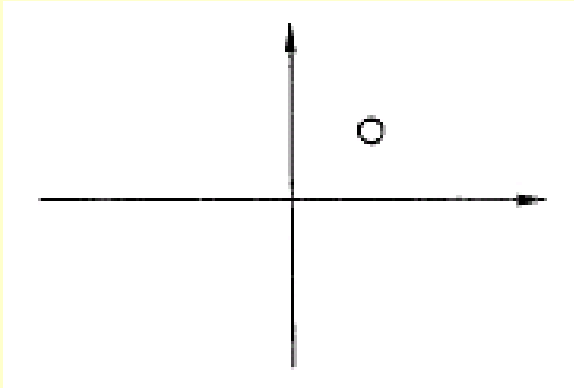
Se alege un singur centru de gravitație al mulțimii de antrenament, Y_0 , $k=0$

2. **Perturbarea centroizilor** Y_j (2^k) cu o mică valoare, transformând fiecare centroid Y_j în $Y_j + \epsilon$ și $Y_j - \epsilon$, unde ϵ este un vector dat de normă mică

3. **Convergența** - se aplică algoritmul Lloyd cu cei 2^{k+1} centroizi astfel constituiți ca dicționar de plecare

4. **Condiție sfârșit** - $k=k+1$, dacă $k > k_0$: stop, dacă nu, se revine la 2.

- LBG VQ (<http://www.data-compression.com/vqanim.html>)



Metoda LBG consolidată (Enhanced LBG)

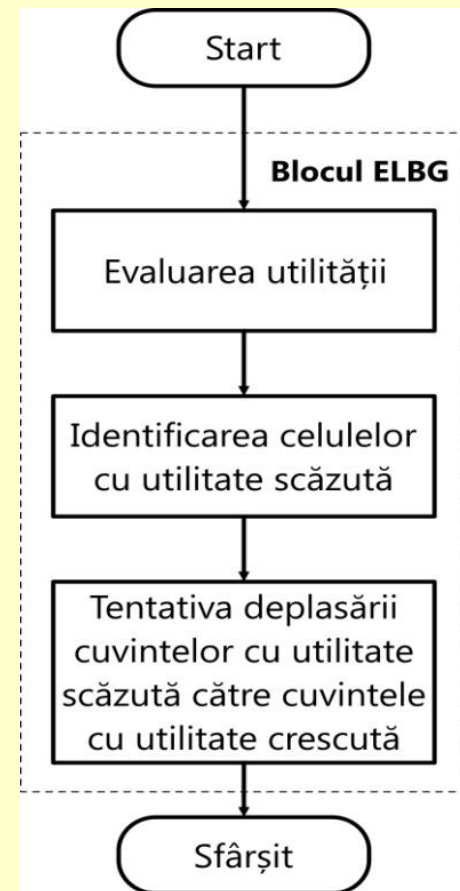
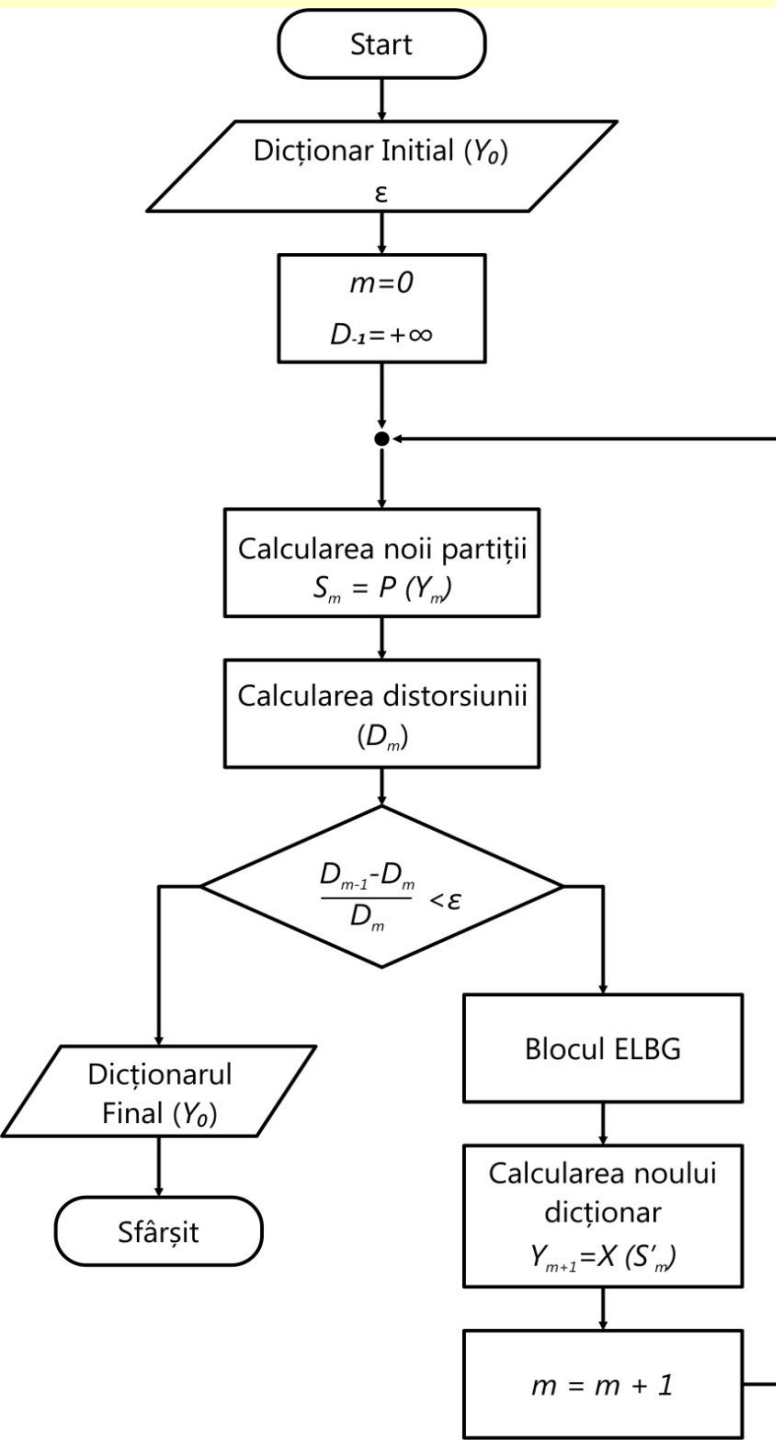
- are scopul de a soluționa dezavantajele metodei clasice LBG
- Se introduce o nouă noțiune, care poartă numele de “utilitatea unui cuvânt de cod”
- După evaluarea utilității cuvintelor de cod, se pot identifica cele cu o utilitate scăzută.
- Această informație este utilă pentru pasul următor, deplasarea cuvintelor de cod. Se încearcă deplasarea tuturor cuvintelor de cod de utilitate scăzută, care se află în preajma celor de utilitate crescută.
- Pentru un cuantizor vectorial optim, conform teoremei lui Gersho, fiecare celulă are aceeași contribuție la distorsiunea totală și astfel indicele de utilitate este 1 pentru toate celulele.
- Partea principală a algoritmului ELBG constă în secvența de executare a tentativelor de deplasare a cuvintelor de cod (SoCA's – Shifting of Codewords Attempts)

- Pentru început, se calculează distorsiunea medie pentru toate celulele:

$$D_m = \frac{1}{N} \sum_{i=1}^N D_i$$

- Se definește o măsură a utilității , care depinde de distorsiunea celulei i , D_i și normalizată cu valoarea medie D_m .

$$U_i = \frac{D_i}{D_m} , \quad i=1,N$$



Organigrama algoritmului ELBG

Modele MARKOV Ascunse (HMM)

- Bazele teoretice ale lanturilor Markov se cunosc de peste 100 de ani
- Metoda DTW se făcea o comparatie de tip formă-formă, fiind nevoie de o normalizare neliniară a esantioanelor (dilatare-contractare), pentru calculul distantei dintre forma de test si model ;
- la metodele stochastice *distanța este înlocuită cu probabilități* calculate în etapa de antrenare;
- Forma semnalului vocal este considerată ca un *semnal continuu , observabil în timp în anumite puncte de observatie*;
- Modelul descrie aceste stări cu ajutorul probabilităților de tranzitie dintr-o stare în alta si de probabilitatea de observare din fiecare stare;
- Compararea constă în căutarea în acest graf de stări a drumului cel mai probabil corespunzător unei succesiuni de evenimente observate în lantul de intrare ;
- Aceste metode sunt robuste si fiabile datorită existentei unor buni algoritmi de antrenare;
- Recunoasterea este rapidă pentru că metodele au în general putine stări si volumul de calcul este redus.

MODELE STOCHASTICE

- Un model stochastic este un proces aleator care poate schimba starea s_i , $i=1, \dots, N$, aleator în momentele $t=1,2,\dots, T$.
- În fiecare stare este o variabilă aleatoare $X(t)=s_i$, $t \in [1,T]$, care ia valori într-o multime de observatie care se face în timpul analizei fenomenului
- Evolutia sistemului este o succesiune de tranzitii de stare plecând de la $X(1)=S_1$:
$$s_1 \rightarrow s_2 \rightarrow \dots \rightarrow s_T$$
- Legea de evolutie a sistemului constă în a avea lantul de tranzitii $s_1, s_2, \dots, s_m$, $\forall m \leq T$ cu probabilitatea $P(s_1, s_2, \dots, s_T)$, care se definește din aproape în aproape, astfel :

$$\begin{aligned} P(s_1, s_2, \dots, s_T) &= P(s_1, s_2, \dots, s_{T-1}) \bullet P(s_T \mid s_1, s_2, \dots, s_{T-1}) = \\ &= P(s_1, s_2, \dots, s_{T-2}) \bullet P(s_{T-1} \mid s_1, s_2, \dots, s_{T-2}) \bullet P(s_T \mid s_1, s_2, \dots, s_{T-1}) = \\ &= P(s_1) \bullet P(s_2 \mid s_1) \bullet P(s_3 \mid s_2, s_1) \bullet \dots \bullet P(s_T \mid s_1, s_2, \dots, s_{T-1}) . \end{aligned} \quad (1)$$

- Un proces stocastic verifică proprietatea Markov dacă :

$$\forall t \quad P(X(t)=s_i \mid X(t-1)=s_j, X(t-2)=s_k, \dots) = P(X(t)=s_i \mid X(t-1)=s_j) \quad (2)$$

- Dacă notăm $X(t)=q_t$, formula devine :

$$\forall t \quad P(q_t=s_i \mid q_{t-1}=s_j, q_{t-2}=s_k, \dots) = P(q_t=s_i \mid q_{t-1}=s_j) \quad (3)$$

de unde $P(q_1, q_2, \dots, q_T) = P(q_1) \cdot P(q_2 \mid q_1) \cdot \dots \cdot P(q_T \mid q_{T-1})$.

-Evolutia este total determinată de probabilitatea inițială si de cele de tranziție succesivă

- un proces Markov astfel definit este de ordinul întâi

- Un model Markov este stationar dacă :

$$\forall t, k \quad P(q_t=s_i \mid q_{t-1}=s_j) = P(q_{t+k}=s_i \mid q_{t+k-1}=s_j) \quad (4)$$

- se definește matricea de probabilitate de tranziție $A = [a_{ij}]$:

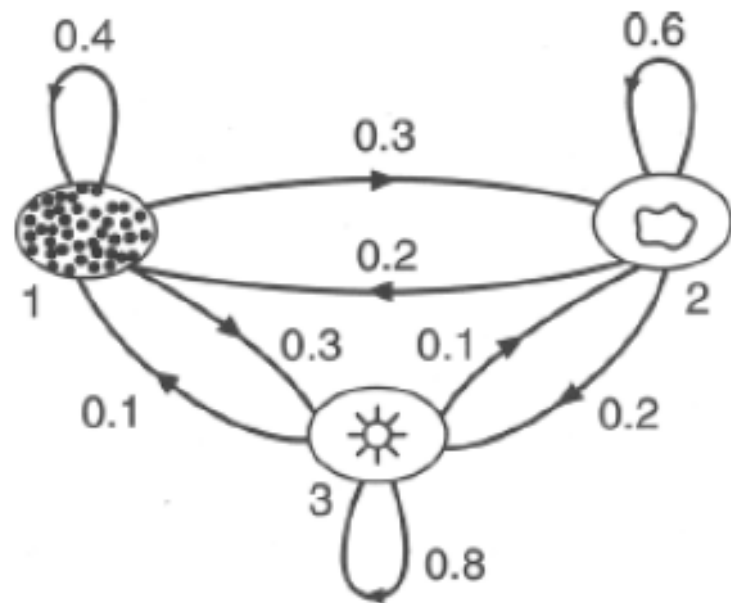
$$a_{ij} = P(q_t=s_j \mid q_{t-1}=s_i) \quad , \quad i, j = 1, \dots, N \quad (5)$$

$$\text{si } \forall i, j \quad , \quad a_{ij} \geq 0 \quad , \quad \forall i \quad \sum_{j=1}^N a_{ij} = 1 \quad .$$

Problema:

Considerand ca azi vremea e insorita (S3), care este probabilitatea (conforma modelului ca timpul in urmatoarele 7 zile sa fie “soare-soare-ploaie-ploaie-soare-noros-soare” ?

$$A = \{a_{ij}\} = \begin{bmatrix} 0.4 & 0.3 & 0.3 \\ 0.2 & 0.6 & 0.2 \\ 0.1 & 0.1 & 0.8 \end{bmatrix}$$



Solutie: Definim secventa de observatii, O, ca:

$$O = \{ S_3, S_3, S_3, S_1, S_1, S_3, S_2, S_3 \}$$

Si calculam ca $P(O|Model)$.

$$P(O|Model) = P[S_3, S_3, S_3, S_1, S_1, S_3, S_2, S_3|Model]$$

$$P(O | Model) = P[S_3, S_3, S_3, S_1, S_1, S_3, S_2, S_3 | Model]$$

$$= P[S_3]P[S_3 | S_3]^2 P[S_1 | S_3]P[S_1 | S_1]$$

$$\cdot P[S_3 | S_1]P[S_2 | S_3]P[S_3 | S_2]$$

$$= \pi_3 (a_{33})^2 a_{31} a_{11} a_{13} a_{32} a_{23}$$

$$= 1(0.8)^2 (0.1)(0.4)(0.3)(0.1)(0.2)$$

$$= 1.536 \cdot 10^{-04}$$

$$\boxed{\pi_i = P[q_1 = S_i], \quad 1 \leq i \leq N}$$

Modele Markov ascunse

(MMA = HMM - Hidden Markov Models)

- MMA este un model stochastic particular si prezintă o formă (semnal, fenomen...) prin două siruri de variabile aleatoare : unul ascuns si celălalt observabil.

Sirul ascuns corespunde succesiunii de stări $q_1, q_2, \dots, q_T, Q(1:T)$, unde q_i iau valori prin multimea de stări $\{s_1, s_2, \dots, s_N\}$

Sirul observabil corespunde succesiunii observatiilor $O_1, O_2, \dots, O_T$ notate cu $O(1:T)$, unde O_i sunt functie de timp si se găsesc într-o multime de M simboluri observabile $\{v_1, v_2, \dots, v_M\}$.

- Diferenta între un model Markov observabil si unul ascuns se poate sesiza în ex. următor, de lansare a unei monede si explicarea succesiunii de observatii față /stemă

Modelul aruncarii monedei

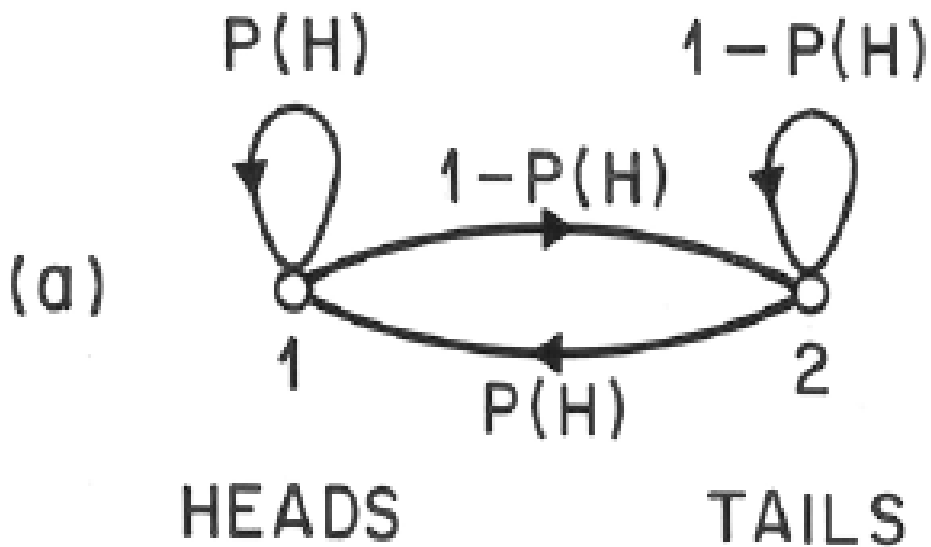
- Se efectueaza o serie de experimente de aruncare a monedei;
- Numarul de monede este necunoscut, se stiu doar rezultatele experimentelor. Seria de observatii este:

$$O = O_1 O_2 O_3 \dots O_T = HHTTTHTTTH \dots H$$

- **Problema:** sa se determine un model MMA care sa explice secventa de observatii.

Cerinte:

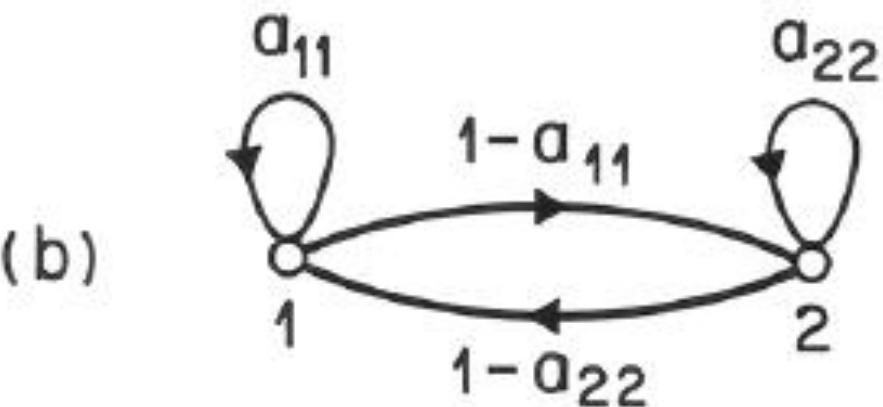
1. Care sunt starile modelului?
2. Cate stari sunt necesare?
3. Care sunt probabilitatile de tranzitie a starilor?



Modelul cu o moneda

MM observabil

O = H H T T H T H H T T H ...
S = 1 1 2 2 1 2 1 1 2 2 1 ...



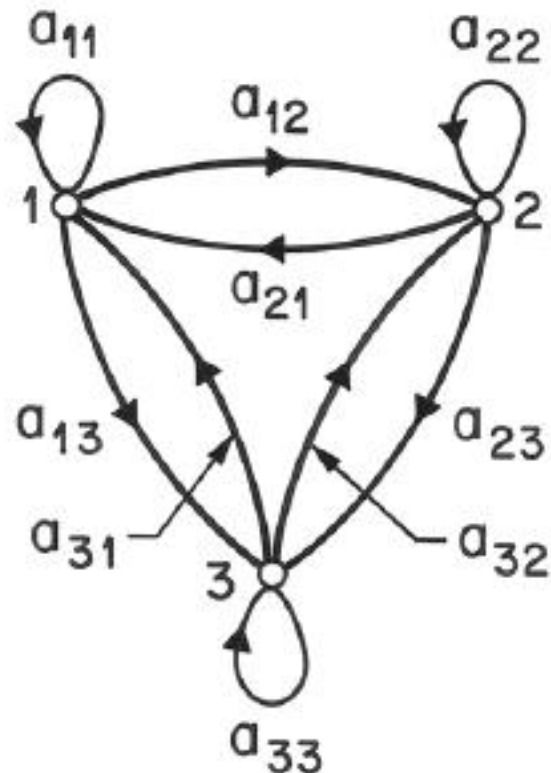
Modelul cu 2 monede

MM Ascuns

O = H H T T H T H H T T H ...
S = 2 1 1 2 2 2 1 2 2 1 2 ...

$$\begin{array}{ll} P(H) = P_1 & P(H) = P_2 \\ P(T) = 1 - P_1 & P(T) = 1 - P_2 \end{array}$$

(c)



Modelul cu 3 monede

MM Ascuns

O = H H T T H T H H T T H ...

S = 3 1 2 3 3 1 1 2 3 1 3 ...

STATE

	1	2	3
P(H)	P_1	P_2	P_3
P(T)	$1 - P_1$	$1 - P_2$	$1 - P_3$

EX. Consideram o reprezentare a experimentului aruncarii monedei printr-un model MMA (HMM) (model λ) – varianta cu 3 monede avand probabilitatile de tranzitie si cele initiale egale cu 1/3:

	Stare 1	Stare 2	Stare 3
P(H)	0.5	0.75	0.25
P(T)	0.5	0.25	0.75

Observam secventa : **O=H H H T H T T T T**

- a. Care este secventa de stari cea mai probabila? Care este probabilitatea de observare a secvenței și această succesiune de stari, cea mai probabila?
- b. Care este probabilitatea ca secvența de observare sa provina în întregime din starea 1 (numai moneda 1)?
- c. Considerand secventa de observatii: **O1= HTTHTHHTTH**, cum se modifica cerintele a, b?

SOLUTIE.

a) Avand sirul de obs. **O=HHHHTHTTTT**, cea mai probabila secventa de stari este aceea ptr. care probabilitatea obs. individuale este maxima.

Astfel ptr. fiecare H, cea mai probabila stare este S2 si pentru fiecare T cea mai probabila este starea S3.

Astfel secventa de stari cea mai probabila este:

S= S2 S2 S2 S2 S3 S2 S3 S3 S3 S3

	Stare 1	Stare 2	Stare 3
P(H)	0.5	0.75	0.25
P(T)	0.5	0.25	0.75

Probabilitatea de O si S (date de model) este:

$$P(O, S | \lambda) = (0.75)^{10} \left(\frac{1}{3} \right)^{10}$$

b) Probabilitatea ca observatiile sa provina numai din S1 adica:

S= S1 S1 S1 S1 S1 S1 S1 S1 S1 S1

$$P(O, \hat{S} \mid \lambda) = (0.50)^{10} \left(\frac{1}{3} \right)^{10}$$

	Stare 1	Stare 2	Stare 3
P(H)	0.5	0.75	0.25
P(T)	0.5	0.25	0.75

$$R = \frac{P(O, S \mid \lambda)}{P(O, \hat{S} \mid \lambda)} = \left(\frac{3}{2} \right)^{10} = 57.67$$

c) Fiind data secventa de obs. **O1=HTTHTHHTTH** avand acelasi nr. de stari ca si anterior H=T=5, rezultatele nu se schimba.

TEMA. Daca probabilitatile de tranzitie intre stari se modifica astfel:

$$\begin{aligned} a_{11} &= 0.9, & a_{21} &= 0.45, & a_{31} &= 0.45 \\ a_{12} &= 0.05, & a_{22} &= 0.1, & a_{32} &= 0.45 \\ a_{13} &= 0.05, & a_{23} &= 0.45, & a_{33} &= 0.1 \end{aligned}$$

Care sunt noile rezultate pentru cerintele a-c ale problemei ?

Elementele unui MMA: $\lambda = (A, B, \Pi)$; se definesc următoarele elemente :

N - numărul de stări reprezentate prin $S = \{s_1, s_2, \dots, s_N\}$, o stare la un moment dat este notată cu q_t ($q_t \in S$).

M - numărul de simboluri observabile în fiecare stare. Multimea de observatii este $V = \{v_1, v_2, \dots, v_M\}$
Un element O_t din V desemneaza un simbol observat la momentul t .

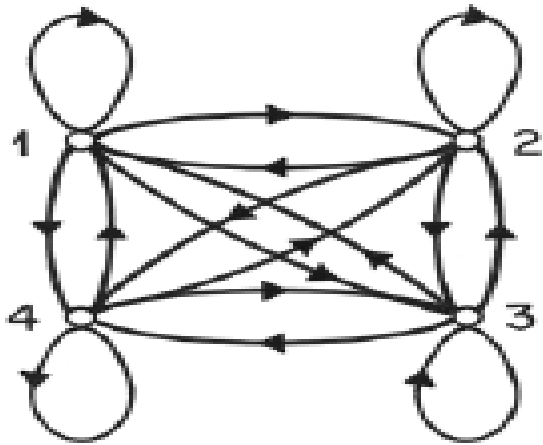
A - o matrice de probabilități de tranziție între stările modelului: a_{ij} - reprezentând probabilitatea cu care modelul evoluează din starea i spre j , în cazul general :
 $a_{ij} = A(i, j) = P(q_{t+1} = s_j \mid q_t = s_i) \quad \forall i, j \in [1, N], \forall t \in [1, T], a_{ij} \geq 0$ si $\forall i, j, \sum a_{ij} = 1$.

B – o matrice de probabilități de observare a simbolurilor în fiecare stare a modelului : $b_j(k)$ care reprezintă probabilitatea ca să observăm simbolul v_k când modelul se află în starea j :

$$b_j(k) = P(O_t = v_k \mid q_t = s_j), \quad 1 \leq j \leq N \quad 1 \leq k \leq M$$
$$b_j(k) \geq 0, \quad \forall j, k \text{ si } \sum b_j(k) = 1$$

Π - o multime a densității de probabilitate initiale : $\Pi = \{\pi_i\}, i = \overline{1, \dots, N}$, unde π_i reprezintă probabilitatea ca starea de pornire a modelului să fie i , astfel :

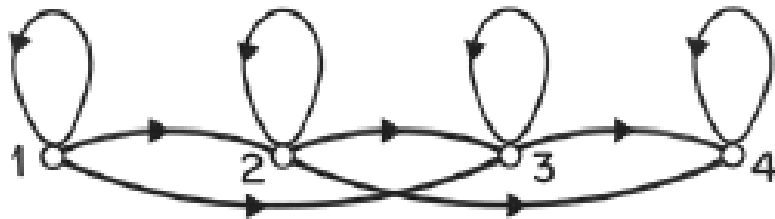
$$\pi_i = P(q_1 = s_i) \quad i = 1, \dots, N \text{ si } \pi_i \geq 0, \sum \pi_i = 1.$$



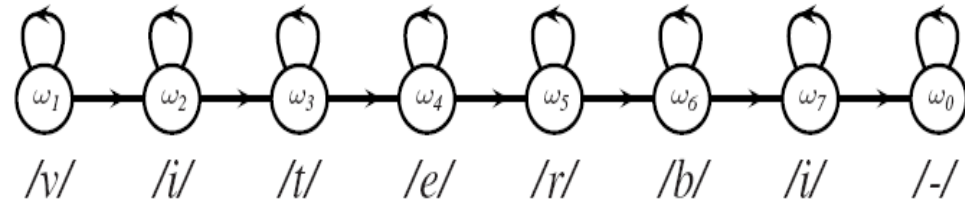
(a)

a) Ergodic

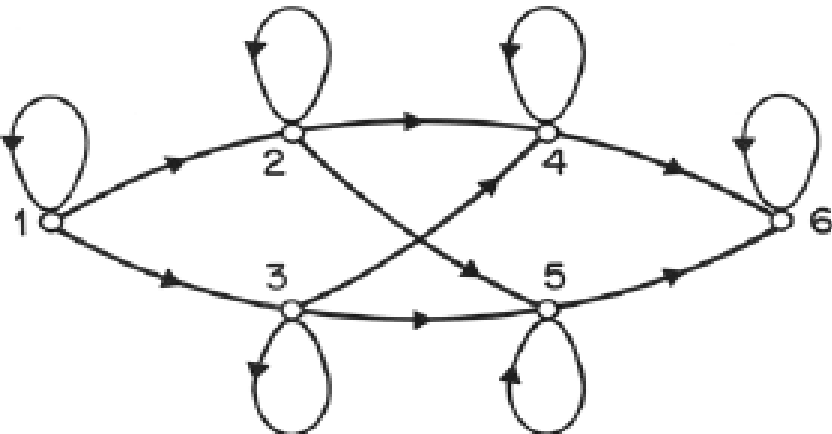
b) Stanga-dreapta



c) Mixt

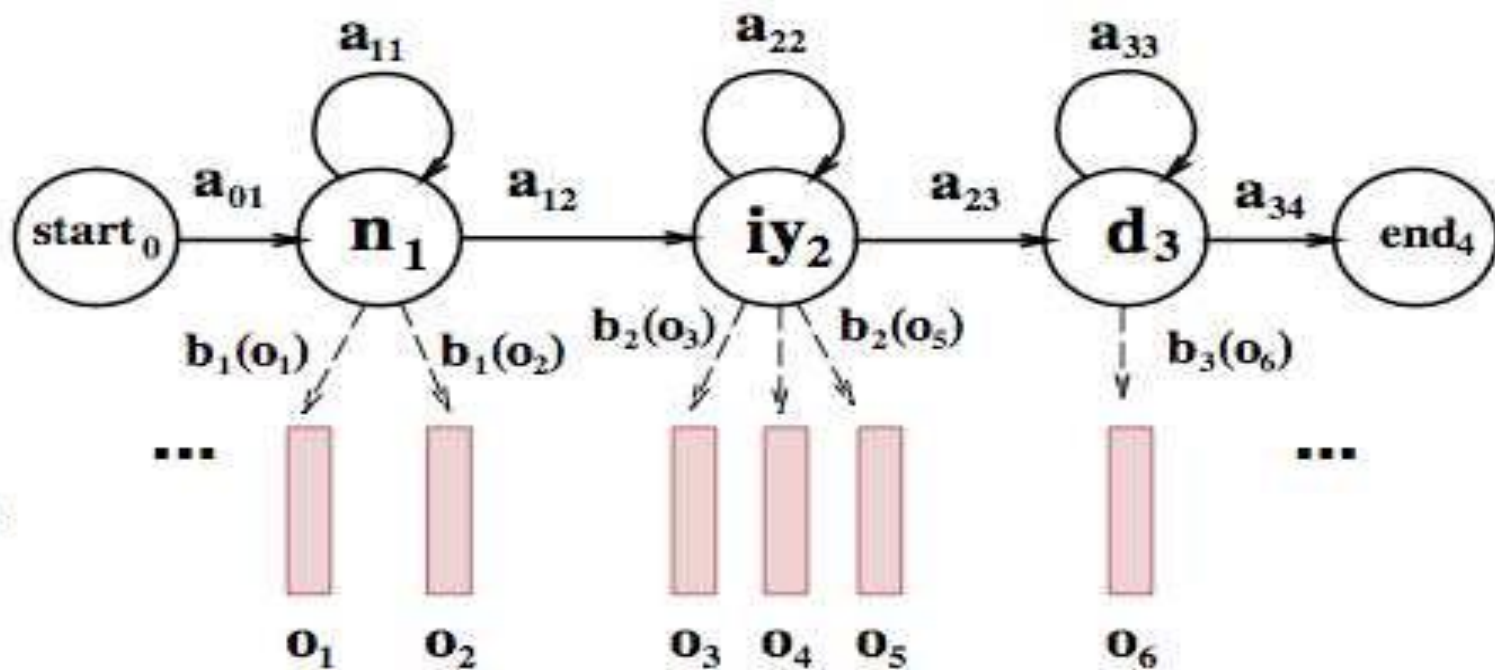


(c)



Word Model

Observation
Sequence
(spectral feature
vectors)



Procedură de generare a observatiilor

Un MMA se poate folosi ca o procedură de generare a observatiilor, iterativă cu patru etape :

1. Alegerea stării de plecare : $t=1$, se alege starea inițială $q_1=s_i$ cu π_i .
2. Observatia în starea selectată : se alege observatia $O_t=v_k$ cu $b_i(k)$;
3. Selectia stării următoare : se trece la starea următoare $q_{t+1}=s_j$ cu ajutorul lui a_{ij} .
4. Schimbarea efectivă de stare :
Schimbăm pasul: $t=t+1$, dacă $t < T$ se trece la 2 , dacă nu se termină.

Problemele de rezolvat la MMA

Sunt trei probleme de bază ale unui MMA care trebuie solutionate :

1. Evaluarea probabilităților de observare

Fiind data succesiunea de observatii: $O = O_1, O_2, \dots, O_T$ si modelul $\lambda = (A, B, \Pi)$, se pune problema determinării probabilităților $P(O|\lambda)$, ca succesiunea observatiilor să fie generată de model.

-Solutia - **Algoritmul Forward-backward**

2. Problema descoperirii secvenței de stări ascunse

- Fiind dată succesiunea observatiilor $O = O_1, O_2, \dots, O_T$ si modelul λ : să se determine secvența de stări, $Q = q_1, q_2, \dots, q_T$, cea mai probabilă din care au rezultat observatiile O_t .
- Se obtin informatii legate de partea ascunsă a modelului (succesiunea de stări).
- Se foloseste un criteriu optimal pentru descoperirea secvenței de stări

-Solutia – **Algoritmul Viterbi**

3. Problema antrenării

- Cum trebuie modificati parametrii modelului $\lambda = (A, B, \Pi)$ pentru a maximiza $P(O|\lambda)$ - reprezintă problema antrenării modelului.

-Solutia – **algoritmul Baum-Welch**

REFERINTE

- L. R. Rabiner, "A Tutorial on Hidden Markov Models and Selected Applications in Speech Recognition," Proc. of the IEEE, Vol.77, No.2, pp.257--286, 1989. Available from <http://ieeexplore.ieee.org/iel5/5/698/00018626.pdf?arnumber=18626>
- Hidden Markov Model Toolkit (HTK)
<http://htk.eng.cam.ac.uk/>
- Hidden Markov Model (HMM) Toolbox for Matlab
<http://www.cs.ubc.ca/~murphyk/Software/HMM/hmm.html>