

## FÖRELÄSNING 1 SPRÅKTEKNOLOGI

# STATISTISKA METODER 2

- MARKOVPROCESSER
- HMM - GÖMDA MARKOVMODELLER
- INFORMATIONSTEORI (ENTROPI)
- TILLÄMPNING: ORDPREDIKTION
  - UTVÄRDERING (PERPLEXITET)
  - MARKOVMODELL
  - ORDFREKVENSER, ORDKLASSTATISTIK
  - HEURISTISKA FÖRBÄTTRINGAR
  - LAGRING AV DATA

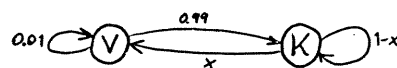
## MARKOVPROCESS

EXEMPEL: STOKASTISK PROCESS  $X_1, X_2, \dots$  SOM LÄSER EN BOKSTAV I TAGET FRÅN ETT ORD OCH AVGÖR OM DEN ÄR EN VOKAL ELLER KONSONANT.

FREKVENSRÄKNING I KORPUS GER:

$$P[X_i = \text{vokal}] = 0.36 \text{ och } P[X_i = \text{vokal} | X_{i-1} = \text{vokal}] = 0.01$$

MODELLERA MED MARKOVKEDJA:



$$(P_{ij}) = \begin{pmatrix} 0.01 & 0.99 \\ x & 1-x \end{pmatrix}$$

$$P[X_i = v] = 0.36, P[X_i = k] = 0.64$$

$$\begin{pmatrix} 0.36 & 0.64 \end{pmatrix} \begin{pmatrix} 0.01 & 0.99 \\ x & 1-x \end{pmatrix} = \begin{pmatrix} 0.36 & 0.64 \end{pmatrix}$$

$$\Rightarrow 0.36 \cdot 0.01 + 0.64x = 0.36 \Rightarrow x = \frac{0.36(1-0.01)}{0.64} \approx 0.56$$

EN MARKOVPROCESS HAR INGET MINNE.

MODELLEN ÄR DÄRFÖR LÄMPAD BARA OM SANNOLIKHETEN FÖR VOKAL INTE PÅVERKAS AV MER ÄN FÖREGÅENDE BOKSTAV, OCH DET GÖR DEN!

## HMM - GÖMDA MARKOVMODELLER

GIVET: • TILLSTÄND  $q_1, \dots, q_n$

• ÖVERGÅNGSSANNOLIKHETSMATRIIS  $P = (p_{ij})$

DÄR  $p_{ij} = P[\text{ÖVERGÅNG FRÅN } q_i \text{ TILL } q_j]$

• STARTSANNOLIKHETER  $v = (v_1, \dots, v_n)$

• SIGNALER  $\sigma_1, \dots, \sigma_m$

• SIGNALMATRIIS  $A = (a_{ij})$

DÄR  $a_{ij} = P[\sigma_j \text{ GENERERAS} \mid \text{TILLSTÄNDET ÄR } q_i]$

EN HMM ÄR TILLSTÄNDEN OSYNLIGA OCH BARA SIGNALERNA KAN OBSERVERAS.

EXEMPEL: TAGGNING AV EN MENING

TILLSTÄND  $\leftrightarrow$  TAGGAR, SIGNALER  $\leftrightarrow$  ORD

PROBLEM: HITTA TROLIGASTE PROCESSEN (TILLSTÄNDSFÖLJEN) SOM GENERERAR MENINGEN

# INFORMATIONSTEORI

INFORMATIONEN HÖR EN SIGNAL  $\sigma$  SOM HAR SANNOLIKHETEN  $p$  ÄR  $I_\sigma = \log_2 \frac{1}{p}$

EXEMPEL:

$$P[\sigma] = 0.25 \Rightarrow I_\sigma = \log_2 \frac{1}{0.25} = \log_2 4 = 2$$

$$P[\sigma] = 1 \Rightarrow I_\sigma = \log_2 \frac{1}{1} = \log_2 1 = 0$$

INFORMATIONSMÄNGDEN MÄTS I BITAR.

VANLIG SIGNAL  $\Rightarrow$  LITEN INFORMATIONSMÄNGD

# ENTROPIN

ENTROPIN ÄR VÄNTEVÄRDET FÖR INFORMATIONEN:

$$H[X] = E\left[\log_2 \frac{1}{p}\right] = \sum p_i \log_2 \frac{1}{p_i} = -\sum p_i \log_2 p_i$$

DÄR  $X$  ÄR EN STOKASTISK VARIABEL,  $P[X=\sigma_i] = p_i$

EXEMPEL:

|                                                 |                 |        |
|-------------------------------------------------|-----------------|--------|
| Låt $Y = \begin{cases} a \\ b \\ c \end{cases}$ | med sannolikhet | $0.5$  |
|                                                 | "               | $0.25$ |
|                                                 | "               | $0.25$ |

$$H[Y] = 0.5 \cdot \log_2 \frac{1}{0.5} + 0.25 \cdot \log_2 \frac{1}{0.25} + 0.25 \cdot \log_2 \frac{1}{0.25} = 0.5 + 0.25 \cdot 2 + 0.25 \cdot 2 = 1.5$$

ENTROPIN ÄR HÖGST DÄR ALLA  $n$  SIGNALER ÄR LIKA SANNOLIKA:

$$H_{\max} = \sum_{i=1}^n \frac{1}{n} \log_2 \frac{1}{1/n} = n \cdot \frac{1}{n} \log_2 n = \log_2 n$$

EXEMPEL:

$$n=3 \Rightarrow H_{\max} = \log_2 3 \approx 1.58$$

REDUNDANSEN ÄR DEN RELATIVA SKILLNADEN MELLAN  $H_{\max}$

och  $H$ :  $R = \frac{H_{\max} - H}{H_{\max}} = 1 - \frac{H}{H_{\max}}$

EXEMPEL:  $R = 1 - \frac{1.5}{1.58} \approx 0.051 = 5.1\%$

## EXEMPEL: ENTROPIN FÖR SVENSKA

ENTROPIN FÖR SVENSKA BOKSTÄVER ÄR

$$H = \sum_{a=A}^{\bar{o}} P[a] \cdot \log_2 \frac{1}{P[a]} = 4.34$$

OM ALLA BOKSTÄVER VAR LIKA VANLIGA VORE ENTROPIN

$$H_{\max} = \log_2 29 \approx 4.86$$

$\Rightarrow$  REDUNDANSEN ÄR  $1 - \frac{4.34}{4.86} \approx 11\%$

FÖRSÖK TA HÄNSYN TILL ATT BOKSTÄVERNA INTE FÖREKOMMER HELT SLUMPMÄSSIGT I SVENSKA ORD.

ENTROPIN FÖR SVENSKA ORD ÄR

$$H = \sum_{x \in \{\text{SVENSKA ORD}\}} P[x] \cdot \log_2 \frac{1}{P[x]} = 8.3$$

MEDELORDLÄNGDEN = 5.3 BOKSTÄVER

ENTROPIN UTSLAGEN PÅ BOKSTÄVERNA ÄR  $\frac{8.3}{5.3} \approx 1.6$

REDUNDANSEN BLIR DÄR  $1 - \frac{1.6}{4.86} \approx 67\%$

# ORDPREDIKTION

EXAMEN AV  
JOHAN CARLBERGER  
1997

PROBLEM: GIVET ORD  $w_1, w_2, \dots, w_{i-1}$

GISSA ORDET  $w_i$

IDÉ: HITTA  $X$  SOM MAXIMERAR  $P[X=w_i | w_1, w_2, \dots, w_{i-1}]$

## UTVÄRDERING AV EN ORDPREDIKTIONSMETOD:

INFORMATIONSTEORETISK:

PERPLEXITETEN FÖR  $X$  ÄR

$$\text{PERP}[X] = 2^{\text{H}[X]}$$

MOTSVARAR ANTALET ORD PREDIKTÖRN HAR ATT VÄLJA MELLAN.

OM ALLA N ORD ÄR LIKA SANNOLIKA BLIR

$$\text{PERP}[X] = 2^{\log_2 N} = N.$$

MÅL: MINIMERA PERPLEXITETEN.

PRAKTISK:

MÄT ANTALET TANGENTTRYCKNINGAR SOM BEHÖVS FÖR  
ATT GENERERA EN TESTTEXT MED ORDPREDIKTÖRN.

MÅL: MAXIMERA ANTALET SPARADE TANGENTTRYCKNINGAR

TESTA INTE ETT PROGRAM PÅ SAMMA  
TEXT SOM DET HAR TRÄNATS PÅ!

# STATISTIKINSAMLING

SAMLA IN FÖLJANDE STATISTIK FRÅN EN STOR KORPUS:

|                        |                          |
|------------------------|--------------------------|
| ORDFREKVENSER          | $C_o(w)$                 |
| ORDBIGRAMSFREKVENSER   | $C_{ob}(w_1, w_2)$       |
| TAGGFREKVENSER         | $C_t(t)$                 |
| TAGGBIGRAMSFREKVENSER  | $C_{tb}(t_1, t_2)$       |
| TAGGTRIGRAMSFREKVENSER | $C_{tob}(t_1, t_2, t_3)$ |
| ORDTAGGPARSFREKVENSER  | $C_{ot}(w, t)$           |
| ANTAL ORD I KORPUSEN   | $C$                      |

BERÄKNA SEDAN (OCH LAGRA):

$$f_o(w) = \frac{C_o(w)}{C}, \quad f_{ob}(w_2 | w_1) = \frac{C_{ob}(w_1, w_2)}{C_o(w_1)}$$

$$f_t(t) = \frac{C_t(t)}{C}, \quad f_{tb}(t_2 | t_1) = \frac{C_{tb}(t_1, t_2)}{C_t(t_1)}, \quad f_{tob}(t_3 | t_1, t_2) = \frac{C_{tob}(t_1, t_2, t_3)}{C_{tb}(t_1, t_2)}$$

$$f_{ot}(t | w) = \frac{C_{ot}(w, t)}{C_o(w)}$$

## TVÅ MARKOVMODELLER

### TAGGMODELLEN

ANDRA ORDNINGENS MARKOVMODELL FÖR TAGGAR.

TILLSTÄNDEN BESKRIVS AV PAR AV TAGGAR  $(t_1, t_2)$ .

$$P_{(t_1, t_2), (t_2, t_3)} = P[t_3 | t_1, t_2] = q_0 f_{ttt3}(t_1, t_2, t_3) + q_1 f_{ttt}(t_2 | t_1, t_2) + q_2 f_t(t_3)$$

DÄR  $q_0, q_1, q_2$  ÄR LÄMPLIGT VALDA KONSTANTER

### ORDMODELLEN

FÖRSTA ORDNINGENS MARKOVMODELL FÖR ORD.

TILLSTÄNDEN BESKRIVS AV ORDEN I KORPUSEN.

$$P_{w_1, w_2} = P[w_2 | w_1] = r_0 f_{oo}(w_2 | w_1) + r_1 f_o(w_2) + (r_2 f_{ob}(w_1, w_2) + r_3 f_o(w_2)) \cdot P_o(w_1)$$

DÄR  $r_0, r_1, r_2, r_3$  ÄR LÄMPLIGT VALDA KONSTANTER OCH

$$P_o(w_2) = f_{ot}(t_2 | w_2) \cdot P[t_2 | t_0, t_1]$$

# LATMANHASHNING

HASHA BARA PÅ DOM TRE FÖRSTA BOKSTÄVERNA I SÖKNYCKELN. ANVÄND SEDAN BINÄRSÖKNING.

LÄMPLIGT FÖR SÖKNING MED FÅ DISKACCESSER I STOR TEXT NÄR INDEXET INTE KAN LIGGA I PRIMÄRMINNET.

EXEMPEL: INDEXERA STORT SVENSK-ENGELSKT LEXIKON.

GIVET: • STOR FIL  $L$  MED HELA LEXIKONTEXTEN.  
• INDEX  $I$  DÄR VARJE RAD ÄR: SÖKORD POSITION I  $L$

SKAPA SÖKINDEX: SORTERA  $I$  MED SORT.

SKAPA EN INDEXARRAY  $A[abc]$  SOM FÖR VARJE  $abc$  ANGER VAR I  $I$  FÖRSTA ORDET SOM BÖRJAR SÅ FINNS.

FÖRBEREDELSE INFÖR SÖKNINGAR:

LÄS IN  $A$  FRÅN FIL. ÖPPNA FILERNA  $I$  OCH  $L$ .

SÖKNINGSALGORITM( $w$ ) =

```
wprefix ← FÖRSTA 3 BOKST. I W
i ← A[wprefix]
j ← A[wprefix+1]
WHILE j-i > 1000 DO
    m ← ⌊(i+j)/2⌋
    GÅ TILL POSITION m I I
    LÄS IN NÄSTA ORD FRÅN I I S
    IF S ≤ w THEN i ← m
    ELSE j ← m
```

```
GÅ TILL POSITION i I I
WHILE TRUE DO
    LÄS IN NÄSTA ORD FRÅN I I S
    IF S = w THEN
        LÄS IN POSITIONEN I X
        RETURN X
    IF S > w THEN
        RETURN NOTFOUND
```

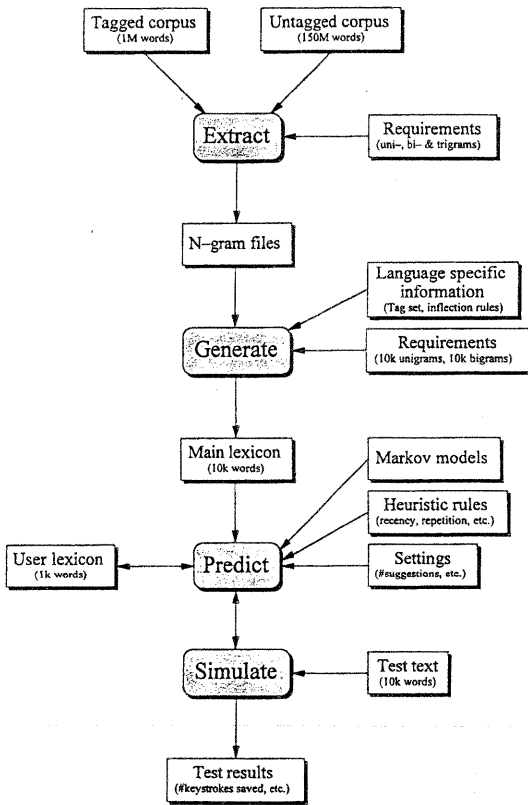


Figure 3. A system for word predictor construction.

## Viterbis algoritmen

Anta att vi vet sannolikheten  $P(x, y)$  för att ordklassen  $x$  ska följas av ordklassen  $y$  i en mening. Vi har också två speciella ordklasser *start* för meningens början och *slut* för meningens slut.  $P(\text{start}, x)$  är då sannolikheten för att en mening inleds med ordklassen  $x$  och  $P(y, \text{slut})$  är sannolikheten för att en mening avslutas med ordklassen  $y$ .

Anta också att vi har ett lexikon som talar om vilka ordklasser varje ord kan ha. Exempelvis säger lexikonet att boken bara kan vara substantiv men att gula både kan vara adjektiv och substantiv.

Den troligaste ordklassatagningen av en mening är den taggning som maximerar produkten av dom ingående sannolikheterna.

Indata består av en mening med  $n$  ord. Det finns  $m$  stycken ordklasser (inklusive *start* och *slut*).

Viterbis algoritmen löser detta problem i linjär tid i antalet ord  $n$  med hjälp av algoritmkonstruktionsmetoden dynamisk programmering.

Låt  $M[i, j]$  vara lösningen (maximala produkt sannolikheten) för delproblemet att tagga meningen fram till ord  $i$  så att ord  $i$  får taggen  $j$ .  $M[i, j]$  kan definieras rekursivt som

$$M[i, j] = \begin{cases} 1 & \text{om } i = 0 \text{ och } j = \text{start}, \\ 0 & \text{om } i = 0 \text{ och } j \neq \text{start}, \\ \max_{1 \leq k \leq m} M[i-1, k] \cdot P[k, j] & \text{annars,} \end{cases}$$

Vilket  $k$ -värde som ger maximum i rekursionen lagras i  $\text{father}[i, j]$  så att algoritmen ska kunna följa den bästa taggningen tillbaka från *slut* till *start*. Den bästa taggningen lagras sedan i  $\text{POS}[1..n]$ .

```
Indata:  $P[1..m, 1..m]$ ,  $w[1..n]$ , Lexikon: ord  $\rightarrow$  set of integer
for  $j \leftarrow 1$  to  $m$  do  $M[0, j] \leftarrow 0$ 
 $M[0, \text{start}] \leftarrow 1$ 
for  $i \leftarrow 1$  to  $n+1$  do
    if  $i \leq n$  then  $S \leftarrow \text{Lexikon}(w[i])$  else  $S \leftarrow \{\text{slut}\}$ 
    for  $j \leftarrow 1$  to  $m$  do  $M[i, j] \leftarrow 0$ 
    foreach  $j \in S$  do
        for  $k \leftarrow 1$  to  $m$  do
            tmp  $\leftarrow M[i-1, k] \cdot P[k, j]$ 
            if  $\text{tmp} > M[i, j]$  then  $M[i, j] \leftarrow \text{tmp}$ ,  $\text{father}[i, j] \leftarrow k$ 
```

```
c ← slut
for  $i \leftarrow n$  downto 1 do
     $\text{POS}[i] \leftarrow c \leftarrow \text{father}[i+1, c]$ 
    for  $i \leftarrow 1$  to  $n$  do
        write  $w[i]$ ,  $\text{POS}[i]$ 
```

I värsta fallet kan varje ord ha  $m$  ordklasser. Dom tre nästlade for-looparna går då  $n+1$ ,  $m$  respektive  $m$  varv, vilket ger tidskomplexiteten  $O(m^3)$ .