

## Učení, adaptace, ladění

Módelky se podaří napoprvé konstruovat složitější síť správně.

Máme-li k dispozici dostatečně velké soubory dat (database of cases), můžeme učením stanovit strukturu sítě i potenciály (tj. pravděpodobnosti příslušné uzlům)

Dostaneme-li dodatečně další data, můžeme adaptaci modifikovat síť tak, aby lépe reprezentovala nové zkušenosti.

Ladění je proces modifikace parametrů (tj. vstupních pravděpodobností) tak, aby byl splněn nějaký explicitní požadavek na aposteriorní pravděpodobnosti.

(Je to podobné jako trénování neuronové sítě)

Tyto procesy jsou dosud předmětem aktivního výzkumu a vývoje, takže zatím nejsou standardní součástí Bayesovského softwaru kromě jiného i proto, že i když je třeba metoda známa, tak má příliš velkou časovou a prostorovou složitost.

### Účení - příklad

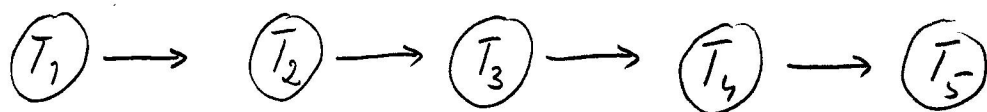
Vytváříme síť pro slova = posloupnosti symbolů  $a, b$  délky 5

Z frekvencí analýzy je známo, že :

| první 2<br>písmena | poslední 3 písmena |       |       |       |       |       |       |       |
|--------------------|--------------------|-------|-------|-------|-------|-------|-------|-------|
|                    | aaa                | aab   | aba   | abb   | baa   | bab   | bb a  | bbb   |
| aa                 | 0.017              | 0.021 | 0.019 | 0.019 | 0.045 | 0.068 | 0.045 | 0.068 |
| ab                 | 0.033              | 0.04  | 0.037 | 0.038 | 0.011 | 0.016 | 0.010 | 0.015 |
| ba                 | 0.011              | 0.014 | 0.010 | 0.01  | 0.031 | 0.046 | 0.031 | 0.045 |
| bb                 | 0.050              | 0.060 | 0.056 | 0.057 | 0.016 | 0.023 | 0.015 | 0.023 |

Hledáme model pro proměnné - znaky  $T_1, T_2, T_3, T_4, T_5$

Jako první náš asi napadne:



označíme ho  $M_{Simp}$

V tomto modelu počítáme podle řetězového pravidla

$$P^*(\bar{T}_1, T_2, \bar{T}_3, \bar{T}_4, T_5) = P(\bar{T}_1) P(T_2 | \bar{T}_1) P(\bar{T}_3 | T_2) P(\bar{T}_4 | \bar{T}_3) P(T_5 | \bar{T}_4)$$

Zřejmě  $P(\bar{T}_1 = a) = P(T_1 = b) = \frac{1}{2}$

$$P(\bar{T}_2 | \bar{T}_1) = \frac{P(\bar{T}_2, \bar{T}_1)}{P(\bar{T}_1)}$$

např.  $P(T_2 = a, T_1 = a) = P(aa) = \text{součet pravděpodobností,}$

$\text{v řádku aa ve frekvencní tabulce} = 0.3$

$$P(T_2 = a | T_1 = a) = \frac{0.3}{0.5} = 0.6$$

podobně  $P(\bar{T}_2 = a | \bar{T}_1 = b) = 0.4$

$$P(\bar{T}_2 = b | \bar{T}_1 = a) = 0.4$$

$$P(\bar{T}_2 = b | \bar{T}_1 = b) = 0.6$$

Analogicky se napočítají další pravděpodobnosti:

$$P(\bar{T}_2 = a) = 0.5$$

(= součet 1. a 3. řádku frekvencní tabulky)

$$P(\bar{T}_3 = a, \bar{T}_2 = a) = \text{součet prvků 1.-4. sloupců}$$

$$\text{pro řádky 1. a 3.} = 0.121$$

$$P(\bar{T}_3 = a | \bar{T}_2 = a) = \frac{0.121}{0.5} = 0.24$$

$$P(\bar{T}_3 = a \mid \bar{T}_2 = b) = 0.74$$

$$P(\bar{T}_3 = b \mid \bar{T}_2 = a) = 0.76$$

$$P(\bar{T}_3 = b \mid \bar{T}_2 = b) = 0.26$$

atd.

Z toho nám vyjde tabulka psti'  $P^*$ :

| první 2<br>znaky | poslední 3 znaky |       |       |       |       |       |       |       |
|------------------|------------------|-------|-------|-------|-------|-------|-------|-------|
|                  | aaa              | aab   | aba   | abb   | baa   | bab   | bba   | bbb   |
| aa               | 0.016            | 0.023 | 0.018 | 0.021 | 0.044 | 0.067 | 0.05  | 0.061 |
| ab               | 0.03             | 0.044 | 0.033 | 0.041 | 0.011 | 0.015 | 0.012 | 0.014 |
| ba               | 0.01             | 0.016 | 0.012 | 0.014 | 0.029 | 0.045 | 0.033 | 0.041 |
| bb               | 0.044            | 0.067 | 0.052 | 0.061 | 0.016 | 0.023 | 0.017 | 0.021 |

Na pohled je vidět, že tabulky nejsou totožné.

Model  $M_{\text{Simp}}$  můžeme přijmout, pokud se moc neliší - k tomu

ale potřebujeme nějakou míru vzdálenosti - je více

možností, ale vezměme obyčejnou Euklidovskou vzdálenost

$$\text{mezi vektory } x \text{ a } y : \text{dist}_Q(x, y) = \sum_i (x_i - y_i)^2$$

(kvadratickou, bez odmocňování)

dostaneme  $\text{dist}_Q(P, P^*) = 0.00033\%$

Je to hodně nebo málo? Potřebujeme ještě nějakou hranici pro rozhodování.

Předpokládáme, že pro zádne 2 pravděpodobnosti necháme rozdíl větší než 0.004,

Pak rozumná hranice pro rozdíl tabulek je

$$32 \times 0.004^2 = 0.000512,$$

takže model  $M_{\text{Simp}}$  je z tohoto hlediska akceptovatelný.

Máme ale i jiná kritéria:

sítě nesmí být moc velké, aby se s ní dalo rozumně rychle (a v rozumném prostoru počítat).

Definujeme velikost sítě  $M$  jako  $\text{Size}(M) = \sum_{A \in U} S_p(A)$ ,

kde  $U = \{A_1, \dots, A_m\}$  jsou všechny proměnné vsiti,

$\text{pa}(A)$  jsou rodiče  $A$

$S_p(A)$  je počet prvků (vstupů) v tabulce pro  $P(A|\text{pa}(A))$

Máme  $\text{Size}(M_{\text{Simp}}) = 2 + 4 + 4 + 4 + 4 = 18.$

Definujeme míru akceptace jako

$$\text{Acc}(P, M^*) = \text{Size}(M^*) + k \cdot \text{dist}(P, P^*)$$

kde  $P^*$  je pravděpodobnost příslušná modelu  $M^*$

$k$  je kladné reálné číslo

Např. pro  $k = 10000$  dostaneme pro model  $M_{\text{Simp}}$

$$\text{Acc}(M_{\text{Simp}}) = 21.37$$

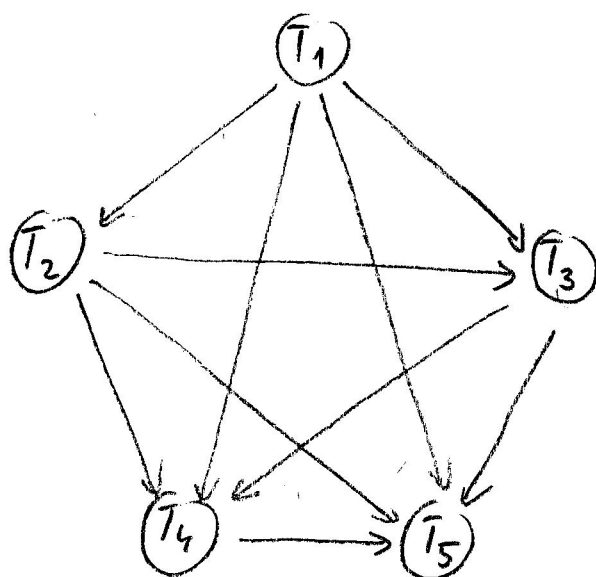
Úlohou teď bude najít model, který minimalizuje  $\text{Acc}$

při hranici vzdálenosti 0.0005 (s naší konstantou  $k = 10000$ ).

V principu bychom měli vyhodnotit všechny možné orientované grafy s vrcholy  $T_1, T_2, T_3, T_4, T_5$ . Je přirozené se omezit na případ, kdy hrana z  $T_i$  do  $T_j$  vede jen když  $i < j$ .

Začneme s maximálním tokovým grafem a postupně budeme ubírat hrany, ale nesmíme se přitom dostat přes hranici vzdálenosti.

Model  $M_{Max}$ :



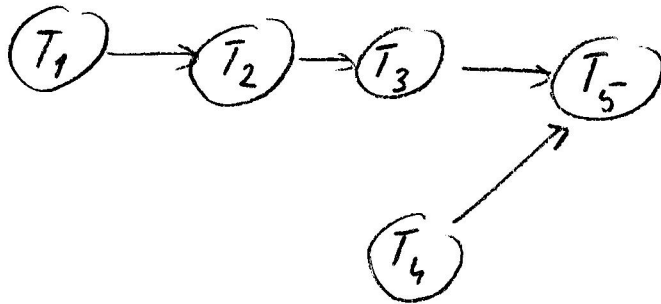
Označme  $P(\text{word} | M_{Max})$  pravdepodobnostni rozdeľeni' určene' modelom  $M_{Max}$ . Platí  $P(\text{word} | M_{Max}) = P(\text{word})$

$$\begin{aligned}
 \text{Dle: } P(\text{word}) &= P(T_1, T_2, T_3, T_4, T_5) = \\
 &= P(T_5 | T_1, T_2, T_3, T_4) P(T_1, T_2, T_3, T_4) = \\
 &= P(T_5 | T_1, T_2, T_3, T_4) P(T_4 | T_1, T_2, T_3) P(T_1, T_2, T_3) = \\
 &= P(T_5 | T_1, T_2, T_3, T_4) P(T_4 | T_1, T_2, T_3) P(T_3 | T_1, T_2) \times \\
 &\quad P(T_1, T_2) \\
 &= P(T_5 | T_1, T_2, T_3, T_4) P(T_4 | T_1, T_2, T_3) P(T_3 | T_1, T_2) P(T_2 | T_1) \\
 &\quad \times P(T_1) \\
 &= P(\text{word} | M_{Max})
 \end{aligned}$$

tedy distance tohoto modelu od skutečnosti je nulová

$$Acc(M_{Max}) = Size(M_{Max}) = 2 + 4 + 8 + 16 + 32 = 62$$

Výsledkem postupného ubírání hran bude graf  $M_{Min}$ :



$$P^*(T_1, T_2, T_3, T_4, T_5) = P(T_1)P(T_4)P(T_2|T_1)P(T_3|T_2)P(T_5|T_3, T_4)$$

Výsledná tabulka pravděpodobnosti bude

musíme  
deprecitovat  
%

| první 2<br>znaky | poslední 3 znaky |              |              |       |              |              |             |       |
|------------------|------------------|--------------|--------------|-------|--------------|--------------|-------------|-------|
|                  | aaa              | aab          | aba          | abb   | baa          | bab          | bba         | bbb   |
| aa               | 0.014            | 0.021        | 0.019        | 0.019 | 0.045        | 0.068        | 0.045       | 0.068 |
| ab               | <u>0.034</u>     | 0.04         | 0.034        | 0.038 | <u>0.01</u>  | <u>0.015</u> | 0.01        | 0.015 |
| ba               | 0.011            | 0.014        | 0.010        | 0.010 | 0.031        | <u>0.045</u> | <u>0.03</u> | 0.045 |
| bb               | <u>0.051</u>     | <u>0.062</u> | <u>0.054</u> | 0.054 | <u>0.015</u> | 0.023        | 0.015       | 0.023 |

Pouze na podtržených místech se liší od správné frekvenci tabulky



%

$$P(\bar{T}_5 = a | \bar{T}_3 = a, \bar{T}_4 = a) = \frac{P(\bar{T}_5 = a, \bar{T}_4 = a, \bar{T}_3 = a)}{P(\bar{T}_4 = a, \bar{T}_3 = a)} =$$

$$= \frac{\text{sumet sloypce aaa frekvenci' tabulky}}{\text{sumet sloypce aaa s aab}} = \frac{0.111}{0.111 + 0.135} =$$

$$= \frac{0.111}{0.245} = 0,453 \doteq 0,45$$

$$P(\bar{T}_5 | \bar{T}_4, \bar{T}_3):$$

| $\bar{T}_4$ | $\bar{T}_3$  |            |
|-------------|--------------|------------|
|             | a            | b          |
| a           | (0.45, 0.55) | (0.4, 0.6) |
| b           | (0.5, 0.5)   | (0.4, 0.6) |

$$(T_5 = a, T_5 = b)$$

Vypočteme  $\text{dist}_Q(P, P^*) = \sum_i (x_i - y_i)^2 =$

$$= 0.001^2 + 0.001^2 + 0.001^2 + 0.001^2 + 0.001^2 + \\ + 0.001^2 + 0.002^2 + 0.001^2 + 0.001^2 =$$

$$= 0.000008 + 0.000004 = 0.000012$$

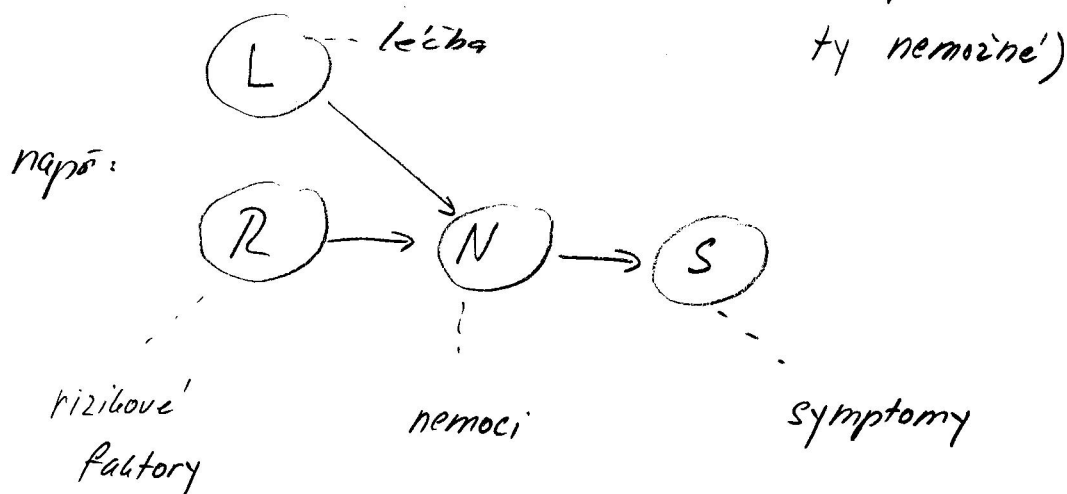
$$\text{Acc}(M_{\text{Min}}) = \text{Size}(M_{\text{Min}}) + 10000 \times \text{dist}_Q(P, P^*) =$$

$$= 2 + 2 + 4 + 4 + 8 = 20 + 0.12 = 20,12$$

Takže vybereme model  $M_{\text{Min}}$ .

Další možnosti:

určit nejistou hierarchii pro klustery proměnných, která  
na'm pomůže určit možné hrany (nebo spíš vyloučit



rozhodně od proměnných z kategorie symptomů nepovedou  
hrany do rizikových faktorů atd.

No a tohle je vlastně Data-Mining - dát databázi strukturu Bayesovské sítě a tím určit závislosti mezi veličinami.

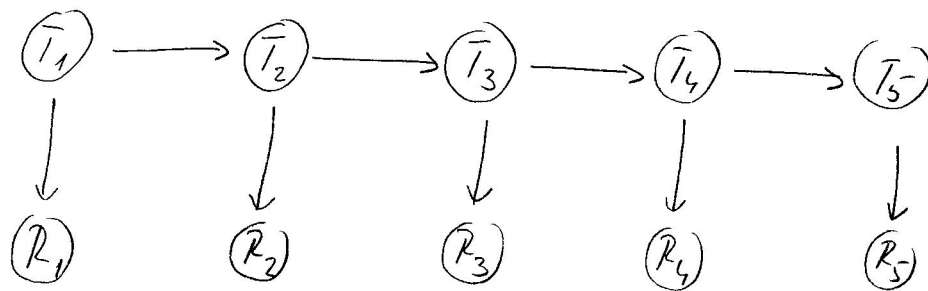
Ma' to řadu praktických problémů - velkou časovou a prostorovou náročnost, chybějící hodnoty (pozorování) v databázích apod.

### Adaptace

Model, který jsme vybrali v procesu učení, nemusí platit jednou provždy. Během práce se sítí dostáváme nová data a chceme, aby se systém učil i z nich.

Adaptace se týká jak podmíněných pravděpodobností, tak struktury sítě.

Adaptace pravděpodobností - příklad: Přenos symbolů  
máme síť:



$T_i$  jsou vysílané znaky, nabývají hodnot  $a, b$

$R_i$  jsou přijaté znaky, — " —  $a, b, c$

( $c$  slouží jako oddělovací slovo - v našem případě vysílané slovo je pětipísmenné, proto se  $c$  mezi  $T_i$  nevyskytuje, ale v důsledku zhrzení přenosovým kanálem mezi přijatými znaky být může)

Podmíněné pravděpodobnosti  $P(R_i | T_i)$  byly odhadnuty na základě zkušenosti se statistickými vlastnostmi zkreslení při přenosu:

|         | $T = a$ | $T = b$ |
|---------|---------|---------|
| $R = a$ | 0.8     | 0.15    |
| $R = b$ | 0.1     | 0.8     |
| $R = c$ | 0.1     | 0.05    |

Pravděpodobnosti  $P(T_{i+1} | T_i)$  se získají z tabulky frekvencí analýzy - to už jsme měli.

Ted si představme, že při každém desátém slově můžeme nějakým způsobem zkontrolovat, co bylo skutečně odesláno.

Tabulku  $P(R | T)$  můžeme interpretovat tak, že při nějakém počátečním výběru velikosti 100 byly pozorovány četnosti:

|         | $T = a$ | $T = b$ |
|---------|---------|---------|
| $R = a$ | 80      | 15      |
| $R = b$ | 10      | 80      |
| $R = c$ | 10      | 5       |

Když nyní získáme další pozorování, dejme tomu jev  $(R=a | T=a)$ , pak se četnost tohoto jevu zvýší na 81, ostatní četnosti zůstanou stejné a velikost výběru bude 101.

Tedy můžeme modifikovat  $P(R=a | T=a) = \frac{81}{101} = 0.8019$

$$P(R=b | T=a) = \frac{10}{101} = 0.099$$

$$P(R=c | T=a) = \frac{10}{101} = 0.099$$

(Klíč se tomu "fractional updating")

Tento postup se může po každém dalším pozorování opakovat - má to ale jednu vadu. Velikost souboru se nám bude rychle

zvětšovat, takže přepočtené pravděpodobnosti budou rychle konvergovat k nějakým číslům, která se dál už téměř nebudou měnit (budou to limity), a tudíž téměř nebudou reagovat na další přidání pozorování.

Aby se tomu zabránilo, udržuje se tzv. efektivní velikost souboru, ze kterého se odhady počítají, označíme ji  $s^*$ .

Měla by mít tu vlastnost, že podíl  $\frac{s^*-1}{s^*} = q$  bude stále konstantní. ( $q$  se nazývá fading factor - faktor slábnutí?)

Zvolíme-li třeba  $s^* = 100$ , pak  $q = 0.99$

Tímto faktorem se budou naobit všechny četnosti i velikost souboru po každém novém pozorování, tj.

$$\text{velikost souboru } S := Sq + 1 = 100 \times 0.99 + 1 = 100$$

(efektivní velikost souboru se nemění)

$$\text{četnost } (R=a | T=a) = 80 \times q + 1 = 80 \times 0.99 + 1 = 80.2$$

$$(R=b | T=a) = 10 \times q = 9.9$$

$$(R=c | T=a) = 10 \times q = 9.9$$

pravděpodobnosti pak dostaneme vydělením 100

Předpokládejme, že jsme získali jedno zkontrolované slovo a víme, že odesláno bylo baaba a přijato bylo baaca.

Můžeme tedy 3x modifikovat  $P(R|a)$  a 2x  $P(R|b)$ .

$$\text{četnost } (R=a | T=a) = ((80 \times 0.99 + 1) \times 0.99 + 1) \times 0.99 + 1 = 80.6$$

$$(R=b | T=a) = ((10 \times 0.99) \times 0.99) \times 0.99 = 9.7$$

$$(R=c | T=a) = 10 \times 0.99^3 = 9.7$$

$$(R=a | T=b) = (15 \times 0.99) \times 0.99 = 14.7$$

$$(R=b | T=b) = (80 \times 0.99 + 1) \times 0.99 = 79.4$$

$$(R=c | T=b) = (5 \times 0.99) \times 0.99 + 1 = 5.9$$

Pravděpodobnosti:

|       | $T=a$ | $T=b$ |
|-------|-------|-------|
| $R=a$ | 0.806 | 0.147 |
| $R=b$ | 0.094 | 0.794 |
| $R=c$ | 0.094 | 0.059 |

Takže třeba  $P(a|a)$  se mírně zvětšila, protože všechny 3 znaky  $a$  se přenesly správně,

$P(b|b)$  se zmenšila, protože ze 2 znaků  $b$  se jeden přenesl špatně.

Pravděpodobnosti  $P(T_{i+1} | T_i)$  se také dají adaptovat a mohou se k tomu použít všechna přijata' slova, nejen ta kontrolní.

Jiná' možnost, jak adaptovat pravděpodobnosti, je např. formálně přidat k síti jednu (nebo několik málo) rodičovskou proměnnou, jejíž rozdělení bude reagovat na nová data, a bude se přizpůsobovat ke všem proměnným, které jsou na ni závislé.

## Adaptace struktury

Bohužel neexistuje žádná síťová metoda pro inkrementální adaptaci struktury.

Důvod: změny struktury jsou složitější a nedají se zakládat na jednotlivém příjmu (třeba na jednom přijatém písmenu), musí být založeny na akumulovaných datech

Jsou 2 způsoby:

akumulovat nová data a čas od času spustit nové učení na celém souboru

ponechat starou síť a vytvořit (naučit) novou na nových datech

Mezi starou a novou sítí se pak buď přepíná, nebo se použijí nějaké váhy.

Mějme modely  $M_1$  a  $M_2$  a váhy  $w_1, w_2$ . Když přijde nějaké nové pozorování (evidence), budeme mít v modelu  $M_1$  pravděpodobnosti  $P_1(A|e)$ , v modelu  $M_2$   $P_2(A|e)$ , a zkombinujeme je na  $P(A|e) = w_1 P_1(A|e) + w_2 P_2(A|e)$

Interpretace:  $w_i$  je pravděpodobnost modelu  $M_i$

Váhy se <sup>vždy</sup> po novém pozorování přepočítají:

$$w_i := P(M_i|e) = \frac{P(e|M_i)P(M_i)}{P(e)} = \frac{w_i P_i(e)}{\sum_j w_j P_j(e)}$$

Problém je, že musíme mít uloženo více modelů, což je paměťově náročné, a vytváření každého nového je zase časově náročné

Jak se da' ušetřit čas:

Předpokláda' se, že nova' síť bude asi dost podobna' té' staré,  
může se lišit třeba jen v několika málo hranách.

Vytipujeme si tyto hrany a označíme je za sporné, tj. v  
síti někdy byt' mohou, někdy ne.

Vytváření nové struktury pak spočíva' vlastně jen v přidávání  
nebo ubírání těchto sporných hran.

A ještě naivní přístup -

doufáme, že pečlivá adaptace pravděpodobnosti vykompenzuje  
i mírně nekorektní strukturu sítě

k adaptaci struktury se přistupuje až když je nevyhnutelná

### Ladění

To je složitější, takže ho zatím vynecháme.

V podstatě jde o to vyladit síť tak, aby pravděpodobnosti

$P(A|e)$  vyhovovaly nějakým předem daným požadavkům.