

Statistical methods in NLP

Unsupervised and semisupervised methods



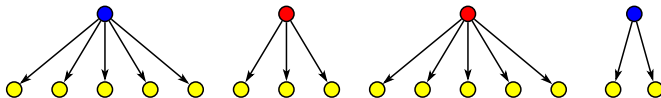
**UNIVERSITY OF
GOTHENBURG**

Richard Johansson

March 3, 2016

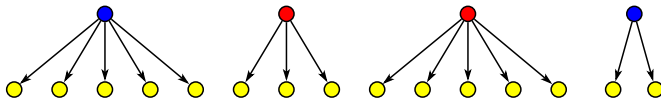
what kind of information is available?

- ▶ **supervised** learning: the labels are given

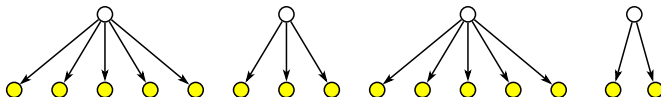


what kind of information is available?

- ▶ **supervised** learning: the labels are given

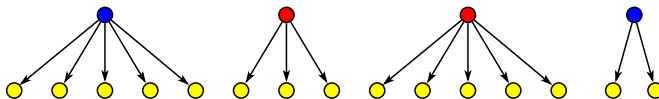


- ▶ **unsupervised** learning (**clustering**): the labels are not given

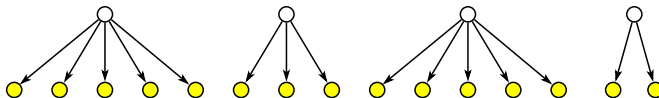


what kind of information is available?

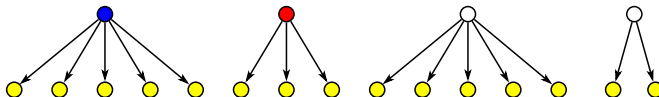
- ▶ **supervised** learning: the labels are given



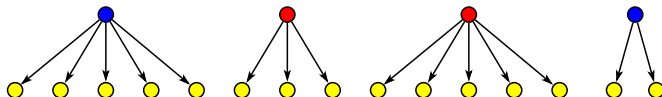
- ▶ **unsupervised** learning (**clustering**): the labels are not given



- ▶ **semisupervised** learning: some of the labels are given



estimation in Naive Bayes, revisited



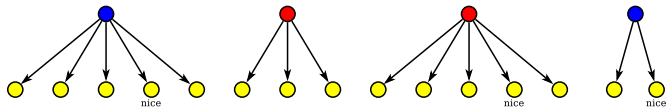
- ▶ Naive Bayes:

$$P(\text{document}, \text{label}) =$$

$$\begin{aligned} P(f_1, \dots, f_n, \text{label}) &= P(\text{label}) \cdot P(f_1, \dots, f_n | \text{label}) \\ &= P(\text{label}) \cdot P(f_1 | \text{label}) \cdot \dots \cdot P(f_n | \text{label}) \end{aligned}$$

- ▶ how do we estimate the probabilities?
 - ▶ **maximum likelihood**: set the probabilities so that the probability of the data is maximized

estimation in Naive Bayes: supervised case



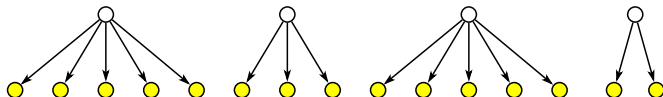
- how do we estimate $P(\text{positive})$?

$$P_{\text{MLE}}(\text{positive}) = \frac{\text{count}(\text{positive})}{\text{count}(\text{all})} = \frac{2}{4}$$

- how do we estimate $P(\text{"nice"}|\text{positive})$?

$$P_{\text{MLE}}(\text{"nice"}|\text{positive}) = \frac{\text{count}(\text{"nice", positive})}{\text{count}(\text{any word, positive})} = \frac{2}{7}$$

what if we are missing labeled data?



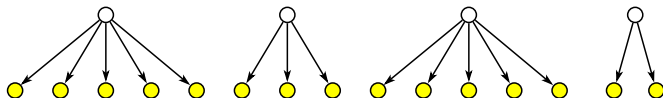
- ▶ we stay in the maximum likelihood framework
- ▶ in the supervised case, we maximize

$$P(\text{doc}_1, \text{lbl}_1) \cdots P(\text{doc}_n, \text{lbl}_n)$$

- ▶ in the unsupervised case, we do the same, only that we don't observe the document labels, so we instead maximize

$$P(\text{doc}_1) \cdots P(\text{doc}_n)$$

maximizing the likelihood in the unsupervised case



- ▶ what is the probability of a document?

- ▶ sum over all possible labels
- ▶ e.g. with sentiment labels:

$$P(\text{doc}) = P(\text{doc}, \text{POS}) + P(\text{doc}, \text{NEG})$$

- ▶ ...so the likelihood is:

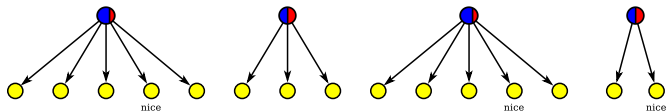
$$(P(\text{doc}_1, \text{POS}) + P(\text{doc}_1, \text{NEG})) \cdots (P(\text{doc}_n, \text{POS}) + P(\text{doc}_n, \text{NEG}))$$

- ▶ this formula is more complex and if we want to maximize it, there is no nice and clean solution as in the supervised case

the Expectation–Maximization algorithm

- ▶ **Expectation–Maximization** is an algorithm that tries to maximize likelihood when there are unobserved variables
- ▶ it is a “circular” two-step algorithm:
 - ▶ **Expectation**: using our current estimates, compute label probabilities for each documents and use them as “soft counts”
 - ▶ for instance, if a document has a probability of 40% of being positive, then we count it as 40% of a positive document
 - ▶ **Maximization**: using the soft counts, compute probability estimates with the normal procedure
- ▶ a chicken-and-egg problem
 - ▶ ...so we need to initialize either the soft counts or the probabilities, and then start at E or M

soft counts example, Naive Bayes

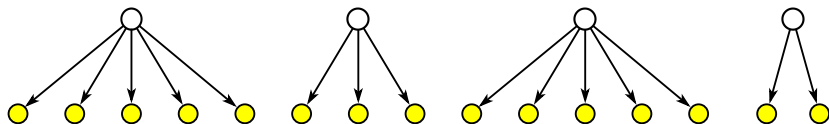


- ▶ let's assume that our classifier assigned the probabilities 0.8, 0.5, 0.7, and 0.6, respectively, for the documents belonging to the positive class

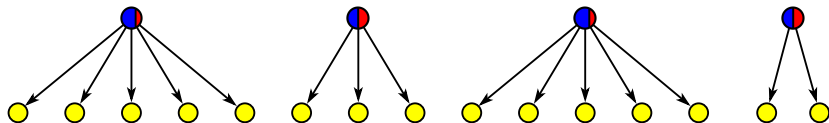
- ▶ then:
$$P_{\text{MLE}}(\text{positive}) = \frac{\text{count}(\text{positive})}{\text{count}(\text{all})} = \frac{0.8 + 0.5 + 0.7 + 0.6}{4}$$

$$\begin{aligned} P_{\text{MLE}}(\text{"nice"}|\text{positive}) &= \frac{\text{count}(\text{"nice", positive})}{\text{count}(\text{any word, positive})} \\ &= \frac{0.8 + 0.7 + 0.6}{0.8 \cdot 5 + \dots + 0.6 \cdot 2} \end{aligned}$$

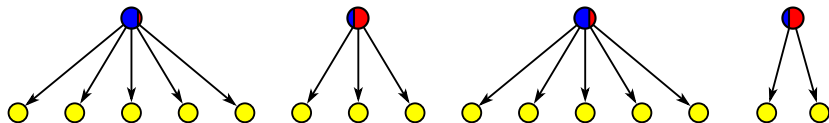
EM iteration for Naive Bayes, example



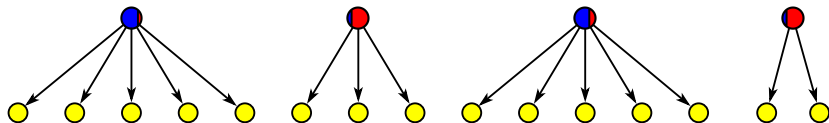
EM iteration for Naive Bayes, example



EM iteration for Naive Bayes, example

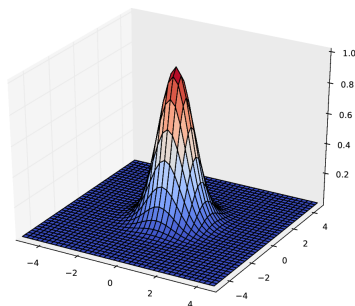


EM iteration for Naive Bayes, example



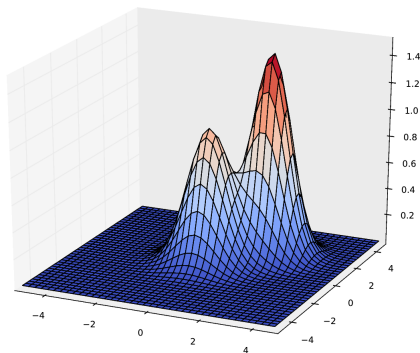
nice property of EM

- ▶ **theorem:** each time we carry out the E and M steps, the likelihood won't decrease
- ▶ EM will converge (stop) at some point
 - ▶ we can climb “uphill” until we reach the top



not-so-nice property of EM

- ▶ ...but the likelihood may end up in a **local maximum** rather than the true maximum
 - ▶ there might be tops that are higher



- ▶ initialization will determine where we end up!

unsupervised experiment

- ▶ DVD reviews vs. music reviews
- ▶ accuracy (assuming category A is DVDs): 0.938

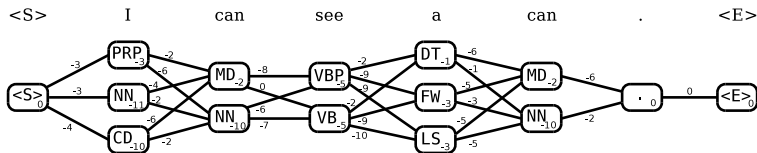
category A:

```
episodes 3.639010348
novel 3.61869012429
sequel 3.57220166914
seagal 3.33618534235
suspense 3.23906073358
vampire 3.19579018682
buffy 3.17306201864
plot 3.16664944584
thriller 3.14954723164
premise 3.12519741486
```

category B:

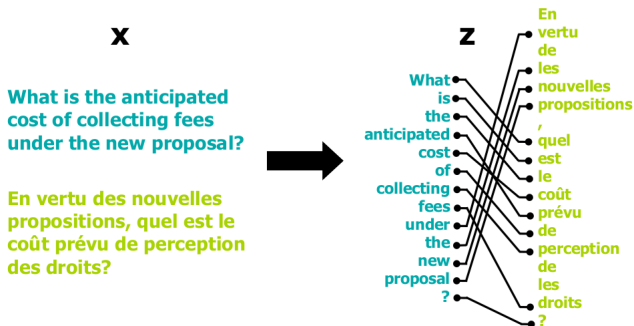
```
vocals -4.15311759228
albums -4.12576126936
album -3.98345604473
catchy -3.98199815327
acoustic -3.86295373766
punk -3.80731612078
guitars -3.79254346553
lyrics -3.65137587064
remix -3.64054296621
lp -3.6211479113
```


tagging example



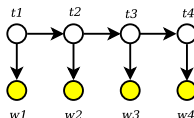
next lecture and VG assignment 2

- ▶ word alignment for machine translation
- ▶ the observed part: sentence pairs
- ▶ the unseen part: word-to-word alignments



Brown clustering

- ▶ the **Brown** algorithm is a clustering method for **words**:
 - ▶ start by putting each word into a cluster
 - ▶ merge clusters to increase HMM language model probability of our corpus



- ▶ there is a popular implementation by Percy Liang:
 - ▶ <http://cs.stanford.edu/~pliang/software/> – under “Word clustering”

Brown et al. *Class-Based n -gram Models of Natural Language*. Computational Linguistics, 1992.

example: clusters from book reviews

- ▶ I created a corpus of 865,000 Amazon book reviews and ran Liang's software to create 2,000 word clusters
- ▶ ...after a few days, it finished
- ▶ here are a couple of examples of clusters

cluster 1000010010111:

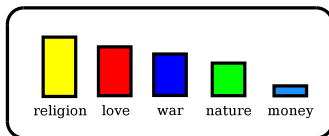
time-consuming	225
repugnant	234
unnerving	240
objectionable	243
anticlimactic	244
reprehensible	258
anti-climactic	270
deceiving	289
disrespectful	299
dissappointing	308

```
cluster 1011111100011:
```

maine	1758
turkey	1796
manhattan	1860
boston	3704
florida	3764
chicago	4535
london	8383
paris	6329
heaven	5864
california	7094

generative story in LDA

- Tolstoy's preferred topics:



religion ■	
God	0.06
faith	0.04
believe	0.03
sacrifice	0.02
pray	0.02
priest	0.02
inner	0.01
devotion	0.01
...	

love ■	
love	0.05
marry	0.04
wedding	0.04
heart	0.03
miss	0.03
dear	0.02
propose	0.02
kiss	0.01
...	

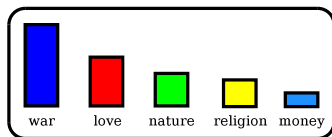
war ■	
war	0.05
soldier	0.04
battle	0.03
gun	0.03
brave	0.03
wound	0.03
resist	0.02
kill	0.02
...	

nature ■	
forest	0.07
flower	0.05
grass	0.04
open	0.04
lake	0.03
tundra	0.03
tree	0.02
sprin	0.01
...	

money ■	
money	0.05
ruble	0.04
gold	0.04
pay	0.03
debt	0.03
earn	0.02
coin	0.02
contract	0.01
...	

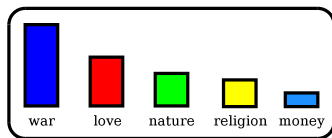
generative story in LDA

- ... Tolstoy first selected a topic composition randomly



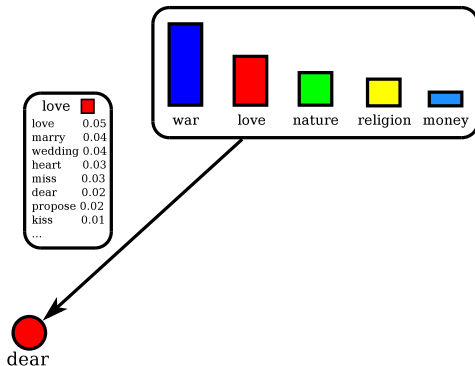
generative story in LDA

- ▶ he then selected a topic for the first word



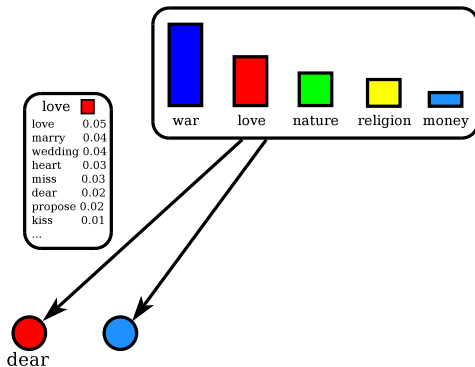
generative story in LDA

- ...and then decided which word to write



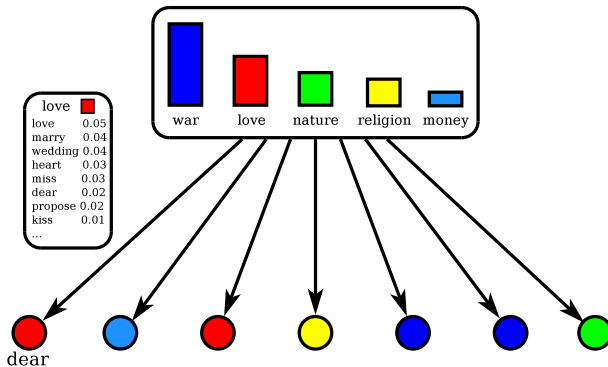
generative story in LDA

- ▶ he then selected the topic for the second word



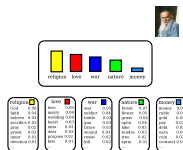
generative story in LDA

- ...and so on

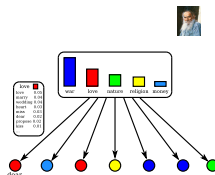


estimation in topic models

- ▶ given a corpus of documents, the LDA estimation process tries to “reverse-engineer” the generative model
- ▶ reconstruct the topics and their overall distributions



- ▶ reconstruct the composition of topics in each document



example: topics for some Swedish Wikipedia articles

- ▶ **Gothenburg**: “city buildings” 0.70, “Swedish bureaucracy” 0.15, “culture” 0.08, “geography” 0.05
- ▶ **jazz**: “music” 0.50, “language and theory” 0.24, “records” 0.18, “Anglo-Saxons” 0.05
- ▶ **natural language processing**: “language and theory” 0.60, “computers” 0.24
- ▶ **Python**: “programming” 0.55, “computers” 0.42
- ▶ **Silvio Berlusconi**: “politics” 0.46, “Catholicism and southern Europe” 0.10, “film and TV” 0.08, “commerce” 0.07
- ▶ **Zlatan Ibrahimovic**: “sport” 0.88, “family” 0.05, “commerce” 0.02

