

Natural Language Processing

Text classification

Yulia Tsvetkov

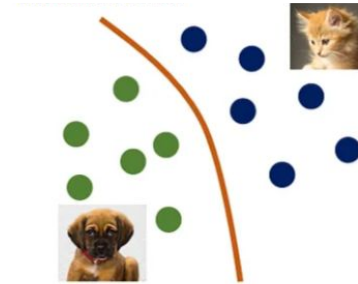
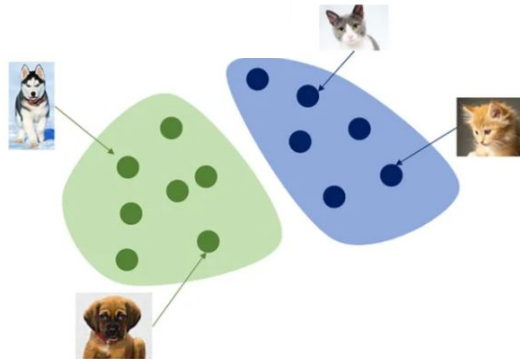
yuliats@cs.washington.edu

Announcements

- Quiz 1 – good luck!

Generative vs discriminative models

Learns the input distribution
Maximizes the joint probability: $P(X, Y)$
Estimates $P(X Y)$ to find $P(Y X)$ using Bayes' rule
Can generate new data
Typically, they are NOT used to solve classification tasks
Generative models possess discriminative properties



Learns the decision boundary between classes
Maximizes the conditional probability: $P(Y X)$
Directly estimates $P(Y X)$
Cannot generate new data
Specifically meant for classification tasks
Discriminative models don't possess generative properties

Logistic Regression	Random Forests	SVMs
Neural Networks	Decision Tree	kNN

Hidden Markov Models	Naive Bayes	Gaussian Mixture Models
Gaussian Discriminant Analysis	LDA	Bayesian Networks

<https://blog.dailydoseofds.com/p/an-intuitive-guide-to-generative>
<https://medium.com/@jordi299/about-generative-and-discriminative-models-d8958b67ad32>

Generative and discriminative models

- Generative text classification: Learn a model of the joint $P(X, y)$, and find

$$\hat{y} = \operatorname{argmax}_{\tilde{y}} P(X, \tilde{y})$$

- Discriminative text classification: Learn a model of the conditional $P(y | X)$, and find

$$\hat{y} = \operatorname{argmax}_{\tilde{y}} P(\tilde{y} | X)$$

Finding the correct class c from a document d in Generative vs Discriminative Classifiers

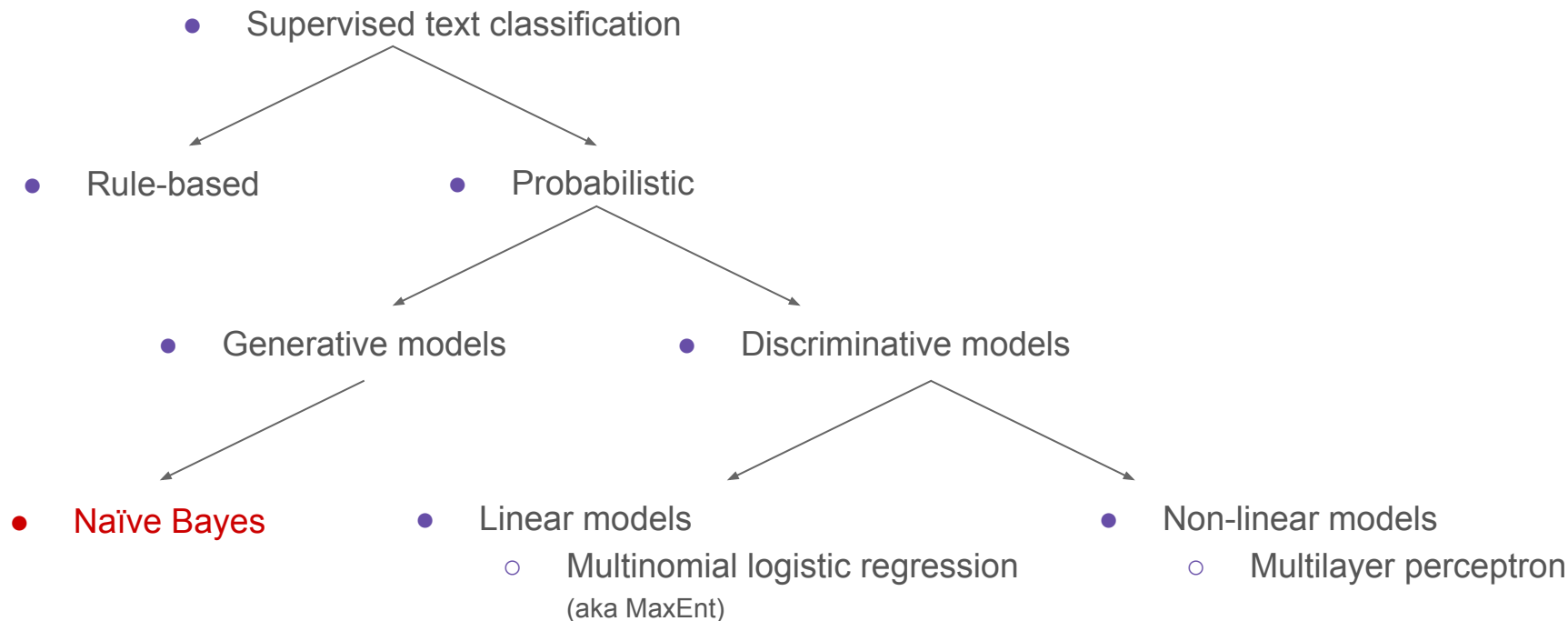
- Naive Bayes

$$\hat{c} = \operatorname{argmax}_{c \in C} \overbrace{P(d|c)}^{\text{likelihood}} \overbrace{P(c)}^{\text{prior}}$$

- Logistic Regression

$$\hat{c} = \operatorname{argmax}_{c \in C} \overbrace{P(c|d)}^{\text{posterior}}$$

We'll consider alternative models for classification



Generative text classification: naïve Bayes

- Simple classification method
 - based on the Bayes rule
- Relies on very simple (naïve) representation of a documents
 - Conditional independence assumption:
the features are conditionally independent, given the target class
(hence the name “naïve”)
 - bag-of-words, no relative order
- A good baseline for more sophisticated models

Andrew Y. Ng and Michael I. Jordan, On discriminative vs. generative classifiers: A comparison of logistic regression and naïve Bayes, Advances in Neural Information Processing Systems 14 (NIPS), 2001.

Naïve Bayes

Sentiment analysis: movie reviews

- Given a document d (e.g., a movie review)
- Decide which class c it belongs to: positive, negative, neutral
- Compute $P(c | d)$ for each c
 - $P(\text{positive} | d)$, $P(\text{negative} | d)$, $P(\text{neutral} | d)$
 - select the one with max P

Bag-of-Words (BOW)

- Given a document d (e.g., a movie review) – how to represent d ?

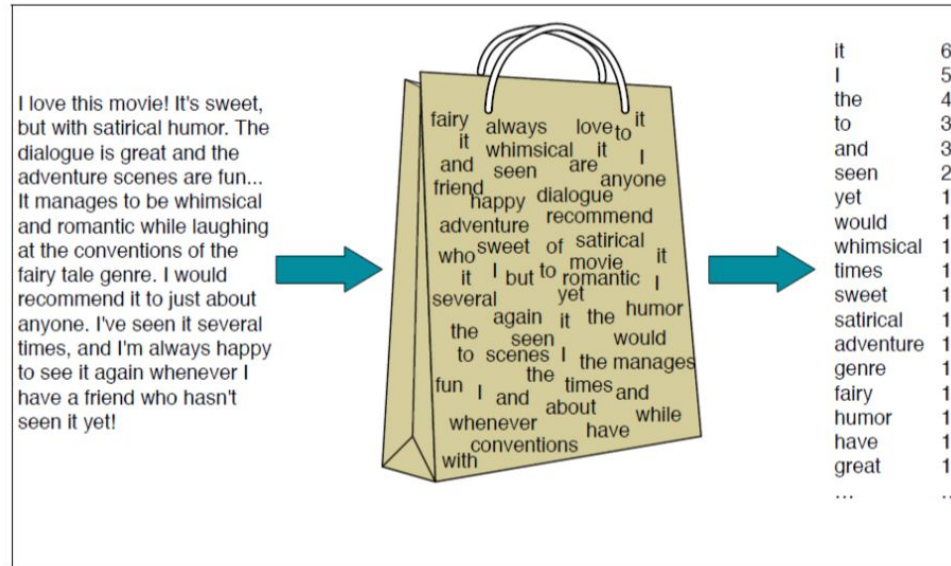


Figure 7.1 Intuition of the multinomial naive Bayes classifier applied to a movie review. The position of the words is ignored (the *bag of words* assumption) and we make use of the frequency of each word.

Figure from J&M 3rd ed. draft, sec 7.1

Naïve Bayes

- Given a document **d** and a class **c**, use Bayes' rule:

$$P(c|d) = \frac{P(d|c)P(c)}{P(d)}$$



posterior

Naïve Bayes

- Given a document **d** and a class **c**, Bayes' rule:

$$P(c|d) = \frac{P(d|c)P(c)}{P(d)}$$

$$P(\text{'positive'}|d) \propto P(d|\text{'positive'})P(\text{'positive'})$$

↓
posterior

↓
likelihood

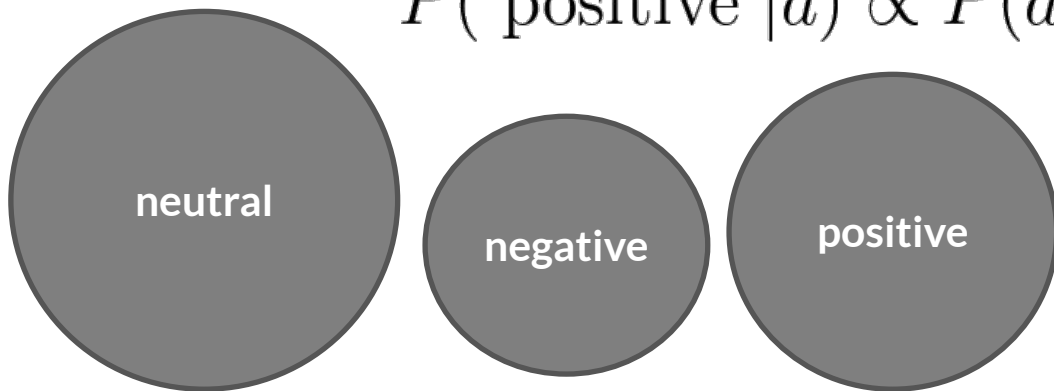
↓
prior

Naïve Bayes

- Given a document d and a class c , Bayes' rule:

$$P(c|d) = \frac{P(d|c)P(c)}{P(d)}$$

$$P(\text{'positive'}|d) \propto P(d|\text{'positive'})P(\text{'positive'})$$



↓
prior

Naïve Bayes

- Given a document **d** and a class **c**, Bayes' rule:

$$P(c|d) = \frac{P(d|c)P(c)}{P(d)}$$

$$P(\text{'positive'}|d) \propto P(d|\text{'positive'})P(\text{'positive'})$$



likelihood

Naïve Bayes independence assumptions

$$P(w_1, w_2, \dots, w_n | c)$$

- **Bag of Words assumption:** Assume position doesn't matter
- **Conditional Independence:** Assume the feature probabilities $P(w_i | c_j)$ are independent given the class c

$$P(w_1, w_2, \dots, w_n | c) = P(w_1 | c) \times P(w_2 | c) \times P(w_3 | c) \times \dots \times P(w_n | c)$$

Document representation

I love this movie. It's sweet but with satirical humor. The dialogue is great and the adventure scenes are fun... it manages to be whimsical and romantic while laughing at the conventions of the fairy tale genre. I would recommend it to just about anyone. I've seen it several times, and I'm always happy to see it again whenever I have a friend who hasn't seen it yet!



**bag of words
(BOW)**



it	6
I	5
the	4
to	3
and	3
seen	2
yet	1
would	1
whimsical	1
times	1
sweet	1
satirical	1
adventure	1
genre	1
fairy	1
humor	1
have	1
great	1
...	...

Document representation

I love this movie. It's sweet but with satirical humor. The dialogue is great and the adventure scenes are fun... it manages to be whimsical and romantic while laughing at the conventions of the fairy tale genre. I would recommend it to just about anyone. I've seen it several times, and I'm always happy to see it again whenever I have a friend who hasn't seen it yet!



**bag of words
(BOW)**



it	6
I	5
the	4
to	3
and	3
seen	2
yet	1
would	1
whimsical	1
times	1
sweet	1
satirical	1
adventure	1
genre	1
fairy	1
humor	1
have	1
great	1
...	...

$$P(d|c) = P(w_1, w_2, \dots, w_n|c) = \prod_i P(w_i|c)$$

Generative text classification: Naïve Bayes

$$C_{NB} = \operatorname{argmax}_c P(c|d) = \operatorname{argmax}_c \frac{P(d|c)P(c)}{P(d)} \propto \text{Bayes rule}$$

$$\operatorname{argmax}_c P(d|c)P(c) = \text{same denominator}$$

$$\operatorname{argmax}_c P(w_1, w_2, \dots, w_n|c)P(c) = \text{representation}$$

$$\operatorname{argmax}_{c_j} P(c_j) \prod_i P(w_i|c) \quad \text{conditional independence}$$

Underflow prevention: log space

- Multiplying lots of probabilities can result in floating-point underflow
- Since $\log(xy) = \log(x) + \log(y)$
 - better to sum logs of probabilities instead of multiplying probabilities
- Class with highest un-normalized log probability score is still most probable

$$C_{NB} = \operatorname{argmax}_{c_j} P(c_j) \prod_i P(w_i|c)$$

$$C_{NB} = \operatorname{argmax}_{c_j} \log(P(c_j)) + \sum_i \log(P(w_i|c))$$

- Model is now just max of sum of weights

Learning the multinomial naïve Bayes

- How do we learn (train) the NB model?

Learning the multinomial naïve Bayes

- How do we learn (train) the NB model?
- We learn $P(c)$ and $P(w_i|c)$ from training (labeled) data

$$C_{NB} = \operatorname{argmax}_{c_j} \log(\underline{P(c_j)}) + \sum_i \log(\underline{P(w_i|c)})$$

Parameter estimation for NB

- Parameter estimation during training
- Concatenate all documents with category c into one mega-document
- Use the frequency of w_i in the mega-document to estimate the word probability

$$C_{NB} = \operatorname{argmax}_{c_j} \log(P(c_j)) + \sum_i \log(P(w_i|c))$$

$$\hat{P}(c_j) = \frac{\text{doccount}(C = c_j)}{N_{doc}}$$

$$\hat{P}(w_i|c_j) = \frac{\text{count}(w_i, c_j)}{\sum_{w \in V} \text{count}(w, c_j)}$$

Parameter estimation for NB

$$\hat{P}(w_i|c_j) = \frac{\text{count}(w_i, c_j)}{\sum_{w \in V} \text{count}(w, c_j)}$$

- fraction of times word w_i appears among all words in documents of topic c_j
- Create mega-document for topic j by concatenating all docs in this topic
 - Use frequency of w in mega-document

Problem with Maximum Likelihood

- What if we have seen no training documents with the word “fantastic” and classified in the topic **positive**?

Problem with Maximum Likelihood

- What if we have seen no training documents with the word “fantastic” and classified in the topic **positive**?

$$\hat{P}(\text{“fantastic”} | c = \text{positive}) = \frac{\text{count}(\text{“fantastic”}, \text{positive})}{\sum_{w \in V} \text{count}(w, \text{positive})} = 0$$

- Zero probabilities cannot be conditioned away, no matter the other evidence!

$$\operatorname{argmax}_{c_j} P(c_j) \prod_i P(w_i | c)$$

Laplace (add-1) smoothing for naïve Bayes

$$\hat{P}(w_i|c_j) = \frac{\text{count}(w_i, c_j) + 1}{\sum_{w \in V} (\text{count}(w, c_j) + 1)}$$

Laplace (add-1) smoothing for naïve Bayes

$$\begin{aligned}\hat{P}(w_i|c_j) &= \frac{\text{count}(w_i, c_j) + 1}{\sum_{w \in V} (\text{count}(w, c_j) + 1)} \\ &= \frac{\text{count}(w_i, c_j) + 1}{(\sum_{w \in V} \text{count}(w, c_j)) + |V|}\end{aligned}$$

- Note about log space

Multinomial naïve Bayes : learning

- From training corpus, extract *Vocabulary*
- Calculate $P(c_j)$ terms
 - For each c_j do
 - $docs_j \leftarrow$ all docs with class = c_j
 - $P(c_j) \leftarrow \frac{|docs_j|}{total \# documents}$

Multinomial naïve Bayes : learning

- From training corpus, extract *Vocabulary*
- Calculate $P(c_j)$ terms
 - For each c_j do
 - $docs_j \leftarrow$ all docs with class = c_j
 - $P(c_j) \leftarrow \frac{|docs_j|}{total \# documents}$
- Calculate $P(w_i | c_j)$ terms
 - $Text_j \leftarrow$ single doc containing all docs_j
 - For each word w_i in *Vocabulary*
 - $n_i \leftarrow$ # of occurrences of w_i in $Text_j$
 - $P(w_j | c_j) \leftarrow \frac{n_i + \alpha}{n + \alpha |Vocabulary|}$

Example

	Doc	Words	Class
Training	1	Chinese Beijing Chinese	c
	2	Chinese Chinese Shanghai	c
	3	Chinese Macao	c
	4	Tokyo Japan Chinese	j
Test	5	Chinese Chinese Chinese Tokyo Japan	?

Example

$$\hat{P}(c) = \frac{N_c}{N}$$

Priors:

$$P(c) = \frac{3}{4}$$

$$P(j) = \frac{1}{4}$$

	Doc	Words	Class
Training	1	Chinese Beijing Chinese	c
	2	Chinese Chinese Shanghai	c
	3	Chinese Macao	c
	4	Tokyo Japan Chinese	j
Test	5	Chinese Chinese Chinese Tokyo Japan	?

Example

$$\hat{P}(c) = \frac{N_c}{N}$$

$$\hat{P}(w|c) = \frac{\text{count}(w,c) + 1}{\text{count}(c) + |V|}$$

Priors:

$$P(c) = \frac{3}{4}$$

$$P(j) = \frac{1}{4}$$

Conditional Probabilities:

$$P(\text{Chinese}|c) = (5+1) / (8+6) = 6/14 = 3/7$$

$$P(\text{Tokyo}|c) = (0+1) / (8+6) = 1/14$$

$$P(\text{Japan}|c) = (0+1) / (8+6) = 1/14$$

$$P(\text{Chinese}|j) = (1+1) / (3+6) = 2/9$$

$$P(\text{Tokyo}|j) = (1+1) / (3+6) = 2/9$$

$$P(\text{Japan}|j) = (1+1) / (3+6) = 2/9$$

	Doc	Words	Class
Training	1	Chinese Beijing Chinese	c
	2	Chinese Chinese Shanghai	c
	3	Chinese Macao	c
	4	Tokyo Japan Chinese	j
Test	5	Chinese Chinese Chinese Tokyo Japan	?

Example

$$\hat{P}(c) = \frac{N_c}{N} \quad \hat{P}(w|c) = \frac{\text{count}(w,c)+1}{\text{count}(c)+|V|}$$

Priors:

$$P(c) = \frac{3}{4}$$

$$P(j) = \frac{1}{4}$$

Conditional Probabilities:

$$\begin{aligned} P(\text{Chinese}|c) &= (5+1) / (8+6) = 6/14 = 3/7 \\ P(\text{Tokyo}|c) &= (0+1) / (8+6) = 1/14 \\ P(\text{Japan}|c) &= (0+1) / (8+6) = 1/14 \\ P(\text{Chinese}|j) &= (1+1) / (3+6) = 2/9 \\ P(\text{Tokyo}|j) &= (1+1) / (3+6) = 2/9 \\ P(\text{Japan}|j) &= (1+1) / (3+6) = 2/9 \end{aligned}$$

	Doc	Words	Class
Training	1	Chinese Beijing Chinese	c
	2	Chinese Chinese Shanghai	c
	3	Chinese Macao	c
	4	Tokyo Japan Chinese	j
Test	5	Chinese Chinese Chinese Tokyo Japan	?

Choosing a class:

$$P(c|d5) \propto \frac{3}{4} * \left(\frac{3}{7}\right)^3 * \frac{1}{14} * \frac{1}{14} \approx 0.0003$$

$$P(j|d5) \propto \frac{1}{4} * \left(\frac{2}{9}\right)^3 * \frac{2}{9} * \frac{2}{9} \approx 0.0001$$

Summary: naïve Bayes is not so naïve

- Naïve Bayes is a probabilistic model
- Naïve because it assumes features are independent of each other for a class
- Very fast, low storage requirements
- Robust to Irrelevant Features
 - Irrelevant Features cancel each other without affecting results
- Very good in domains with many equally important features
 - Decision Trees suffer from fragmentation in such cases – especially if little data
- Optimal if the independence assumptions hold: If assumed independence is correct, then it is the Bayes Optimal Classifier for problem
- A good dependable baseline for text classification
 - But we will see other classifiers that give better accuracy