

LOGISTIC REGRESSION

David Kauchak
CS158 – Fall 2025

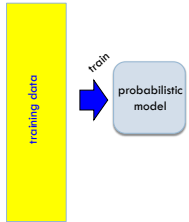
1

Admin

Assignment 7

2

MLE estimation for NB



$$p(y) \prod_{j=1}^m p(x_j | y)$$

$p(y)$ $p(x_i | y)$

What are the MLE estimates for these?

3

Maximum likelihood estimates

$$p(y) = \frac{\text{count}(y)}{n}$$

number of examples with label y
total number of examples

$$p(x_i | y) = \frac{\text{count}(x_i, y)}{\text{count}(y)}$$

number of examples with label y with feature $x_i = 1$
number of examples with label

4

Maximum likelihood estimates

$$p(y) = \frac{\text{count}(y)}{n}$$

x1	x2	label
1	1	1
1	0	1
1	1	1
0	1	-1
0	0	-1

$p(1) = ?$
 $p(-1) = ?$

5

Maximum likelihood estimates

$$p(y) = \frac{\text{count}(y)}{n}$$

x1	x2	label
1	1	1
1	0	1
1	1	1
0	1	-1
0	0	-1

$p(1) = 3/5$
 $p(-1) = 2/5$

6

Maximum likelihood estimates

$$p(x_i | y) = \frac{\text{count}(x_i, y)}{\text{count}(y)}$$

x1	x2	label
1	1	1
1	0	1
1	1	1
0	1	-1
0	0	-1

$p(x_1 = 1 1)$	?
$p(x_1 = 0 1)$	?
$p(x_2 = 1 1)$	?
$p(x_2 = 0 1)$	?

7

Maximum likelihood estimates

$$p(x_i | y) = \frac{\text{count}(x_i, y)}{\text{count}(y)}$$

x1	x2	label
1	1	1
1	0	1
1	1	1
0	1	-1
0	0	-1

$p(x_1 = 1 1)$	3/3
$p(x_1 = 0 1)$	0/3
$p(x_2 = 1 1)$	2/3
$p(x_2 = 0 1)$	1/3

8

Maximum likelihood estimates

$$p(x, y) = p(y) \prod_{j=1}^m p(x_j | y)$$

x1	x2	label
1	1	1
1	0	1
1	1	1
0	1	-1
0	0	-1

$p(x_1 = 1 1)$	3/3
$p(x_1 = 0 1)$	0/3
$p(x_2 = 1 1)$	2/3
$p(x_2 = 0 1)$	1/3

$$p(x_1 = 1, x_2 = 1, y = 1) = ?$$

9

Maximum likelihood estimates

$$p(x, y) = p(y) \prod_{j=1}^m p(x_j | y)$$

x1	x2	label
1	1	1
1	0	1
1	1	1
0	1	-1
0	0	-1

$p(x_1 = 1 1)$	3/3
$p(x_1 = 0 1)$	0/3
$p(x_2 = 1 1)$	2/3
$p(x_2 = 0 1)$	1/3

$$\begin{aligned} p(x_1 = 1, x_2 = 1, y = 1) &= p(1)p(x_1 = 1 | 1)p(x_2 = 1 | 1) \\ &= \frac{3}{5} * 1 * \frac{2}{3} \\ &= \frac{6}{15} \end{aligned}$$

10

Maximum likelihood estimates

$$p(x, y) = p(y) \prod_{j=1}^m p(x_j | y)$$

x1	x2	label
1	1	1
1	0	1
1	1	1
0	1	-1
0	0	-1

$p(x_1 = 1 1)$	3/3
$p(x_1 = 0 1)$	0/3
$p(x_2 = 1 1)$	2/3
$p(x_2 = 0 1)$	1/3

$$p(x_1 = 0, x_2 = 1, y = 1) = ?$$

11

Maximum likelihood estimates

$$p(x, y) = p(y) \prod_{j=1}^m p(x_j | y)$$

x1	x2	label
1	1	1
1	0	1
1	1	1
0	1	-1
0	0	-1

$p(x_1 = 1 1)$	3/3
$p(x_1 = 0 1)$	0/3
$p(x_2 = 1 1)$	2/3
$p(x_2 = 0 1)$	1/3

$$\begin{aligned} p(x_1 = 0, x_2 = 1, y = 1) &= p(1)p(x_1 = 0 | 1)p(x_2 = 1 | 1) \\ &= \frac{3}{5} * 0 * \frac{2}{3} \\ &= 0 \quad \ominus \end{aligned}$$

12

Maximum likelihood estimates

Full model trained!

$$p(1) = 3/5$$

$$p(-1) = 2/5$$

x1	x2	label
1	1	1
1	0	1
1	1	1
0	1	-1
0	0	-1

$$p(x_1 = 1 \mid 1) = 3/3$$

$$p(x_2 = 1 \mid 1) = 2/3$$

$$p(x_1 = 1 \mid -1) = 0/2$$

$$p(x_2 = 1 \mid -1) = 1/2$$

13

Basic steps for probabilistic modeling

Step 1: pick a model

Step 2: figure out how to estimate the probabilities for the model

Step 3 (optional): deal with overfitting

Probabilistic models

Which model do we use, i.e. how do we calculate $p(\text{feature}, \text{label})$?

How do train the model, i.e. how do we **estimate the probabilities** for the model?

How do we deal with overfitting?

14

Likelihood distribution

The likelihood of the data, given the model parameters (θ) is:

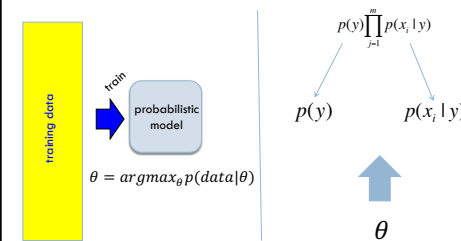
$$p(\text{data}|\theta) \quad \text{likelihood of the data, given the current model}$$

Maximum Likelihood Estimation (MLE), picks the parameter as:

$$\theta = \operatorname{argmax}_{\theta} p(\text{data}|\theta)$$

15

Training: MLE



16

MLE problems

$p(x_1 = 1 1)$	3/3
$p(x_1 = 0 1)$	0/3
$p(x_2 = 1 1)$	2/3
$p(x_2 = 0 1)$	1/3

$$\theta = \operatorname{argmax}_{\theta} p(\text{data}|\theta)$$

The MLE is completely driven by the training data

17

Priors

Coin1 data: 3 Heads and 1 Tail

Coin2 data: 30 Heads and 10 tails

Coin3 data: 2 Tails

Coin4 data: 497 Heads and 503 tails

If someone asked you what the probability of heads was for each of these coins, what would you say?

18

Posterior distribution

We want to incorporate a prior belief of what the probabilities might be

$$p(\theta|\text{data}) \quad \text{given the data, how likely are different model parameters}$$

Maximum a posteriori principle, pick parameters that maximize the posterior distribution:

$$\theta = \operatorname{argmax}_{\theta} p(\theta|\text{data})$$

19

MLE vs. Maximum posteriori

MLE $\theta = \operatorname{argmax}_{\theta} p(\text{data}|\theta)$

Maximum posteriori $\theta = \operatorname{argmax}_{\theta} p(\theta|\text{data})$

Both are methods for picking model parameters

20

Maximum posteriori

$$p(\theta|data) = ? \quad (\text{Hint: Bayes' rule})$$

21

Maximum posteriori

$$p(\theta|data) = \frac{p(data|\theta)p(\theta)}{p(data)}$$

22

Maximum posteriori

What are each of these probabilities?

$$p(\theta|data) = \frac{p(data|\theta)p(\theta)}{p(data)}$$

23

Maximum posteriori

likelihood of the data
under the modelprobability of different parameters,
call the **prior**

$$p(\theta|data) = \frac{p(data|\theta)p(\theta)}{p(data)}$$

probability of seeing the data
(regardless of model)

24

Priors

$$\theta = \operatorname{argmax}_{\theta} \frac{p(\text{data} | \theta) p(\theta)}{p(\text{data})}$$

Does $p(\text{data})$ matter for the argmax?

25

Priors

likelihood of the data
under the model

probability of different parameters,
call the **prior**

$$\theta = \operatorname{argmax}_{\theta} p(\text{data} | \theta) p(\theta)$$

What does MLE assume for a prior on the
model parameters?

26

Priors

likelihood of the data
under the model

probability of different parameters,
call the **prior**

$$\theta = \operatorname{argmax}_{\theta} p(\text{data} | \theta) p(\theta)$$

- Assumes a **uniform prior**, i.e. all θ are equally likely!
- Relies solely on the likelihood

27

A better approach

$$\theta = \operatorname{argmax}_{\theta} p(\text{data} | \theta) p(\theta)$$

$$\text{likelihood}(\text{data}) = \prod_{i=1}^n p_{\theta}(x_i)$$

We can use any distribution we'd like.

This allows us to impart additional **bias**
into the model

28

Another view on the prior

Remember, the max is the same if we take the log:

$$\theta = \operatorname{argmax}_{\theta} \log(p(\text{data} | \theta)) + \log(p(\theta))$$

$$\text{log-likelihood} = \sum_{i=1}^n \log(p(x_i))$$

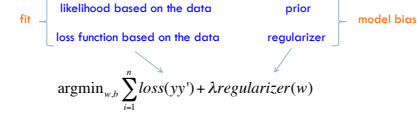
We can use any distribution we'd like.
This allows us to impart addition **bias** into the model

Does this look like something we've seen before
(Hint, think gradient descent)?

29

Regularization vs prior

$$\theta = \operatorname{argmax}_{\theta} \log(p(\text{data} | \theta)) + \log(p(\theta))$$



30

Prior for NB

$$\theta = \operatorname{argmax}_{\theta} \log(p(\text{data} | \theta)) + \log(p(\theta))$$

Uniform prior



$$p(x_i | y) = \frac{\text{count}(x_i, y)}{\text{count}(y)}$$

Dirichlet prior



$$p(x_i | y) = \frac{\text{count}(x_i, y) + \lambda}{\text{count}(y) + \text{possible_values_of_} x_i * \lambda}$$

31

Prior: another view

$$p(x_1, x_2, \dots, x_m, y) = p(y) \prod_{j=1}^m p(x_j | y)$$

$$\text{MLE: } p(x_i | y) = \frac{\text{count}(x_i, y)}{\text{count}(y)}$$

What happens to our likelihood if, for one of the labels, we never saw a particular feature?

Goes to 0!

32

Prior: another view

$$p(x_i | y) = \frac{\text{count}(x_i, y)}{\text{count}(y)}$$



$$p(x_i | y) = \frac{\text{count}(x_i, y) + \lambda}{\text{count}(y) + \text{possible_values_of_}x_i * \lambda}$$

Adding a prior avoids this!

33

Smoothing



$$p(x_i | y) = \frac{\text{count}(x_i, y)}{\text{count}(y)}$$



$$p(x_i | y) = \frac{\text{count}(x_i, y) + \lambda}{\text{count}(y) + \text{possible_values_of_}x_i * \lambda}$$

for each label, pretend like we've seen each feature value occur in λ additional examples

Sometimes this is also called **smoothing** because it is seen as smoothing or interpolating between the MLE and some other distribution

34

Priors

Coin1 data: 3 Heads and 1 Tail
 Coin2 data: 30 Heads and 10 tails
 Coin3 data: 2 Tails
 Coin4 data: 497 Heads and 503 tails

$$p(\text{heads}) = \frac{\text{count}(\text{heads}) + \lambda}{\text{totalflips} + 2\lambda}$$

Does this do the right thing in these cases?

35

Priors

Coin1 data: 3 Heads and 1 Tail:

$$p(\text{heads}) = \frac{4}{6} = 0.667$$

Coin2 data: 30 Heads and 10 tails:

$$p(\text{heads}) = \frac{31}{42} = 0.738$$

Coin3 data: 2 Tails

$$p(\text{heads}) = \frac{3}{4} = 0.75$$

Coin4 data: 497 Heads and 503 tails

$$p(\text{heads}) = \frac{498}{1002} = 0.497$$

$$p(\text{heads}) = \frac{\text{count}(\text{heads}) + \lambda}{\text{totalflips} + 2\lambda}$$

36

Maximum likelihood estimates

$$p(x_i | y) = \frac{\text{count}(x_i, y) + \lambda}{\text{count}(y) + \text{possible_values_of_}x_i * \lambda} \quad \lambda = 1$$

x1	x2	label
1	1	1
1	0	1
1	1	1
0	1	-1
0	0	-1

$p(x_1 = 1 1)$	?
$p(x_1 = 0 1)$	?
$p(x_2 = 1 1)$	?
$p(x_2 = 0 1)$	?

37

Maximum likelihood estimates

$$p(x_i | y) = \frac{\text{count}(x_i, y) + \lambda}{\text{count}(y) + \text{possible_values_of_}x_i * \lambda} \quad \lambda = 1$$

x1	x2	label
1	1	1
1	0	1
1	1	1
0	1	-1
0	0	-1

$p(x_1 = 1 1)$	4/5
$p(x_1 = 0 1)$	1/5
$p(x_2 = 1 1)$	3/5
$p(x_2 = 0 1)$	2/5

38

Avoids zero probability events!

MLE

$p(x_1 = 1 1)$	3/3
$p(x_1 = 0 1)$	0/3
$p(x_2 = 1 1)$	2/3
$p(x_2 = 0 1)$	1/3

smoothed/prior

$p(x_1 = 1 1)$	4/5
$p(x_1 = 0 1)$	1/5
$p(x_2 = 1 1)$	3/5
$p(x_2 = 0 1)$	2/5

39

Basic steps for probabilistic modeling

Step 1: pick a model

Step 2: figure out how to estimate the probabilities for the model

Step 3 (optional): deal with overfitting

Probabilistic models

Which model do we use, i.e. how do we calculate $p(\text{feature}, \text{label})$?

How do we train the model, i.e. how do we estimate the probabilities for the model?

How do we deal with overfitting?

40

Joint models vs conditional models

We've been trying to model the joint distribution (i.e. the data generating distribution):

$$p(x_1, x_2, \dots, x_m, y)$$

trying to model the data
generating distribution

However, if all we're interested in is classification, why not directly model the conditional distribution:

$$p(y | x_1, x_2, \dots, x_m)$$

direct model for classification

41

A first try: linear

$$p(y | x_1, x_2, \dots, x_m) = x_1 w_1 + w_2 x_2 + \dots + w_m x_m + b$$

Any problems with this?

- Nothing constrains it to be a probability
- Could still have combination of features and weight that exceeds 1 or is below 0

42

The challenge

$$x_1 w_1 + w_2 x_2 + \dots + w_m x_m + b$$

Linear model

$+\infty$



$-\infty$

$$p(y | x_1, x_2, \dots, x_m)$$

probability

1



0

We like linear models!

Can we transform the probability into a function that ranges over all values?

43

Odds ratio

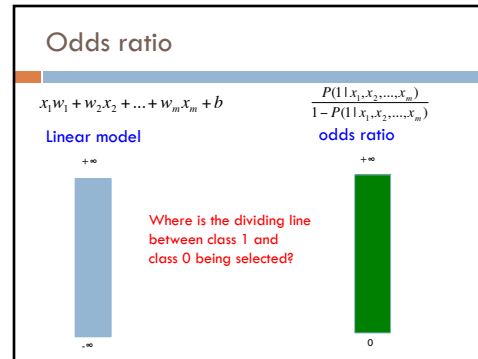
Rather than predict the probability, we can predict the ratio of 1/0 (positive/negative)

Predict the **odds** that it is 1 (true): How much more likely is 1 than 0.

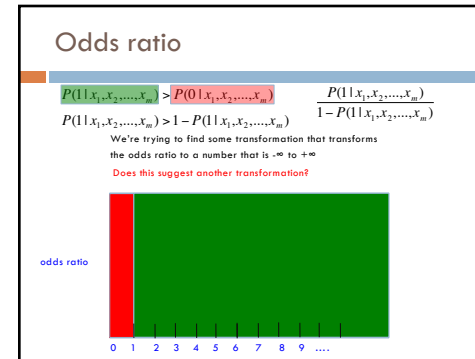
Does this help us?

$$\frac{P(1 | x_1, x_2, \dots, x_m)}{P(0 | x_1, x_2, \dots, x_m)} = \frac{P(1 | x_1, x_2, \dots, x_m)}{1 - P(1 | x_1, x_2, \dots, x_m)} = x_1 w_1 + w_2 x_2 + \dots + w_m x_m + b$$

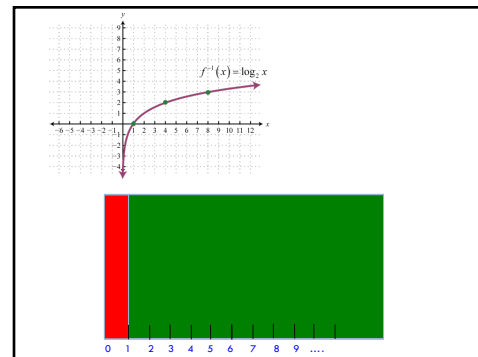
44



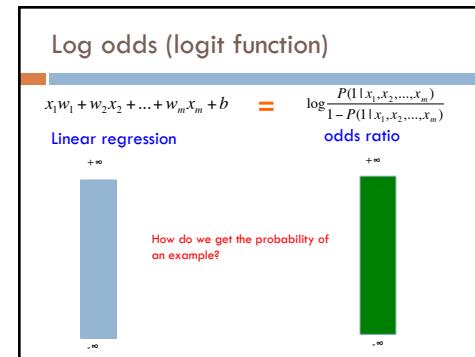
45



46



47



48

Log odds (logit function)

$$\log \frac{P(1|x_1, x_2, \dots, x_m)}{1 - P(1|x_1, x_2, \dots, x_m)} = w_1 x_1 + w_2 x_2 + \dots + w_m x_m + b$$

$$\frac{P(1|x_1, x_2, \dots, x_m)}{1 - P(1|x_1, x_2, \dots, x_m)} = e^{w_1 x_1 + w_2 x_2 + \dots + w_m x_m + b}$$

$$P(1|x_1, x_2, \dots, x_m) = (1 - P(1|x_1, x_2, \dots, x_m)) e^{w_1 x_1 + w_2 x_2 + \dots + w_m x_m + b}$$

...

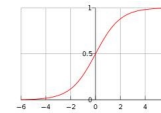
$$P(1|x_1, x_2, \dots, x_m) = \frac{1}{1 + e^{-(w_1 x_1 + w_2 x_2 + \dots + w_m x_m + b)}}$$

anyone
recognize
this?

49

Logistic function

$$\text{logistic} = \frac{1}{1 + e^{-x}}$$



50

Logistic regression

How would we classify examples once we had a trained model?

$$\log \frac{P(1|x_1, x_2, \dots, x_m)}{1 - P(1|x_1, x_2, \dots, x_m)} = w_1 x_1 + w_2 x_2 + \dots + w_m x_m + b$$

If the sum > 0 then $p(1)/p(0) > 1$, so positive

if the sum < 0 then $p(1)/p(0) < 1$, so negative

Still a *linear* classifier (decision boundary is a line)

51

Training logistic regression models

How should we learn the parameters for logistic regression (i.e. the w 's and b)?

$$\log \frac{P(1|x_1, x_2, \dots, x_m)}{1 - P(1|x_1, x_2, \dots, x_m)} = w_1 x_1 + w_2 x_2 + \dots + w_m x_m + b$$

parameters

$$P(1|x_1, x_2, \dots, x_m) = \frac{1}{1 + e^{-(w_1 x_1 + w_2 x_2 + \dots + w_m x_m + b)}}$$

52

MLE logistic regression

Find the parameters that maximize the likelihood (or log-likelihood) of the data:

$$\begin{aligned}\text{log-likelihood} &= \sum_{i=1}^n \log(p(x_i)) \\ &= \sum_{i=1}^n \log\left(\frac{1}{1 + e^{-y_i(w_0x_0 + w_1x_1 + \dots + w_mx_m + 1)}}\right) \quad \text{assume labels } 1, -1 \\ &= \sum_{i=1}^n -\log(1 + e^{-y_i(w_0x_0 + w_1x_1 + \dots + w_mx_m + 1)})\end{aligned}$$

53

MLE logistic regression

$$\text{log-likelihood} = \sum_{i=1}^n -\log(1 + e^{-y_i(w_0x_0 + w_1x_1 + \dots + w_mx_m + 1)})$$

We want to maximize, i.e.

$$\begin{aligned}MLE(\text{data}) &= \operatorname{argmax}_{w,b} \text{log-likelihood}(\text{data}) \\ &= \operatorname{argmax}_{w,b} \sum_{i=1}^n -\log(1 + e^{-y_i(w_0x_0 + w_1x_1 + \dots + w_mx_m + 1)}) \\ &= \operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_0x_0 + w_1x_1 + \dots + w_mx_m + 1)})\end{aligned}$$

Look familiar? Hint: anybody reading the book?

54

MLE logistic regression

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_0x_0 + w_1x_1 + \dots + w_mx_m + 1)})$$

Surrogate loss functions:

Zero/one: $\ell^{(0/1)}(y, \hat{y}) = \mathbf{1}[y\hat{y} \leq 0]$

Hinge: $\ell^{(\text{hinge})}(y, \hat{y}) = \max\{0, 1 - y\hat{y}\}$

Logistic: $\ell^{(\text{log})}(y, \hat{y}) = \frac{1}{\log 2} \log(1 + \exp[-y\hat{y}])$

Exponential: $\ell^{(\text{exp})}(y, \hat{y}) = \exp[-y\hat{y}]$

Squared: $\ell^{(\text{sq})}(y, \hat{y}) = (y - \hat{y})^2$

55

logistic regression: three views

$$\log \frac{P(1|x_1, x_2, \dots, x_m)}{1 - P(1|x_1, x_2, \dots, x_m)} = w_0 + w_1x_1 + w_2x_2 + \dots + w_mx_m \quad \text{linear classifier}$$

$$P(1|x_1, x_2, \dots, x_m) = \frac{1}{1 + e^{-(w_0 + w_1x_1 + w_2x_2 + \dots + w_mx_m)}} \quad \text{conditional model logistic}$$

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_0x_0 + w_1x_1 + \dots + w_mx_m + 1)}) \quad \text{linear model minimizing logistic loss}$$

56

Overfitting

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_1x_1 + w_2x_2 + \dots + w_mx_m + b)})$$

If we minimize this loss function, in practice, the results aren't great and we tend to overfit

Solution?

57

Regularization/prior

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_1x_1 + w_2x_2 + \dots + w_mx_m + b)}) + \lambda \operatorname{regularizer}(w,b)$$

or

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_1x_1 + w_2x_2 + \dots + w_mx_m + b)}) - \log(p(w,b))$$

What are some of the regularizers we know?

58

Regularization/prior

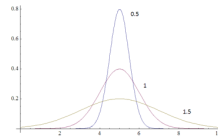
L2 regularization:

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_1x_1 + w_2x_2 + \dots + w_mx_m + b)}) + \lambda \|w\|^2$$

Gaussian prior:

Gaussians are defined by a mean (μ) and a variance (σ^2)

$p(w,b) \sim$



59

Regularization/prior

L2 regularization:

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_1x_1 + w_2x_2 + \dots + w_mx_m + b)}) + \lambda \|w\|^2$$

Gaussian prior:

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_1x_1 + w_2x_2 + \dots + w_mx_m + b)}) + \frac{1}{2\sigma^2} \|w\|^2$$

Does the λ make sense? $\lambda = \frac{1}{2\sigma^2}$

60

Regularization/prior

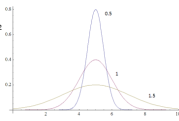
L2 regularization:

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_1x_1 + w_2x_2 + \dots + w_mx_m + b)}) + \lambda \|w\|^2$$

Gaussian prior:

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_1x_1 + w_2x_2 + \dots + w_mx_m + b)}) + \frac{1}{2\sigma^2} \|w\|^2$$

$$\lambda = \frac{1}{2\sigma^2}$$



61

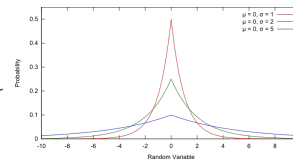
Regularization/prior

L1 regularization:

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_1x_1 + w_2x_2 + \dots + w_mx_m + b)}) + \lambda \|w\|$$

Laplacian prior:

$$p(w,b) \sim$$



62

Regularization/prior

L1 regularization:

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_1x_1 + w_2x_2 + \dots + w_mx_m + b)}) + \lambda \|w\|$$

Laplacian prior:

$$\operatorname{argmin}_{w,b} \sum_{i=1}^n \log(1 + e^{-y_i(w_1x_1 + w_2x_2 + \dots + w_mx_m + b)}) + \frac{1}{\sigma} \|w\|$$

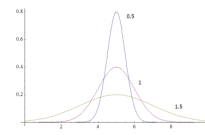
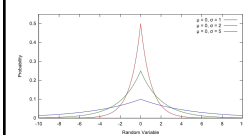
$$\lambda = \frac{1}{\sigma}$$

63

L1 vs. L2

L1 = Laplacian prior

L2 = Gaussian prior



64

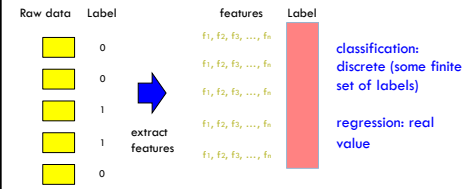
Logistic regression

Why is it called logistic regression?
It is a classifier??

$$\log \frac{P(1|x_1, x_2, \dots, x_m)}{1 - P(1|x_1, x_2, \dots, x_m)} = w_1 x_1 + w_2 x_2 + \dots + w_m x_m + b$$

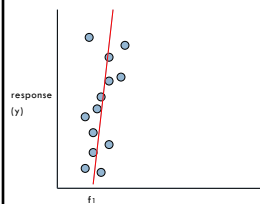
65

A digression: regression vs. classification



66

linear regression



How can we find this line?

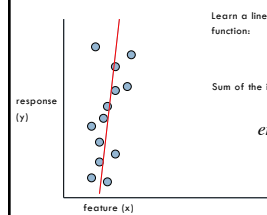
Given some points, find the **line** that best fits/explains the data

Our model is a line, i.e. we're assuming a linear relationship between the feature and the label value

$$h(y) = w_1 x_1 + b$$

67

Linear regression



Learn a line h that minimizes some loss/error function:

$$error(h) = ?$$

Sum of the individual errors:

$$error(h) = \sum_{i=1}^n \underbrace{|y_i - h(f_i)|}_{0/1 \text{ loss}}$$

68

Error minimization

How do we find the minimum of an equation?

$$\text{error}(h) = \sum_{i=1}^n |y_i - h(f_i)|$$

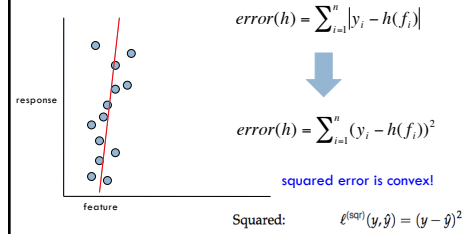
Take the derivative, set to 0 and solve (going to be a min or a max)

Any problems here?

Ideas?

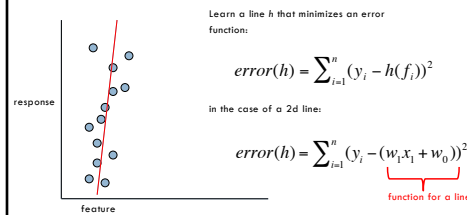
69

Linear regression



70

Linear regression



71

Linear regression

We'd like to *minimize* the error

Find w_1 and w_0 such that the error is minimized

$$\text{error}(h) = \sum_{i=1}^n (y_i - (w_1 f_i + w_0))^2$$

We can solve this in closed form

72

Multiple linear regression

If we have m features, then we have a line in m dimensions

$$h(\vec{f}) = w_0 + w_1 f_1 + w_2 f_2 + \dots + w_m f_m$$

weights

73

Multiple linear regression

We can still calculate the squared error like before

$$h(\vec{f}) = w_0 + w_1 f_1 + w_2 f_2 + \dots + w_m f_m$$

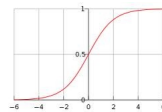
$$error(h) = \sum_{i=1}^n (y_i - (w_0 + w_1 f_{i1} + w_2 f_{i2} + \dots + w_m f_{im}))^2$$

Still can solve this exactly!

74

Logistic function

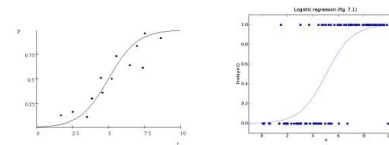
$$\text{logistic} = \frac{1}{1 + e^{-x}}$$



75

Logistic regression

Find the best fit of the data based on a logistic



76

Basic steps for probabilistic modeling

Step 1: pick a model

Step 2: figure out how to estimate the probabilities for the model

Step 3 (optional): deal with overfitting

Probabilistic models

Which model do we use, i.e. how do we calculate $p(\text{feature}, \text{label})$?

How do train the model, i.e. how to we we **estimate the probabilities** for the model?

How do we deal with overfitting?

77

Probabilistic models summarized

Two classification models:

- Naïve Bayes (models **joint** distribution)
- Logistic Regression (models **conditional** distribution)
 - In practice this tends to work better if all you want to do is classify

Priors/smoothing/regularization

- Important for both models
- In theory: allow us to impart some prior knowledge
- In practice: avoids overfitting and often tune on development data

78