

Today: Multinomial (multiclass) logistic regression
Regularization
Convex functions & sets w/ examples

NB: HW2 is posted, due next Monday.

You will need to:

-use Jupyter, Scikit-learn, matplotlib, numpy

Multinomial Logistic Reg

Given features $x \in \mathbb{R}^d$ goal is to predict class $y \in [K] = \{1, \dots, K\}$

Ex: given image $x \in \mathbb{R}^{m \times n}$ determine what object it contains

Model: assume the log-odds (logits or scores) of each class given x relative to class K are linear:

$$\log \frac{P(Y=1 | X=x)}{P(Y=K | X=x)} = b_1 + \omega_1^T x$$

$$\log \frac{P(Y=2|X=x)}{P(Y=k|X=x)} = b_2 + \omega_2^T x$$

$$\log \frac{P(Y=k-1|X=x)}{P(Y=k|X=x)} = b_{k-1} + \omega_{k-1}^T x$$

These $k-1$ equations plus the constraint that $\sum_{i=1}^k P(Y=i) = 1$

imply

$$P(Y=i|X=x) = \frac{\exp(b_i + \omega_i^T x)}{1 + \sum_{l=1}^{k-1} \exp(b_l + \omega_l^T x)} \quad i=1, \dots, k-1$$

$$P(Y=k|X=x) = \frac{1}{1 + \sum_{l=1}^{k-1} \exp(b_l + \omega_l^T x)}$$

For ease of use, we can introduce some unnecessary parameters b_k, ω_k to get the simpler expressions

$$P(Y=i | X=x) = \frac{\exp(b_i + \omega_i^T x)}{\sum_{l=1}^k \exp(b_l + \omega_l^T x)} = \frac{\exp(\omega x + b)_i}{\mathbf{1}^T \exp(\omega x + b)}$$

where $\underline{b} \in \mathbb{R}^k$ and $\omega \in \mathbb{R}^{k \times d}$ $\omega = \begin{bmatrix} \omega_1^T \\ \vdots \\ \omega_k^T \end{bmatrix}$

are the model parameters

and if $\underline{v}$ is a vector then $\exp(\underline{v}) \in \mathbb{R}^k$ is the componentwise

exponentiation of $\underline{v}$, and $\mathbf{1} = \mathbf{1}_k$ is the k -dim vector of all ones

$$P(\varphi = \underline{i} | X = x) = \underline{\text{softmax}}(\underline{W}x + \underline{b})_{\underline{i}}$$

where $\text{softmax} : v \mapsto \frac{\exp(v)}{\mathbf{1}^T \exp(v)}$

$$P(\varphi = \underline{j} | X = x_i) = \frac{\text{softmax}(Wx_i + b)_j}{j}$$

softmax is a continuous approximation of the function that maps a vector v to a vector of zeros and ones where there is a one corresp to the largest entry in v , and a zero everywhere else.

Notice that $\text{softmax}(v)$ is always a probability vector, and softmax is defined on all of $\mathbb{R}^k$. (multidim analog of the sigmoid)

Ex:

$$\begin{bmatrix} 30 \\ 1 \\ 0 \end{bmatrix}$$

logits

$\xrightarrow{\text{softmax}}$

$$\begin{bmatrix} 1 \\ 2.54 \times 10^{-13} \\ 9.36 \times 10^{-14} \end{bmatrix}$$

prob dist over classes

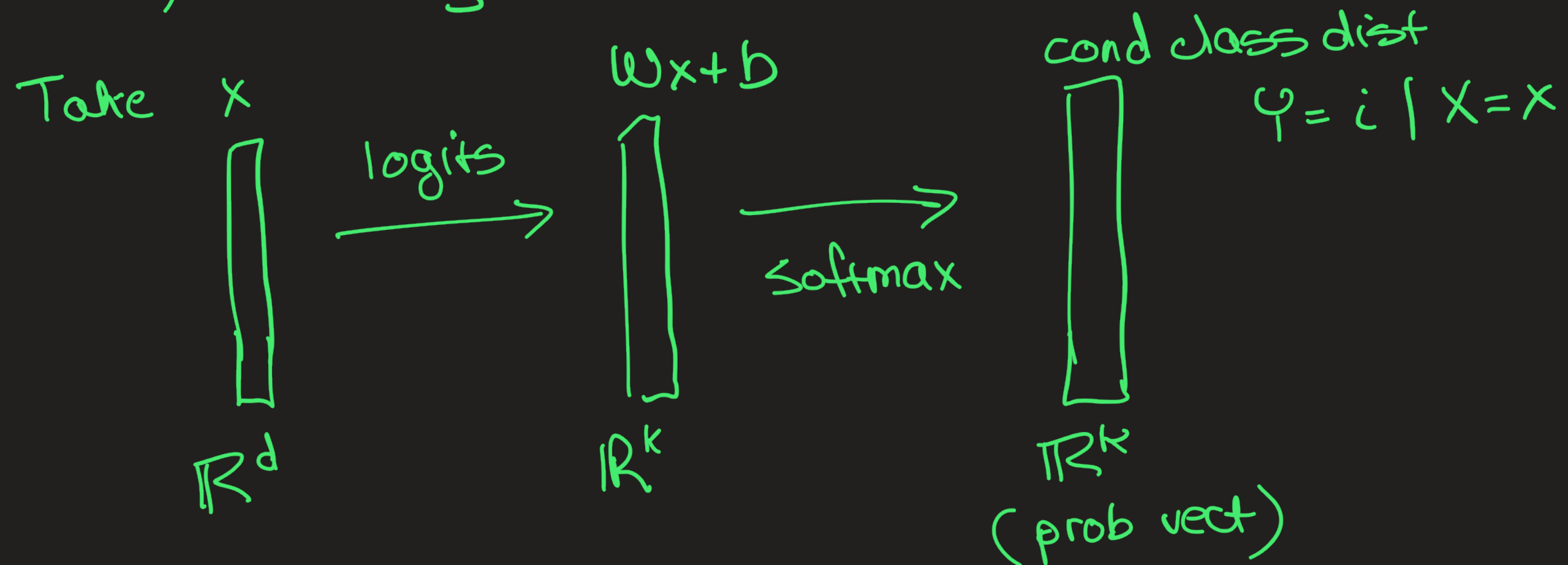
$$\begin{bmatrix} .9 \\ .05 \\ .05 \end{bmatrix}$$

logits

$\rightarrow$

$$\begin{bmatrix} 0.539 \\ 0.230 \\ 0.230 \end{bmatrix}$$

Using multiclass logistic reg. to predict given new data, assuming we know w & b .



Then classify x as having class

$$\begin{aligned} i^* &= \underset{i}{\operatorname{argmax}} \mathbb{P}(\varphi = i \mid X = x) \\ &= \underset{i}{\operatorname{argmax}} \operatorname{softmax}(w x + b)_i \end{aligned}$$

How do we determine w & b ?

As before, we maximize the likelihood of our model parameters given the training data; we do this by minimizing the negative log-likelihood (NLL)

$$\mathcal{L}(w, b) := \sum_{i=1}^n -\log \underbrace{P(Y=y_i | X=x_i; w, b)}$$

$$w^*, b^* = \underset{w, b}{\operatorname{argmin}} \mathcal{L}(w, b)$$

To write the NLL more conveniently, encode the class observations as one-hot vectors. Instead of $y_i \in [K]$ consider $y_i \in \{e_1, \dots, e_K\} \subseteq \mathbb{R}^K$ e.g. if 2 was the class of training example i , $y_i = e_2 = \begin{bmatrix} 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix} \in \mathbb{R}^K$

Now we see that if the i th training example has class j ,
then

$$\log P(Y = y_i | X = x_i) = \log [\text{softmax}(\omega x_i + b)_j]$$

$$= y_i^T \log(\text{softmax}(\omega x_i + b))$$

$$= y_i^T \log\left(\frac{\exp(\omega x_i + b)}{\mathbf{1}^T \exp(\omega x_i + b)}\right)$$

$$\rightarrow = \frac{y_i^T (\omega x_i + b)}{\underbrace{\log\left(\sum_{l=1}^k \exp(\omega x_i + b)_l\right)}_{\text{logsumexp}}}$$

$$= y_i^T (\omega x_i + b) - \text{logsumexp}(\omega x_i + b)$$

where $\text{logsumexp}(v) = \log\left(\sum_{i=1}^k \exp(v_i)\right)$ when $v \in \mathbb{R}^k$

$k=3$

Example

$$y^T \log \left(\frac{\exp(\omega x + b)}{I^T \exp(\omega x + b)} \right) = \left\langle \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} \log(\exp(\omega_1^T x + b_1)/Z) \\ \log(\exp(\omega_2^T x + b_2)/Z) \\ \log(\exp(\omega_3^T x + b_3)/Z) \end{bmatrix} \right\rangle$$

$$= \log(\exp(\omega_1^T x + b_1)/Z)$$

$$= \underbrace{(\omega_1^T x + b_1)}_{\text{logit}} - \underline{\underline{\log(Z)}}$$

We end up with the ERM problem:

$$w^*, b^* = \underset{w, b}{\operatorname{argmin}} \quad - \sum_{i=1}^n \left[y_i^T (w x_i + b) - \log \left(\sum_{l=1}^k \exp(w x_i + b)_l \right) \right]$$

Recall: $x^T y = \sum_{i=1}^k x_i y_i = \langle x, y \rangle$

Again: there is no reason why you have to model the logits/scores as linear in x . We could use an arbitrary function

$$f: \mathbb{R}^d \rightarrow \mathbb{R}^k$$

$$f^* = \underset{f \in \mathcal{F}}{\operatorname{argmin}} \quad - \sum_{i=1}^n \left[y_i^T f(x_i) - \log \left(\sum_{l=1}^k \exp(f(x_i))_l \right) \right]$$

Regularization & Overfitting

In ML we often don't have a strong idea of what is the ideal model for our data, so we try a bunch of different models and choose the one that gives better performance. If we pick too complicated/powerful of a class of models, we can overfit: have

$$\boxed{R_n(\hat{f}_T)} \ll \underline{R(f^*)}$$

a large generalization gap. False confidence that your model has low risk.

To avoid overfitting:

- 1) use Occam's razor: use very simple models when we can
- 2) use regularization

structural minimization

Regularized ERM is obtained by adding a regularization term to your objective to encourage your complicated model to be simple

Ex: we can encourage the model to have small coeffs by adding an ℓ_2 regularization term

$$w_* = \underset{w}{\operatorname{argmin}} \quad \underbrace{\frac{1}{n} \sum_{i=1}^n (y_i - w^T x_i)^2}_{\text{data-fit term}} + \underbrace{\lambda \|w\|_2^2}_{\text{regularizer}}$$

$\lambda \rightarrow$ regularization parameter, > 0

This formulation (ℓ_2 regularization + ℓ_2 loss) is called ridge regression

Recall OLS estimation : Assume $y \sim \mathcal{N}(\underline{\omega^T x}, \sigma^2)$
↓
ordinary least squares

then determine ω^* using MLE

$$\Rightarrow \omega^* = \underset{\omega}{\operatorname{argmin}} \frac{1}{n} \sum_{i=1}^n (y_i - \omega^T x_i)^2$$

$$= \underset{\omega}{\operatorname{argmin}} \frac{1}{n} \|X_{\text{tr}} \omega - y_{\text{tr}}\|_2^2$$

$$= (X_{\text{tr}}^T X_{\text{tr}})^{-1} X_{\text{tr}}^T y_{\text{tr}}$$

$$X_{\text{tr}} = \begin{bmatrix} x_1^T \\ \vdots \\ x_n^T \end{bmatrix} \in \mathbb{R}^{n \times d}$$

$$y_{\text{tr}} = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} \in \mathbb{R}^n$$

notice Bayes optimal estimator (b/c we're using ℓ_2 -loss)
for y given x is $\mathbb{E}[y|x] = \omega^T x$

Ridge Regression (ℓ_2 loss + ℓ_2 regularization)

$$\omega^* = \underset{\omega}{\operatorname{argmin}} \frac{1}{n} \|X_{tr} \omega - y_{tr}\|_2^2 + \lambda \|\omega\|_2^2$$

$$= (X_{tr}^T X_{tr} + n \lambda I)^{-1} X_{tr}^T y_{tr}$$

l_1 -regularization (aka Lasso regularization)

Goal: we want only a few nonzero model weights

$$\|w\|_1 = \sum_{i=1}^d |w_i|$$

LASSO (Least Absolute Shrinkage & Selection Operator)

l_2 loss + l_1 regularization

$$w^* = \arg \min_w \frac{1}{n} \|X_{tr} w - y_{tr}\|_2^2 + \lambda \|w\|_1$$

This does not have a closed-form solution.

Regularization is a powerful technique that is ubiquitous in ML. In general we use the formulation

$$\hat{f}_\lambda = \operatorname{argmin}_{f \in \mathcal{F}} R_n(f) + \lambda \mathcal{J}(f)$$

where $\mathcal{J}: \mathcal{F} \rightarrow \mathbb{R}$ is the regularizer and λ is the regularization parameter. This is called regularized empirical risk minimization (RERM)

Ex: polynomial (quadratic) models vs linear models

regression against $[1; x]$ — linear model

vs

regression against $[1; x; x \otimes x]$ where

$$x \otimes x = \begin{bmatrix} x_1^2 \\ x_1 x_2 \\ x_1 x_3 \\ \vdots \\ x_1 x_d \\ x_2 x_1 \\ \vdots \\ x_2 x_d \\ \vdots \\ x_d x_1 \\ \vdots \\ x_d^2 \end{bmatrix} \in \mathbb{R}^{d^2}$$

See the sample code on the ^{course} website

How to choose the regularization parameter λ ?

In theory as $n \rightarrow \infty$, "let the data speak" by

taking $\lambda = \Theta\left(\frac{1}{\sqrt{n}}\right)$. The idea is that even if $\mathcal{F}$ is complicated, we can avoid overfitting by seeing enough samples from the population $\mathbb{P}$.

In practice try various values of λ and choose the one that minimizes the gap in the risk between the train & test (really validation) data sets.

