

ML & Opt Lecture 6

- Practical considerations:
 - Train/valid/test splits
 - Cross-validation
 - Hyperparameter selection
 - Regularization (overfitting)
- Revisiting binary/multiclass logistic regression, linear separability

Logistic Regression

$x \in \mathbb{R}^d$ features use to predict $y \in \{-1, 1\}$

Assume that the log odds of y being 1 is linear in x ($p = \mathbb{P}(Y=1|x)$)

$$\ln\left(\frac{p}{1-p}\right) = \omega_0 + \sum_{i=1}^d \omega_i x_i = \omega_0 + \omega^T x$$

$$= \tilde{\omega}^T \tilde{x} \quad \text{where}$$

$\uparrow$
equivalent

$$\tilde{x} = \begin{bmatrix} x \\ 1 \end{bmatrix}$$

$$\text{and } \tilde{\omega} = \begin{bmatrix} \omega \\ \omega_0 \end{bmatrix}$$

Interpretation: probability of $Y=1$

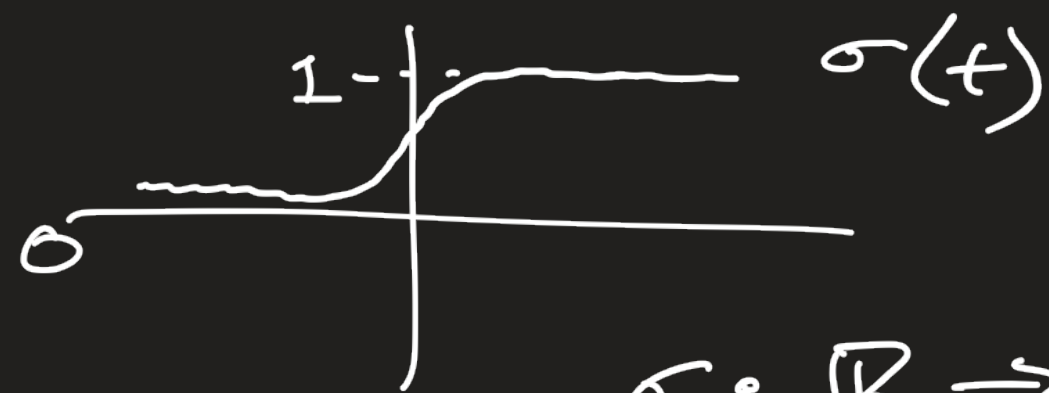
goes up if $\omega_i > 0$ and x_i increases;

goes down if $\omega_i < 0$ and x_i increases;

and ω_0 determines the baseline prob of $Y=1$

$$\Rightarrow \frac{p}{1-p} = \exp(\omega_0 + \omega^T x) \Rightarrow p = \frac{1}{1 + e^{-\omega_0 + \omega^T x}}$$

$$\Rightarrow p = \sigma(\underbrace{\omega_0 + \omega^T x}_{\text{called the logit}}) \quad \text{where} \quad \sigma(t) = \frac{1}{1 + e^{-t}}$$



$$\sigma: \mathbb{R} \rightarrow [0, 1]$$

This is the same as saying

$$y \sim \text{Bern}(\sigma(\omega_0 + \omega^T x))$$

Goal: recover ω_0 & ω given training data so
in the future when we observe x we
can predict y

Notice we predict y by choosing $c \in \{-1, 1\}$ that maximizes

$$p(y=c|x) = \begin{cases} \sigma(\tilde{\omega}^T \tilde{x}) & c=1 \\ 1 - \sigma(\tilde{\omega}^T \tilde{x}) & c=-1 \end{cases}$$

notice

$$1 - \sigma(t) = 1 - \frac{1}{1 + e^{-t}} = \frac{1 + e^{-t}}{1 + e^{-t}} - \frac{1}{1 + e^{-t}}$$

$$= \frac{e^{-t}}{1 + e^{-t}} = \frac{1}{e^t + 1} = \sigma(-t)$$

so

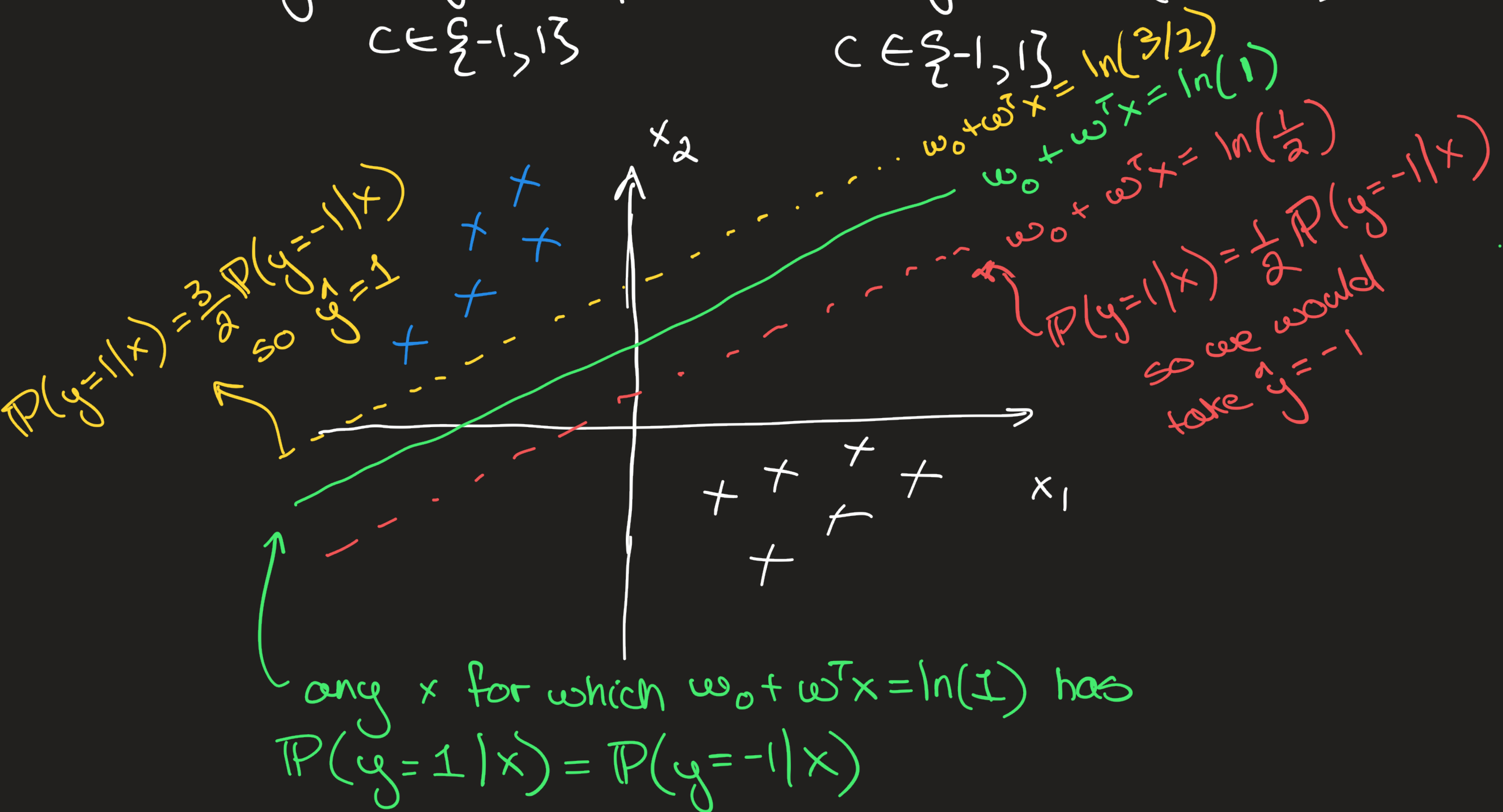
$$\boxed{1 - \sigma(t) = \sigma(-t)}$$

This implies

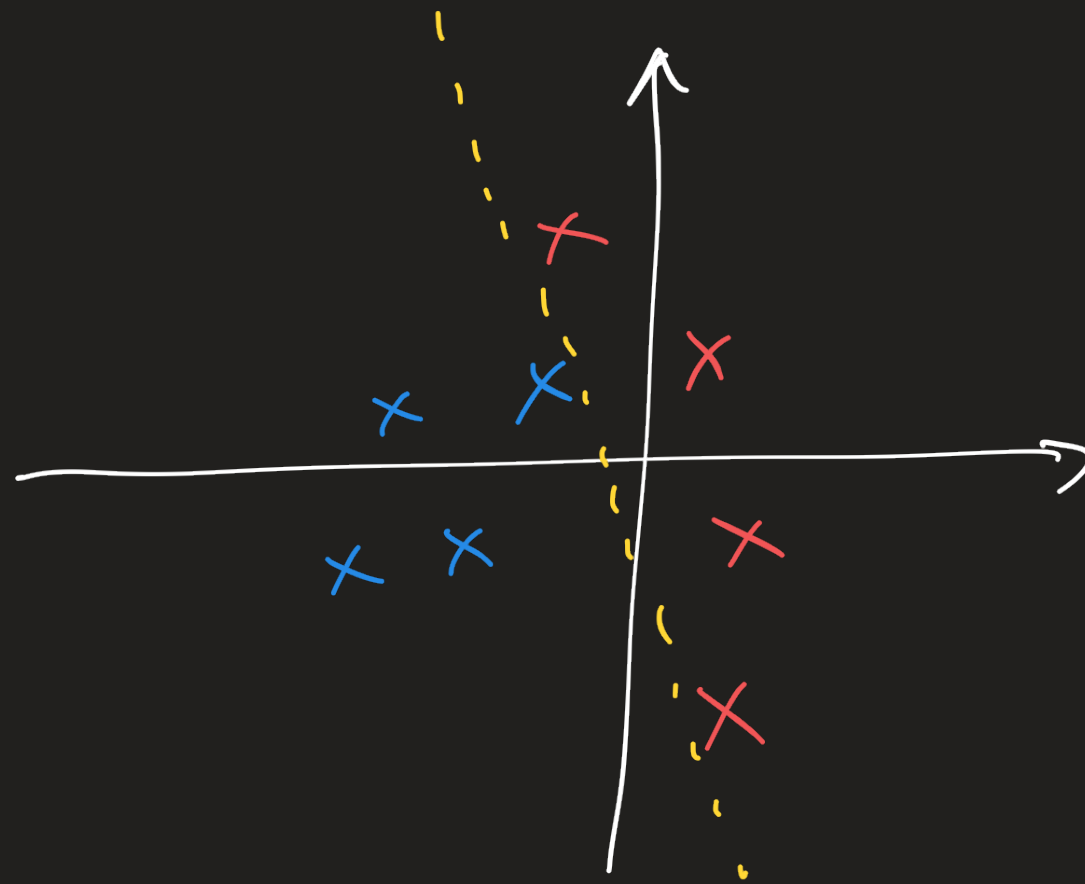
$$p(y=c|x) = \begin{cases} \sigma(\tilde{\omega}^T \tilde{x}) & c=1 \\ \sigma(-\tilde{\omega}^T \tilde{x}) & c=-1 \end{cases} = \sigma(c \tilde{\omega}^T \tilde{x})$$

Now we see that we can predict y , given x , with

$$\hat{y} = \operatorname{argmax}_{c \in \{-1, 1\}} p(c|x) = \operatorname{argmax}_{c \in \{-1, 1\}} \sigma(c \tilde{w}^T \tilde{x})$$



Question: given training data, how to estimate w_0, w



Use MLE to find best separating hyperplane

$$\mathcal{L}(w_0, w) = \frac{1}{n} \sum_{i=1}^n -\log p(y_i | x_i; w_0, w)$$

↑
NLL of the model parameters w_0, w given the observed (training) data

Notice

$$\mathcal{L}(\omega_0, \omega) = \frac{1}{n} \sum_{i=1}^n -\log \sigma(y_i \tilde{\omega}^T \tilde{x}_i)$$

$$= \frac{1}{n} \sum_{i=1}^n \log \left(\frac{1}{\sigma(y_i \tilde{\omega}^T \tilde{x}_i)} \right)$$

$$= \frac{1}{n} \sum_{i=1}^n \log (1 + e^{-y_i \tilde{\omega}^T \tilde{x}_i})$$

take $\ell(u, v) = \log(1 + e^{-uv})$ called the logistic loss

$$= \frac{1}{n} \sum_{i=1}^n \ell(\tilde{\omega}^T \tilde{x}_i, y_i)$$

So fitting the model parameters for logistic regression using MLE is done by solving

$$\hat{w}_0, \hat{w} = \underset{\substack{w_0 \in \mathbb{R} \\ w \in \mathbb{R}^d}}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n \ell(w_0 + w^T x_i, y_i)$$

where $\ell(u, v) \equiv \log(1 + e^{-uv})$ is the logistic loss fn.

- notice logistic regression always has one solution (when it has a finite solution)

ASIDE: if the training data is actually linearly separable, and w_0, w separate the

data, then $\alpha w_0, \alpha w$ also separate the data, and I can choose $\alpha \rightarrow \pm\infty$ so

that $\ell(\alpha w_0, \alpha w) \rightarrow 0$

$$\begin{aligned} w_0 + w^T x &= c \\ \downarrow \\ d w_0 + d w^T x &= d c \end{aligned}$$

Multiclass Logistic Regression

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

Ex: given pixel image $x \in [256]^{m \times n}$ determine the type of image

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

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

$\vdots$

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

Parameters



$b \in \mathbb{R}^{k-1}$



$\omega \in \mathbb{R}^{(k-1) \times d}$

These $k-1$ equations plus the fact that $\sum_{i=1}^k \mathbb{P}(\varphi=i|x) = 1$ gives us

$$\mathbb{P}(\varphi=i | X=x) = \frac{\exp(b_i + \omega_i^\top x)}{1 + \sum_{\ell=1}^{k-1} \exp(b_\ell + \omega_\ell^\top x)} \quad i=1, \dots, k-1$$

$$\mathbb{P}(\varphi=k | X=x) = \frac{1}{1 + \sum_{\ell=1}^{k-1} \exp(b_\ell + \omega_\ell^\top x)}$$

For ease of use we introduce nonzero logits for class k ,
 i.e. note that

$$\begin{aligned}
 P(Y=k|X) &= \frac{1}{1 + \sum_{i=1}^{k-1} \exp(b_i + \omega_i^T x)} \cdot \frac{\exp(b_k + \omega_k^T x)}{\exp(b_k + \omega_k^T x)} \\
 &= \frac{\exp(b_k + \omega_k^T x)}{\exp(b_k + \omega_k^T x) + \sum_{i=1}^{k-1} \exp((b_i + b_k) + (\omega_i + \omega_k)^T x)}
 \end{aligned}$$

similarly we have

$$P(Y=i|X) = \frac{\exp((b_i + b_k) + (\omega_k + \omega_i)^T x)}{\exp(b_k + \omega_k^T x) + \sum_{i=1}^{k-1} \exp((b_i + b_k) + (\omega_i + \omega_k)^T x)}$$

We relabel the shifted biases and weights

$$b_k \leftarrow b_k$$

$$\omega_k \leftarrow \omega_k$$

$$\left. \begin{array}{l} b_i \leftarrow b_i + b_k \\ \omega_i \leftarrow \omega_i + \omega_k \end{array} \right\} \text{ for } i \neq k$$

logit for class i

then we have

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

$$\sum_{l=1}^k \exp(b_l + \omega_l^T x)$$

sum exp
of the logits

$$= \frac{(e^{b + \omega^T x})_i}{\mathbf{1}^T e^{b + \omega^T x}}$$

where $\text{softmax}(v) = \frac{e^v}{\mathbf{1}^T e^v}$

$$= \text{softmax}(b + \omega^T x)_i$$

Recall $\sigma: \mathbb{R} \rightarrow [0, 1]$
logits probability $y = I$

Now for multiclass logistic regression

$\text{softmax}: \mathbb{R}^k \rightarrow$ set of probability vectors
logit vector over the k classes
for k classes

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 corresponding to the largest entry in v , and a zero everywhere else

Ex:

$$\begin{bmatrix} 30 \\ 1 \\ 0 \end{bmatrix} \xrightarrow{\text{softmax}} \begin{bmatrix} 1 \\ 2.54 \times 10^{-13} \\ 9.36 \times 10^{-14} \end{bmatrix}$$

logits class probs

Ex:

$$\begin{bmatrix} 0.9 \\ 0.05 \\ 0.05 \end{bmatrix} \xrightarrow{\text{softmax}} \begin{bmatrix} 0.539 \\ 0.230 \\ 0.230 \end{bmatrix}$$

logits class probs

Given that my data follows the label dist

$$P(Y=i|x) = \text{softmax}(Wx+b)_i$$

how do I predict the class label given x ?

Take x

$Wx+b$

$\mathbb{R}^d$

compute
logits

$\mathbb{R}^k$

compute
softmax
probs

$\mathbb{R}_k$

$$P(Y=i|x) = \text{softmax}(Wx+b)_i$$

and take x as having class

$$\hat{y} = \underset{i \in [K]}{\operatorname{argmax}} P(Y=i|x)$$

$$= \underset{i \in [K]}{\operatorname{argmax}} (Wx+b)_i$$

We determine ω & b given training data by minimizing the NLL (i.e. MLE):

$$\mathcal{L}(\omega, b) := \frac{1}{n} \sum_{i=1}^n -\log \mathbb{P}[Y=y_i | X=x_i; \omega, b]$$

and take

$$\hat{\omega}, \hat{b} = \underset{\omega, b}{\operatorname{argmin}} \mathcal{L}(\omega, b)$$

To write the NLL more conveniently, encode the class observations as one-hot vectors. Instead of $y_i \in [k]$, consider $y_i \in \{e_1, \dots, e_k\} \subseteq \mathbb{R}^k$.

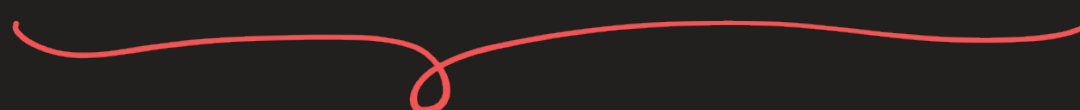
E.g. if training example i is of class 2, we take

$$y_i = e_2 = \begin{bmatrix} 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix} \in \mathbb{R}^k$$

Now we see that if the i th training example has class j , then $y_i = e_j$ and

$$\begin{aligned}\log \mathbb{P}[Y=y_i | X=x_i] &= \log \mathbb{P}(Y=e_j | X=x_i) \\&= \log \text{softmax}(\omega x_i + b)_j \\&= \log e_j^T \text{softmax}(\omega x_i + b) \\&= \log y_i^T \text{softmax}(\omega x_i + b) \\&= y_i^T \log \left(\text{softmax}(\omega x_i + b) \right) \\&= y_i^T \log \left(\frac{\exp(\omega x_i + b)}{1^T \exp(\omega x_i + b)} \right)\end{aligned}$$

$$\Rightarrow \log P[Y=y_i | X=x_i] = y_i^T (\omega x_i + b) - \log \left(\sum_{l=1}^k \exp(\omega x_i + b)_l \right)$$



 $\log \text{sumexp}(\omega x_i + b)$

where $\log \text{sumexp} : v \mapsto \log(\mathbf{1}^T e^v)$

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

Therefore we get the MLE estimates for the model parameters by solving

$$\hat{\omega}, \hat{b} = \underset{\omega, b}{\operatorname{argmin}} -\frac{1}{n} \sum_{i=1}^n \left[y_i^T (\omega x_i + b) - \log \left(\sum_{l=1}^K \exp(\omega x_i + b)_l \right) \right]$$

question: what is the loss function so that

$$\frac{1}{n} \sum_{i=1}^n \ell(y_i, \omega x_i + b) = \mathcal{L}(\omega, b) ?$$