

Logistic Regression

Input: d -dimension feature vector $x \in \mathbb{R}^d$

Output: A class or label $l \in \{1, \dots, k\}$

Ex: Image classification:

Given input image $x \in \mathbb{R}^{n \times n}$, predict a label for the image (e.g. cat/dog, MNIST: $\{0, \dots, 9\}$, CIFAR-10 $\{0, \dots, 9\}$)

Ex: Binary Classification:

Given some medical data $x \in \mathbb{R}^n$
try to predict incidence of heart disease
Label: $\{0/1\}$

Logistic Regression model

Given input feature vector $x \in \mathbb{R}^n$,
the model predicts a prob. distribution over the labels $\{1, \dots, k\}$

$$\text{softmax} \left(\underbrace{\begin{matrix} \mathbb{R}^{k \times n} & \mathbb{R}^n & \mathbb{R}^k \\ \uparrow & \uparrow & \uparrow \\ A & x & b \end{matrix}}_{\text{affine linear map}} \right)$$

affine linear map: $\mathbb{R}^n \rightarrow \mathbb{R}^k \xrightarrow{\text{softmax}} \mathcal{P}(\{1, \dots, k\})$

softmax: Input: $y \in \mathbb{R}^k = \begin{pmatrix} y_1 \\ \vdots \\ y_k \end{pmatrix}$

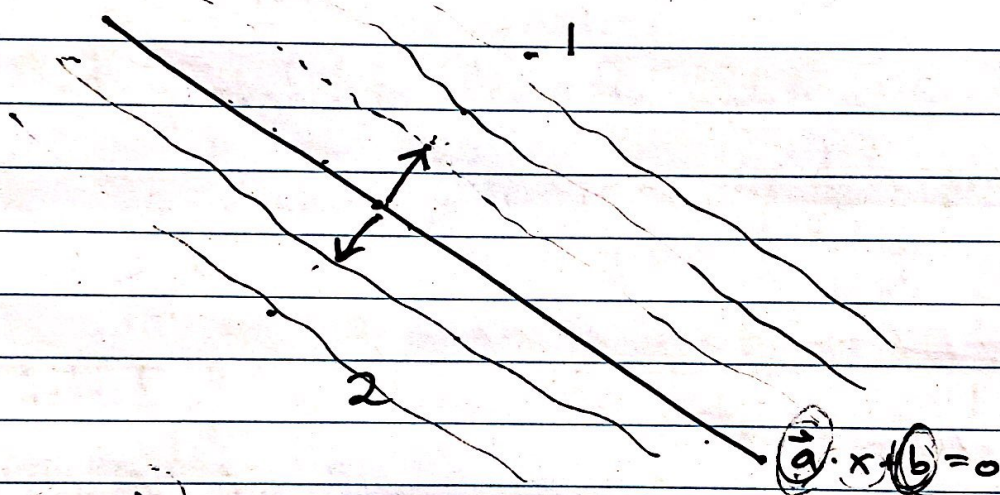
Output: Distribution $p(j) = \frac{e^{y_j}}{\sum_{i=1}^k e^{y_i}}$

"Take exponential and normalize"

Given input $x \in \mathbb{R}^n$, we return the distribution $\text{softmax}(Ax + b)$, A, b are parameters.

Ex: Two classes: Parameters are $\vec{a} \in \mathbb{R}^n, b \in \mathbb{R}$.

$$p(1) = \frac{e^{\vec{a} \cdot x + b}}{e^{\vec{a} \cdot x + b} + 1}, \quad p(2) = \frac{1}{e^{\vec{a} \cdot x + b} + 1}$$



$$\frac{p(1)}{p(2)} = e^{\vec{a} \cdot x + b}, \quad \log\left(\frac{p(1)}{p(2)}\right) \leftarrow \text{logarithm of the odds}$$

Assumption: $\log\left(\frac{p(1)}{p(2)}\right)$ is linear in the feature vector.

Learning the parameters $\vec{a}, b$ from data

Data: $\{(x_1, l_1), \dots, (x_n, l_n)\} = \mathcal{D}$

$\uparrow$ Feature vectors $\nwarrow$ labels

Given this data $\mathcal{D}$, how can we estimate A, b ?

Maximum Likelihood

Given x_1 , model gives softmax $(Ax_1 + b)$

$$\frac{1}{\sum_{i=1}^k e^{a_i \cdot x_1 + b_i}} \begin{pmatrix} e^{\vec{a}_1 \cdot x_1 + b_1} \\ \vdots \\ e^{\vec{a}_k \cdot x_1 + b_k} \end{pmatrix}$$

What probability does the model assign to l_1 ?

$$\frac{1}{\sum_{i=1}^k e^{a_i \cdot x_1 + b_i}} \cdot e^{\vec{a}_{l_1} \cdot x_1 + b_{l_1}} \quad \left(\begin{array}{l} \text{data points are} \\ \text{indep. so} \\ \text{take product} \end{array} \right)$$

Instead of considering this prob, we consider its negative logarithm:

$$\log \left(\sum_{i=1}^k e^{a_i \cdot x_1 + b_i} \right) - (\vec{a}_{l_1} \cdot x_1 + b_{l_1})$$

Logistic Regression loss:

$$L_D(A, b) = \sum_{(x, l) \in D} \left[\log \left(\sum_{i=1}^k e^{a_i \cdot x + b_i} \right) - (\vec{a}_l \cdot x + b_l) \right]$$

$$(A, b) = \min_{A, b} L_D(A, b)$$

Basic Statistical Learning Theory

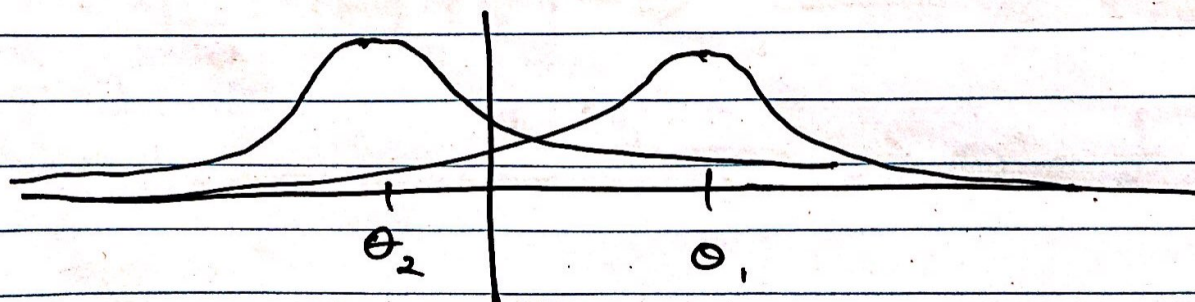
- Goal: Estimate an unknown probability dist. D on a set X from samples (i.i.d) $x_1, \dots, x_n \in X$.
- Introduce a family of distributions p_θ for $\theta \in \Theta$ and try to choose θ to "match" the samples.
- Maximum Likelihood Estimate: Choose θ to maximize the probability of the samples.

Ex: Let $X = \mathbb{R}$, have some samples $x_1, \dots, x_n$ drawn from a distribution D .

• Let p_θ be a Gaussian w/ variance 1, centered at $\theta \in \mathbb{R} = \Theta$, i.e.,

$$\text{density } p_\theta(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{(x-\theta)^2}{2}}$$

↑
centered at θ ↑
variance 1



- Use the samples to find the center θ .

Maximum Likelihood Estimation

• Given $\Theta \in \Theta = \mathbb{R}$, what is the probability of $\{x_j\}_{j=1}^n$?

Likelihood function (function of Θ)

- Samples independent: $p_{\Theta}(\{x_j\}_{j=1}^n) = \prod_{j=1}^n p_{\Theta}(x_j)$.

$$= \frac{1}{(2\pi)^{n/2}} \prod_{j=1}^n e^{-(x_j - \Theta)^2/2} = \frac{1}{(2\pi)^{n/2}} e^{-\sum_{j=1}^n (x_j - \Theta)^2/2}$$

- MLE: Choose Θ to maximize this!

- Often it's useful to consider $\log(p_{\Theta}(\{x_j\}_{j=1}^n))$.

log likelihood function

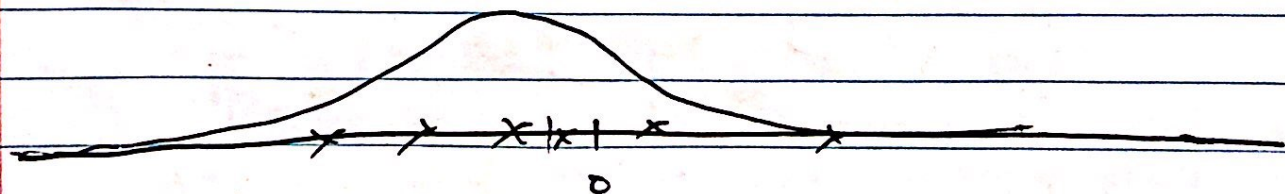
$$\Theta^* = \operatorname{argmax}_{\Theta \in \Theta} \log(p_{\Theta}(\{x_j\}_{j=1}^n))$$

$$\left(\operatorname{argmin}_{\Theta \in \Theta} -\log(p_{\Theta}(\{x_j\}_{j=1}^n)) \right)$$

• For this example: $\log(p_{\Theta}(\{x_j\}_{j=1}^n)) =$

$$= -\log(2\pi) \cdot \left(\frac{n}{2}\right) - \sum_{j=1}^n \frac{(x_j - \Theta)^2}{2}$$

$$\Theta^* = \operatorname{argmin}_{\Theta \in \mathbb{R}} \sum_{j=1}^n \frac{(x_j - \Theta)^2}{2} \Rightarrow \Theta^* = \frac{1}{n} \sum_{j=1}^n x_j$$



Classification / Logistic Regression

$$\begin{aligned}
 p(y|x, w, b) &= \text{softmax}(w x + b) = \\
 &= \frac{1}{\sum_{i=1}^k e^{w_i x + b_i}} \begin{pmatrix} e^{w_1 x + b_1} \\ \vdots \\ e^{w_k x + b_k} \end{pmatrix}
 \end{aligned}$$

$w \in \mathbb{R}^{k \times d}$, $x \in \mathbb{R}^d$, $b \in \mathbb{R}^k$
 $y_i = \begin{pmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ 0 \end{pmatrix}$ (label i)
 e_{y_i}

Need to calculate:

$$\begin{aligned}
 -\log(p(y_i|x_i, w, b)) &= -\log\left(\frac{1}{e^{w x_i + b} \cdot \mathbb{1}} e^{w x_i + b} \cdot y_i\right) \\
 &= \log(e^{w x_i + b} \cdot \mathbb{1}) - \log(e^{w x_i + b} \cdot y_i) \\
 (w, b)^* &= \underset{\substack{w, b \\ \in \mathbb{R}^{k \times d} \times \mathbb{R}^k}}{\text{argmin}} \sum_{i=1}^n \log(e^{w x_i + b} \cdot \mathbb{1}) - \log(e^{w x_i + b} \cdot y_i)
 \end{aligned}$$

Bayesian Approach to Machine Learning

- Goal: Estimate an unknown distribution on X from data $\{x_j\}_{j=1}^n$.

- Build a model:

- Set of parameters Θ
- Family of distributions on X ,
 $p(x|\theta)$
parameter $\in \Theta$

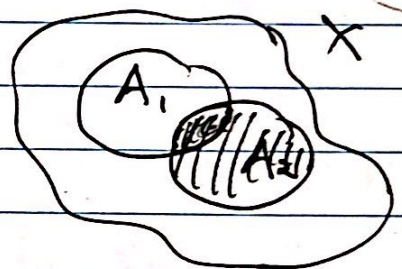
- Prior distribution on the parameters Θ .
 ~~$p(\theta)$~~ $q(\theta)$.

- Use Bayes' Law:

$$p(\theta|x)p(x) = p(x \text{ and } \theta) = p(x|\theta)p(\theta)$$

- Recall if A_1, A_2 events:

$$p(A_1|A_2) = \frac{p(A_1 \cap A_2)}{p(A_2)}$$



$$p(\theta|x) = \frac{p(x|\theta)q(\theta)}{p(x)}$$

$$\begin{array}{ccccc} p(\theta|x) & \sim & p(x|\theta) & q(\theta) \\ \uparrow & & \uparrow & & \uparrow \\ \text{posterior} & & \text{Likelihood} & & \text{prior} \\ \text{distribution} & & \text{function} & & \text{distribution} \end{array}$$

- Start with: prior $q(\theta)$
- Add data x , replace q with the posterior
 $p(\theta|x) \sim p(x|\theta)q(\theta)$

- More data $\rightarrow$ multiply by Likelihood + normalize.

- Left with a posterior distribution
- Sample from posterior dist to approximate

$$P_{\text{pred}}(x) = \int_{\Theta} p(x|\theta) p(\theta) d\theta$$

$\uparrow$
 posterior dist

- Choose θ to maximize posterior distribution:

$$\theta^* = \underset{\theta \in \Theta}{\operatorname{argmax}} p(\theta|x)$$

$$\theta^* = \underset{\theta \in \Theta}{\operatorname{argmin}} -\log(p(\theta|x)) =$$

$$= \underset{\theta \in \Theta}{\operatorname{argmin}} -\log(p(x|\theta)) - \log(q(\theta)) + \log(p(x))$$

$$= \underset{\theta \in \Theta}{\operatorname{argmin}} \underbrace{-\log(p(x|\theta))}_{\text{Negative log likelihood}} - \underbrace{\log(q(\theta))}_{\text{Regularization coming from prior}}$$

Ex: Image Classification / Logistic Regression

- Images $x \in X = \mathbb{R}^d$
- Labels $y \in Y = \{e_1, \dots, e_k\}$, $k = \# \text{ of labels}$

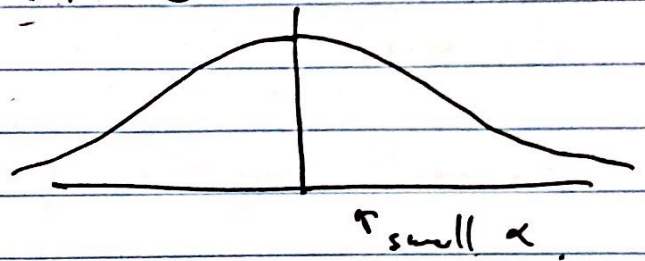
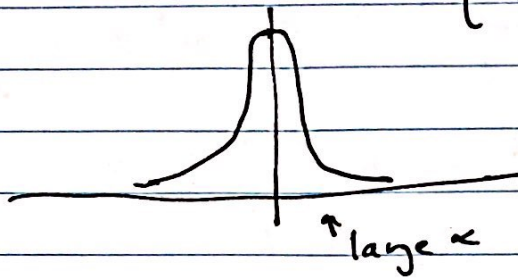
$$\begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix} \quad \begin{pmatrix} 0 \\ \vdots \\ 1 \end{pmatrix}$$

- Data: $\{(x_j, y_j)\}_{j=1}^n$

- Model: $\theta = (W, b) \in \mathbb{R}^{k \times d} \times \mathbb{R}^k$

$$p(y|x, \theta) = \frac{e^{Wx+b} \cdot y}{e^{Wx+b} \cdot \mathbf{1}}$$

• Prior distribution $q(w, b) = C e^{-\alpha(\|w\|_2^2 + \|b\|_2^2)}$



- Calculate the posterior:

$$p(w, b | x, y) = \frac{p(y | x, w, b) \overbrace{p(w, b)}^{q(w, b)}}{p(y | x)}$$

$$\rightarrow = \int_{(w)} p(y | x, w, b) q(w, b) d\theta$$

$$(w, b)^* = \underset{w, b}{\operatorname{argmin}} -\log(p(w, b | \{x_j, y_j\}_{j=1}^n)) =$$

$$= \underset{w, b}{\operatorname{argmin}} -\log(p(\{y_j\}_{j=1}^n | \{x_j\}_{j=1}^n, w, b)) - \log(q(w, b))$$

$$= \underset{w, b}{\operatorname{argmin}} -\log\left(\prod_{j=1}^n \underbrace{p(y_j | x_j, w, b)}_{\frac{e^{w x_j + b} \cdot y_j}{e^{w x_j + b} - 1}}\right) - \log\left(\underbrace{q(w, b)}_{C e^{-\alpha(\|w\|_2^2 + \|b\|_2^2)}}\right)$$

$$= \underset{w, b}{\operatorname{argmin}} \sum_{j=1}^n \log(e^{w x_j + b} - 1) - \log(e^{w x_j + b} \cdot y_j) + \alpha(\|w\|_2^2 + \|b\|_2^2)$$