

7

Logistic Regression

①

~~Def~~ A classification problem asks for a prediction $X \rightarrow \Delta(C)$
 X = input space $\subseteq \mathbb{R}^d$ prob. simplex on C .
 C = classes / labels

$$x \in X \quad x = (x_1, \dots, x_d)$$

pos's $\nwarrow$ $\nearrow$
 dim = features x_i = value @ i^{th} feature.

Def A probabilistic classifier for X, C is a fn

$$P: \Theta \times X \rightarrow \Delta(C) \quad (\theta, \vec{x}) \mapsto p_\theta(\cdot | \vec{x})$$

~~Def~~ Differentiable @ θ for any fixed $\vec{x}$.

Eg. spam filter: $C = \{\text{spam}, \text{not spam}\}$, X = emails

Multinomial logistic regression model.

(2)

$X = \mathbb{R}^d$ $\mathcal{C} = \text{classes / categories}$

$\Delta(\mathcal{C}) = \text{prob. distr.'s on } \mathcal{C}. \subseteq \mathbb{R}^{|\mathcal{C}|}$

$\text{relint } \Delta(\mathcal{C}) = \text{relative interior (in space it spans)}$

$\text{Softmax}: \mathbb{R}^{|\mathcal{C}|} \rightarrow \text{relint } \Delta(\mathcal{C})$

$$\vec{z} \mapsto \frac{e^{\vec{z}}}{\sum_{c=1}^{|\mathcal{C}|} e^{z_c}} \quad \begin{matrix} (\text{vector}) \\ (\#) \end{matrix}$$

vector $\rightarrow$ prob. distr.
w/ weight
on largest
entries

Def A softmax regression model

$$P: \Theta \times X \rightarrow \text{relint } \Delta(\mathcal{C}) \subseteq \mathbb{R}^{|\mathcal{C}|}$$

$$(\theta, \vec{x}) \mapsto \vec{p} = \text{softmax}(\underbrace{W^T \vec{x} + \beta}_{!!})$$

$$\vec{z}_\theta: X \rightarrow \mathbb{R}^{|\mathcal{C}|}$$

"logit"

weights $\downarrow$ biases $\swarrow$

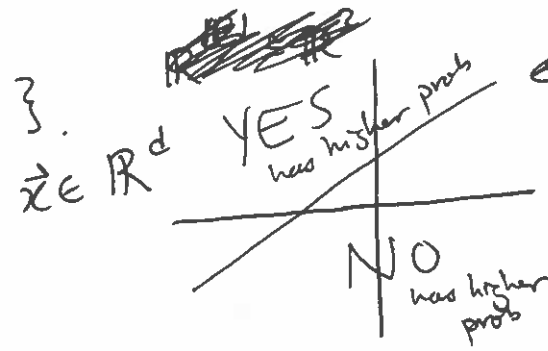
i.e. $\vec{\theta} = (W, \beta)$

parameters

$$\vec{\beta} \in \mathbb{R}^{|\mathcal{C}|}$$

$$W \text{ is } |\mathcal{C}| \times d$$

Eg. If $C = \{\text{yes, no}\}$.



decision boundary (where prob's equal) ③

$$\text{softmax}(W^T \vec{x} + \vec{\beta}) \in \mathbb{R}^2$$

has 1st coord $\geq$ 2nd

i.e. $W^T \vec{x} + \vec{\beta} = \begin{pmatrix} z_1 \\ z_2 \end{pmatrix}$

has slope 1

linear eq'n

Note: $\text{softmax}(1, 1, \dots, 1) = (\underbrace{\frac{1}{n}, \dots, \frac{1}{n}}_n)$

~~$\text{softmax}(\vec{z}) = \frac{e^{z_i}}{e^{z_1} + \dots + e^{z_n}}$~~

$\text{softmax}(\vec{z} + (1, 1, \dots, 1)) = \text{softmax}(\vec{z})$

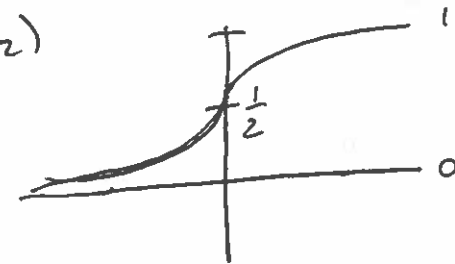
$$= e^{(z_1, z_2)} \left(\frac{1}{e^{z_1} + e^{z_2}} \right)$$

n=2 case

$$\text{softmax}(z_1, z_2) = (1 - \underbrace{\sigma(z_2 - z_1)}, \underbrace{\sigma(z_2 - z_1)})$$

$$= \sigma(z_1 - z_2)$$

with $\sigma(z) = \frac{e^z}{1 + e^z}$



$$\sigma: \mathbb{R} \rightarrow (0, 1)$$

inverse $\text{logit}(p) = \log \frac{p}{1-p}$ "log odds"

$(0, 1) \rightarrow \mathbb{R}$

Loss $l(\underbrace{w, \beta}_{\theta}; \underbrace{\vec{x}, y}_{\substack{\in \mathcal{X} \in \mathcal{C} \\ \text{training data}}}) = -\log p_y = -z_y + \log \sum_{j=1}^{|\mathcal{C}|} e^{z_j}$ (4)

$w/ \quad z = z_{\theta}(\vec{x})$

$\uparrow$
yth entry
 $\downarrow \vec{p}$

For many samples $(\vec{x}_i, y_i)$

$\uparrow$ observation $\nwarrow$ gold label

$$L(w, \beta) = \frac{1}{N} \sum_{i=1}^N l(w, \beta; \vec{x}_i, y_i)$$

We wish to minimize this.

Gradient Descent

$F: \mathbb{R}^k \rightarrow \mathbb{R}$ differentiable $\theta \mapsto F(\theta)$ eg. $(w, \beta) \mapsto L(w, \beta)$

gradient $\nabla F(\theta) := \begin{pmatrix} \frac{\partial F}{\partial \theta_1}(\theta) \\ \vdots \\ \frac{\partial F}{\partial \theta_k}(\theta) \end{pmatrix}$ direction of steepest descent @ θ :

$$-\frac{\nabla F(\theta)}{\|\nabla F(\theta)\|}$$

Gradient Descent Algorithm:

$$\theta_{\text{new}} = \theta_{\text{old}} - \underset{\substack{\uparrow \\ \text{learning rate}}}{\eta} \nabla F(\theta)$$

Gradient of loss

$$\nabla_z l = \vec{p} - \vec{e}_y$$

one hot vector!

$$\nabla_w l = \vec{x} (\vec{p} - \vec{e}_y)^T$$

$$\nabla_b l = \vec{p} - \vec{e}_y$$

Proposition

L is convex as a f" in (w, b)



(5)

Artificial Neurons ("units")

Def An $\nearrow$ is a $f^n: \mathbb{R}^d \rightarrow \mathbb{R}$ of shape $\vec{x} \mapsto \phi(\underbrace{\vec{w} + \vec{x} + b}_{\text{preactivation}})$ } output = activation

$\uparrow$ $\mathbb{R} \rightarrow \mathbb{R}$ activation f^n $\uparrow$ weights $\uparrow$ bias

Choice of activation = type of unit

eg. $H(t) = \begin{cases} 0 & t < 0 \\ 1 & t \geq 0 \end{cases}$ gives a perceptron

Def A Feedforward network or multilayer perceptron (MLP)

w/ layer widths $n_0, n_1, \dots, n_L$ is a f^n $f_\theta: \mathbb{R}^{n_0} \rightarrow \mathbb{R}^{n_L}$

$\nwarrow L = \text{depth}$

defined recursively by:

$$h_0 = \vec{x} \in \mathbb{R}^{n_0}$$

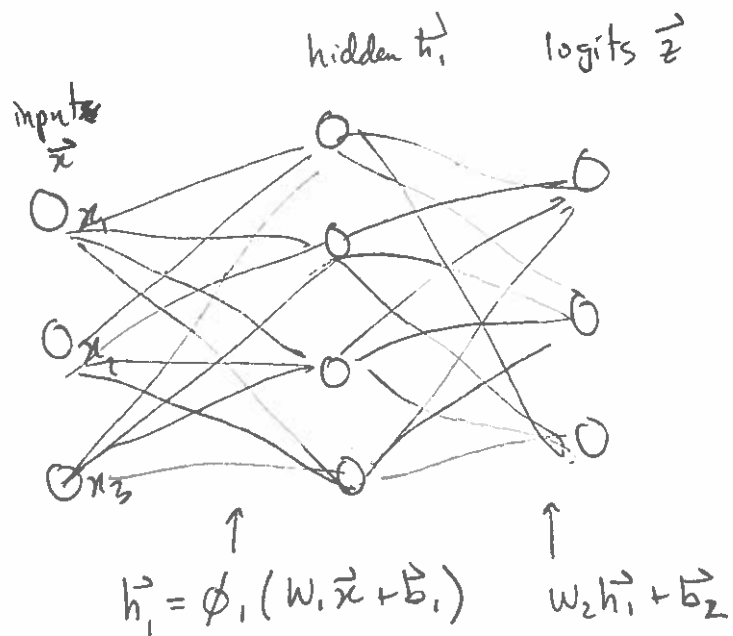
$\swarrow$ applied coordinatewise

$$a_\ell = W_\ell h_{\ell-1} + b_\ell \quad h_\ell = \phi_\ell(a_\ell) \quad 1 \leq \ell < L$$

"logits" $\rightarrow a_L = W_L h_{L-1} + b_L \quad f_\theta(x) = a_L$

with $W_\ell \in \mathbb{R}^{n_\ell \times n_{\ell-1}}, a_\ell, b_\ell \in \mathbb{R}^{n_\ell}, h_\ell \in \mathbb{R}^{n_\ell}$ (~~h~~)





other common activations ϕ :

tanh t

ReLU = $\max\{0, t\}$

others

DEMO: Handwriting rec.

DEMO: Playground.

7

DefFeedforward neural language model consists of:

① encoder $e_{\theta_1}: V^* \rightarrow \mathbb{R}^{n_0}$
 context $\mapsto$ feature vec

② feedf.net $f_{\theta_2}: \mathbb{R}^{n_0} \rightarrow \mathbb{R}^V$

giving $P(\theta, c) = P_{\theta}(\cdot | c) = \text{softmax}(f_{\theta_2}(e_{\theta_1}(c)))$

Example "Bengio-style" "neural n-gram"

encoder: $C \in \mathbb{R}^{V \times d}$ embedding matrix ~~sizes~~ $\left\{ \begin{array}{l} V \rightarrow \mathbb{R}^d \\ C \mapsto \{r_c \mid c \text{ row}\} \end{array} \right\}$

So tokens $a_1, \dots, a_{n-1} \mapsto \vec{x} = \{r_{a_1}, \dots, r_{a_{n-1}}\}$ concatenated into vector
 $\in \mathbb{R}^{(n-1)d}$

ff. net:

one hidden layer:

$$h_{ctx} = \tanh(W_1 \vec{x} + b_1) \quad z = W_2 h_{ctx} + b_2$$

$$\theta = C, W_1, b_1, W_2, b_2$$

$$V, n(n-1)d, h, V \times h, V$$