

Lecture 4 ML & Opt

Logistics:

- Participation 1 due today
- Participation 2 posted later today
- HW1 posted next Tuesday
- Out of town Feb 28 & Mar 3:
guest lectures or recorded lectures

Agenda

- MLE for the Categorical Model
- Practical ML procedure
- Regularization
- Risk decomposition: how well do we need to optimize?

MLE for Categorical Model

$$y|x \sim \text{Categorical}(p_1(x), \dots, p_K(x))$$

$$\text{where } p_i(x) = \frac{e^{\Theta_i^T x}}{Z_{\Theta}(x)} \quad \text{where } Z_{\Theta}(x) = \sum_{j=1}^K e^{\Theta_j^T x}$$

$$\text{recall } p_{\Theta}(x) = \begin{bmatrix} p_1(x) \\ \vdots \\ p_K(x) \end{bmatrix} = \frac{e^{\Theta x}}{1^T e^{\Theta x}} = \text{softmax}(\Theta x)$$

$$\text{where } \Theta = \begin{bmatrix} \Theta_1^T \\ \vdots \\ \Theta_K^T \end{bmatrix}$$

Given $\{(y_i, x_i)\}_{i=1}^n$, where $y_i \in \mathbb{R}^k$ one-hot encodes the class of the i th example (e.g. class 1 encodes as $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and class k as $\begin{bmatrix} 0 \\ \vdots \\ 1 \end{bmatrix}$), our goal is to learn the model parameters $\Theta \in \mathbb{R}^{k \times d}$ using MLE

$$\mathcal{L}(\Theta) = \frac{1}{n} \sum_{i=1}^n -\log p_{\Theta}(y_i | x_i)$$

Recall

$$\begin{aligned} p_{\Theta}(y_i = e_j | x_i) &= e_j^T p_{\Theta}(x_i) = y_i^T p_{\Theta}(x_i) \\ &= y_i^T \text{softmax}(\Theta x_i) \end{aligned}$$

So

$$\mathcal{L}(\Theta) = \frac{1}{n} \sum_{i=1}^n -\log p_{\Theta}(y_i | x_i)$$

$$= \frac{1}{n} \sum_{i=1}^n -\log [y_i^T \text{softmax}(\Theta x_i)]$$

$$= \frac{1}{n} \sum_{i=1}^n - y_i^T \log [\text{softmax}(\Theta x_i)]$$

apply log entrywise to
its vector argument

$$\uparrow ? = \frac{1}{n} \sum_{i=1}^n \ell(\text{softmax}(\Theta x_i), y_i)$$

what is the correct choice of ℓ , to make
the NLL have the form of an empirical risk

Guess: $\mathcal{L}$ is the cross-entropy loss
What is cross-entropy? Easiest to first define
KL-divergence

$$D_{KL}(p \parallel q) = \sum p_i \ln \frac{p_i}{q_i} \quad \text{if } p \text{ and } q \text{ are two probability vectors}$$
$$= \underbrace{\sum p_i \ln p_i}_{\text{Entropy of } p, H(p)} - \underbrace{\sum p_i \ln q_i}_{\text{cross-entropy of } p \text{ and } q, H(p, q)}$$

so minimizing $H(p, q)$ w.r.t. q minimizes
 $D_{KL}(p \parallel q)$ w.r.t. q

Recall the negative log-likelihood w.r.t. the i th sample

$$\begin{aligned} & -y_i^T \log(\text{softmax}(\Theta x_i)) \\ &= \sum_{j=1}^k (y_i)_j \ln(p_{\Theta}(x_i)_j) \\ &= H(y_i, p_{\Theta}(x_i)) \end{aligned}$$

Note that if the i th example is from class j , so $y_i = e_j$, then

$$\begin{aligned} -y_i^T \log(\text{softmax}(\Theta x_i)) &= -e_j^T \ln(p_{\Theta}(x_i)) \\ &= -\ln(p_{\Theta}(x_i)_j) \end{aligned}$$

is the log of the probability assigned to the true class

Upshot is that

$$\begin{aligned} \mathcal{L}(\theta) &= \frac{1}{n} \sum_{i=1}^n H(y_i, \text{softmax}(\theta x_i)) \\ &= \frac{1}{n} \sum_{i=1}^n -\ln[p_{\theta}(x_i)_{y_i}] \end{aligned}$$

So minimizing the NLL learns model parameters $\hat{\theta}$ that maximize the probabilities assigned by the model to the true classes of the observation, and the loss function for this ERM is

$$\mathcal{L}(u, v) = H(u, v)$$

The optimization problem is

$$\hat{\Theta} = \underset{\Theta}{\operatorname{argmin}} \mathcal{L}(\Theta)$$

$$= \underset{\Theta}{\operatorname{argmin}} \frac{1}{n} \sum_{i=1}^n -y_i^T \ln(\operatorname{softmax}(\Theta x_i))$$

$$= \underset{\Theta}{\operatorname{argmin}} \frac{1}{n} \sum_{i=1}^n -y_i^T \ln\left(\frac{e^{\Theta x_i}}{z_{\Theta}(x_i)}\right)$$

$$= \underset{\Theta}{\operatorname{argmin}} \frac{1}{n} \sum_{i=1}^n -y_i^T [\ln(e^{\Theta x_i}) - (\ln z_{\Theta}(x_i)) \mathbf{1}]$$

$$= \underset{\Theta}{\operatorname{argmin}} \frac{1}{n} \sum_{i=1}^n [-y_i^T \Theta x_i + y_i^T \mathbf{1} \ln z_{\Theta}(x_i)]$$

$$= \underset{\Theta}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n \left[-y_i^T \Theta x_i + \ln z_{\Theta}(x_i) \right]$$

$$= \underset{\Theta}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n -y_i^T \Theta x_i + \frac{1}{n} \sum_{i=1}^n \ln z_{\Theta}(x_i)$$

$$= \underset{\Theta}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n -y_i^T \Theta x_i + \frac{1}{n} \sum_{i=1}^n \ln \left(\sum_{j=1}^K e^{\Theta_j^T x_i} \right)$$

This is the optimization problem we solve in practice to learn Θ using MLE for multiclass classification.

Practical Considerations

train vs test vs validation splits



use training data to solve our ERM, that is

$$\hat{f} = \underset{f \in \mathcal{F}}{\operatorname{argmin}} \frac{1}{n_{\text{train}}} \sum_{i=1}^{n_{\text{train}}} \ell(f(x_i), y_i)$$

||
 $R_{n_{\text{train}}}(f)$

→ use fresh samples to estimate the risk

$$R(\hat{f}) = \mathbb{E}_{\mathcal{D}} \ell(\hat{f}(x), y)$$

$$\approx \frac{1}{n_{\text{test}}} \sum_{i=1}^{n_{\text{test}}} \ell(f(x_i), y_i)$$

We use the test set to estimate the "generalization error"

$$R(\hat{f}) - R_{n_{\text{train}}}(\hat{f}) \approx R_{n_{\text{test}}}(\hat{f}) - R_{n_{\text{train}}}(\hat{f})$$

What if we want to choose between different model classes, e.g. $\mathcal{F}_1$ and $\mathcal{F}_2$?

- use training set to find best (empirically) models $\hat{f}_1 \in \mathcal{F}_1$ and $\hat{f}_2 \in \mathcal{F}_2$
- use validation dataset to select best of $\hat{f}_1$ and $\hat{f}_2$ on a fresh independent set of samples
- estimate the generalization gap of this selected model on a fresh independent test set

General flow for supervised ML

- Fix our problem domain, risk (loss) measure, collect appropriate data $\{(x_i, y_i)\}_{i=1}^n$
- Split our data into train/validation/test (60/20/20)
- Determine ahead of time a set models and hyperparameters to try
- Fit each of these using the same training data
- Measure their performance using the validation data and select the best, f_{opt}
- Report an estimate of the gen gap / pop risk of f_{opt} using the test data