

Multi-class linear prediction

COMS 4771 Fall 2025

Multi-class prediction

Multi-class prediction usually means classification problems with $|\mathcal{Y}| \geq 3$ (... but technically also includes $|\mathcal{Y}| = 2$)

- ▶ Assume we use zero-one loss, so corresponding risk is error rate
- ▶ Best prediction of Y given $X = x$ is

$$f^*(x) = \arg \max_{k \in \mathcal{Y}} \Pr(Y = k \mid X = x)$$

1 / 9

	$ \mathcal{Y} = 2$	$ \mathcal{Y} \geq 3$
nearest neighbors	✓	✓
decision trees	✓	✓
generative models	✓	✓
logistic regression	✓	?
Perceptron	✓	?

Linear classifiers are inherently binary output functions;
need some work to extend to $|\mathcal{Y}| \geq 3$

2 / 9

Multi-class logistic regression

Generalize logistic regression model to K classes, $[K] = \{1, 2, \dots, K\}$

$$\Pr(Y = k \mid X = x) = \frac{\exp(x^\top w^{(k)})}{\sum_{c \in [K]} \exp(x^\top w^{(c)})}$$

► K weight vectors $w^{(1)}, \dots, w^{(K)} \in \mathbb{R}^d$ are parameters of the model

$$w^{(k)} = (w_1^{(k)}, \dots, w_d^{(k)})$$

Total of Kd parameters

In this model with parameters $w^{(1)}, \dots, w^{(K)}$, classifier with smallest error rate is

$$f^*(x) = \arg \max_{k \in [K]} \Pr(Y = k \mid X = x) = \arg \max_{k \in [K]} x^\top w^{(k)}$$

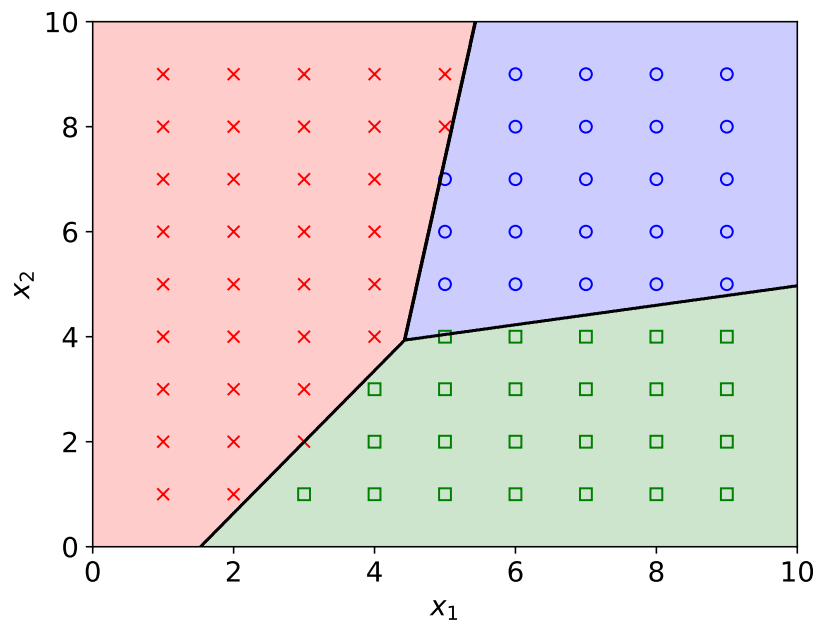
Maximum likelihood estimation: given training data $\mathcal{S}$ from $\mathbb{R}^d \times [K]$, log-likelihood of $w^{(1)}, \dots, w^{(K)}$ is

$$\begin{aligned} \ln L(w^{(1)}, \dots, w^{(K)}) &= \sum_{(x,y) \in \mathcal{S}} \ln \left(\frac{\exp(x^\top w^{(y)})}{\sum_{c \in [K]} \exp(x^\top w^{(c)})} \right) \\ &= \sum_{(x,y) \in \mathcal{S}} \left(x^\top w^{(y)} - \ln \left(\sum_{c \in [K]} \exp(x^\top w^{(c)}) \right) \right) \end{aligned}$$

► Negative log-likelihood turns out to be a convex objective function

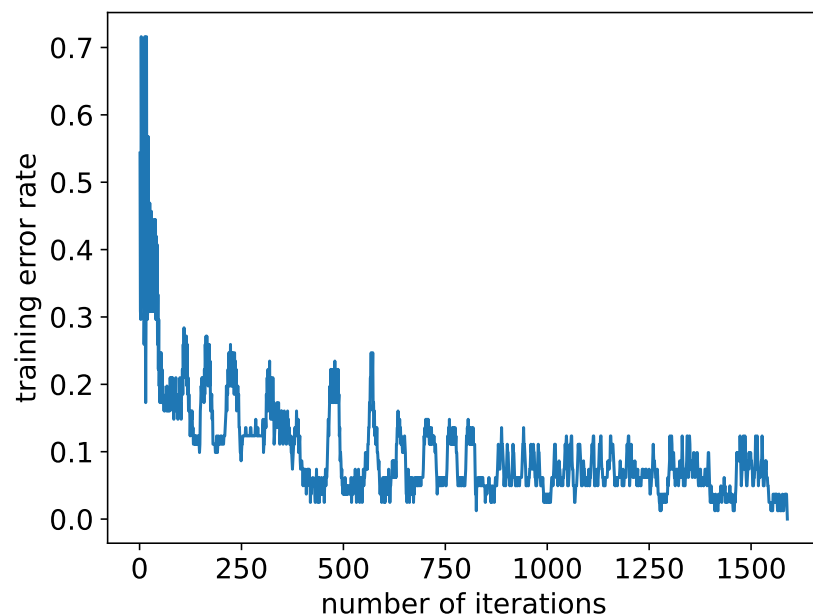
4 / 9

Synthetic example: “linearly separable” dataset



5 / 9

Synthetic example: “linearly separable” dataset



6 / 9

Can also interpret negative log-likelihood as sum of log losses, like in binary case, for prediction function

$$\hat{p}_W(x) = \text{softmax}(Wx)$$

where $\text{softmax}: \mathbb{R}^K \rightarrow \mathbb{R}^K$ ([soft\(arg\)max function](#)) is defined by

$$\text{softmax}(u)_k = \frac{\exp(u_k)}{\sum_{c \in [K]} \exp(u_c)}$$

- $W \in \mathbb{R}^{K \times d}$ is matrix of parameters, one row per class
- Negative log-likelihood is

$$J(W) = \sum_{(x,y) \in \mathcal{S}} -\ln(\hat{p}_W(x)_y)$$

7 / 9

Setup

```
import torch

W = torch.zeros(d, K)
W.requires_grad_(True)

def f(x):
    return torch.softmax(x.matmul(W),
        ↪ dim=1)

loss = torch.nn.NLLLoss(reduction='sum')
```

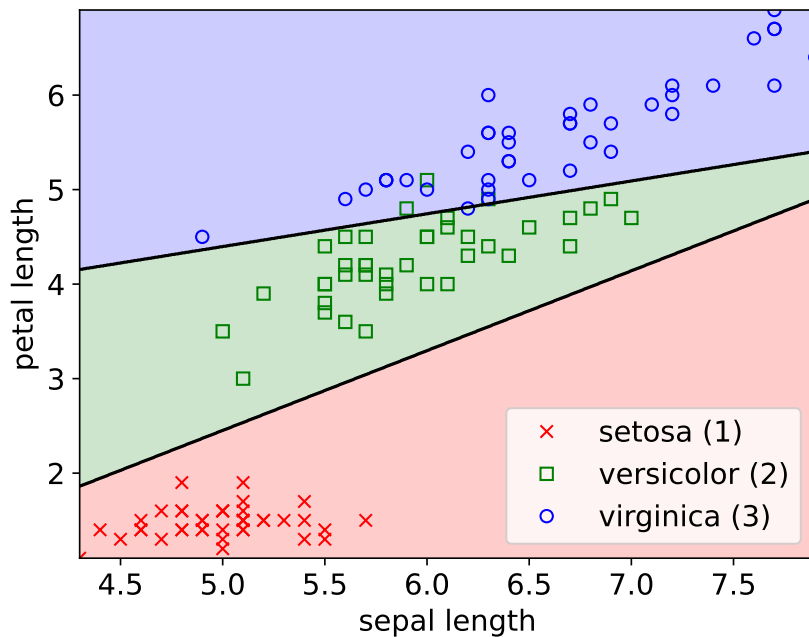
Gradient descent code

```
for t in range(num_iter):
    J = loss(f(x), y)
    J.backward()
    with torch.no_grad():
        W -= eta * W.grad
        W.grad.zero_()
```

8/9

Example: iris dataset

- ▶ Multi-class logistic regression ($\eta = 0.01$, $T = 2^{17}$ iterations)
- ▶ Test error rate: 3.33%



9/9