

# Introduction to Machine Learning

**Lecture 11:**

***Logistic Regression (Classification Algorithm)***

**Prof. Subrahmanya Swamy Peruru**  
Department of Electrical Engineering  
IIT Kanpur

# Probabilistic View of Classification

Assume there is an underlying data distribution  $(X, Y) \sim \mathcal{D}$ .

**Bayes Optimal Classifier (0–1 Loss):**

$$f^* = \operatorname{argmin}_f \mathbb{E}_{(X,Y) \sim \mathcal{D}} [\mathbf{1}\{Y \neq f(X)\}]$$

Upon solving, the optimal classifier can be obtained as:

$$f^*(\mathbf{x}) = \operatorname{arg} \max_{c \in \{1, \dots, C\}} \mathbb{P}(Y = c \mid X = \mathbf{x}).$$

**Example:**

Suppose we have 3 classes and **if the data distribution is given to us:**

For a given input  $\mathbf{x}$ , we compare:

$$P(Y = 1 \mid X = \mathbf{x}) = 0.7, \quad P(Y = 2 \mid X = \mathbf{x}) = 0.2, \quad P(Y = 3 \mid X = \mathbf{x}) = 0.1$$

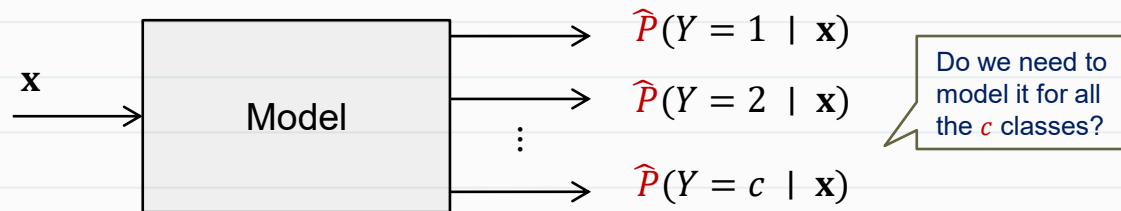
Since  $P(Y = 1 \mid X = \mathbf{x}) = 0.7$ , is the largest, the Bayes optimal prediction is

$$f^*(\mathbf{x}) = 1$$

To find  $f^*$ , we need to know the true data distribution  $\mathcal{D}$ , which is not known to us.

# Learning $P(Y | X)$ from Data

We learn an approximation model  $\hat{P}(Y | X = \mathbf{x})$  from data



Now consider binary classification:  $Y \in \{0,1\}$

It is sufficient to model  $\hat{P}(1 | \mathbf{x})$ .

Since  $\hat{P}(0 | \mathbf{x})$  can be automatically obtained using  $\hat{P}(0 | \mathbf{x}) = 1 - \hat{P}(1 | \mathbf{x})$

## Key question:

- Can we use a linear model for  $\hat{P}(1 | \mathbf{x})$ ?
- Can  $\mathbf{w}^T \mathbf{x} + b$  represent a probability?

# Why Linear Output is Not Enough

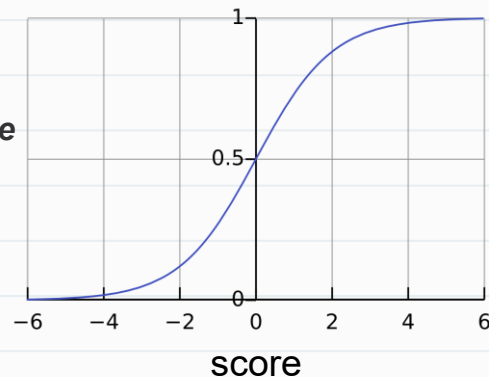
The linear score  $\mathbf{w}^T \mathbf{x} + b$  cannot directly represent probabilities:

- Can be negative
- Can exceed 1

**Idea:** Instead of interpreting  $\mathbf{w}^T \mathbf{x} + b$  as a probability, treat it as a **confidence score**

- Large positive score  $\Rightarrow$  class 1 likely
- Large negative score  $\Rightarrow$  class 0 likely
- Zero score  $\Rightarrow$  not sure of the class

Sigmoid function



We need a function (**sigmoid**) that **converts scores into probabilities**.

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

- Maps  $\mathbb{R} \rightarrow (0,1)$
- $\sigma(0) = \frac{1}{2}$
- $z \rightarrow +\infty \Rightarrow \sigma(z) \rightarrow 1$
- $z \rightarrow -\infty \Rightarrow \sigma(z) \rightarrow 0$

# Logistic Regression Model

Score:

$$\text{score} = \mathbf{w}^T \mathbf{x} + b$$

Probability model:

$$P(Y = 1 | \mathbf{x}; \mathbf{w}, b) = \sigma(\mathbf{w}^T \mathbf{x} + b)$$

Standard Prediction rule:

$$\hat{y} = \begin{cases} 1 & \text{if } P(Y = 1 | \mathbf{x}, \mathbf{w}, b) > \frac{1}{2} \\ 0 & \text{otherwise} \end{cases}$$

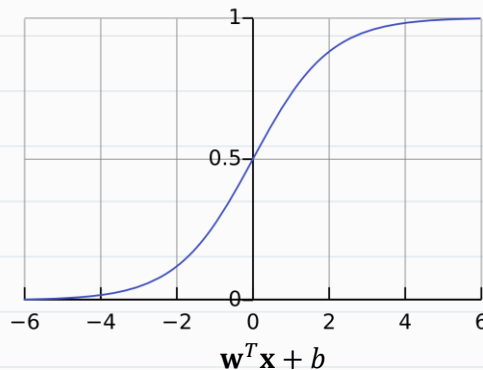
**Equivalent form:** Predict as class 1 if

$$\mathbf{w}^T \mathbf{x} + b > 0$$

Model structure fixed and Decision rule defined.

**Remaining task:** Estimate  $(\mathbf{w}, b)$  from data

$$P(Y = 1 | \mathbf{x}; \mathbf{w}, b) = \sigma(\mathbf{w}^T \mathbf{x} + b)$$



# Maximum Likelihood: Warm-up

**Bernoulli random variable:**

$$P(Y = 1) = \theta, \quad P(Y = 0) = 1 - \theta$$

**Compact form:**

$$P(Y = y) = \theta^y (1 - \theta)^{1-y}$$

**We have data:** The outcome of  $n$  independent coin tosses

$$\{y_1, y_2, \dots, y_n\}$$

**Goal is to estimate the parameter:** Bias of the coin  $\theta$

**Likelihood function of  $\theta$ :** If  $\theta$  is the parameter of the coin, how likely is it to observe that data?

$$L(\theta) = \prod_{i=1}^n P(Y = y_i \mid \theta)$$

**MLE estimate:**

$$\theta_{\text{ML}} = \underset{\theta}{\operatorname{argmax}} L(\theta)$$

Coin Toss



H = 1, T = 0

# Negative Log-Likelihood

## Likelihood

$$L(\theta) = \prod_{i=1}^n P(y_i | \theta)$$

## Log-Likelihood

- Products are hard to optimize; sums are easier
- Log converts products into sums  $\ell(\theta) = \log L(\theta) = \sum_{i=1}^n \log P(y_i | \theta)$
- Logarithm is a strictly increasing (monotonic) function  $\Rightarrow \operatorname{argmax}_{\theta} L(\theta) = \operatorname{argmax}_{\theta} \log L(\theta)$

## From maximization to minimization

Optimization algorithms are usually written as minimization problems, so we define:

$$\text{NLL}(\theta) = -\log L(\theta)$$

## Minimizing Negative log-likelihood

$$\operatorname{argmin}_{\theta} \text{NLL}(\theta)$$

# Example

**Observed data:** 0, 1, 1, 0

**Likelihood:**

$$\begin{aligned} L(\theta) &= \prod_{i=1}^n P(y_i | \theta) \\ &= P(0 | \theta) \times P(1 | \theta) \times P(1 | \theta) \times P(0 | \theta) \\ &= (1 - \theta) \times \theta \times \theta \times (1 - \theta) \\ &= \theta^2(1 - \theta)^2 \end{aligned}$$

**Negative Log-likelihood** (Convex function):

$$\begin{aligned} \text{NNL}(\theta) &= -\log L(\theta) \\ &= -2 \log \theta - 2 \log(1 - \theta) \end{aligned}$$

**Setting the gradient to zero gives:**

$$\begin{aligned} \frac{2}{\theta} - \frac{2}{1 - \theta} &= 0 \\ \Rightarrow \theta_{\text{ML}} &= \frac{1}{2} \end{aligned}$$

# Likelihood for logistic regression

Probability modelled by sigmoid function:

$$P(y = 1 \mid \mathbf{x}; \mathbf{w}, b) = \sigma(\mathbf{w}^T \mathbf{x} + b)$$

$$P(y = 0 \mid \mathbf{x}; \mathbf{w}, b) = 1 - \sigma(\mathbf{w}^T \mathbf{x} + b)$$

Assuming i.i.d. given data  $(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_n, y_n)$

The likelihood function of the parameters unknown parameters  $(\mathbf{w}, b)$ :

$$L(\mathbf{w}, b) = \prod_{i=1}^n P(y_i \mid \mathbf{x}_i; \mathbf{w}, b)$$

Let  $\hat{p}_i = P(y_i = 1 \mid \mathbf{x}_i; \mathbf{w}, b)$  be the probability that  $i^{th}$  data point is predicted as class 1

**Negative log-likelihood (NLL):**

$$\text{NLL}(\mathbf{w}, b) = -\log L(\mathbf{w}, b) = -\sum_{i=1}^n y_i \log \hat{p}_i + (1 - y_i) \log(1 - \hat{p}_i)$$

**Minimizing the NLL loss (solved using Gradient Descent):**

$$\min_{\mathbf{w}, b} \text{NLL}(\mathbf{w}, b)$$

# Log-likelihood and Cross-entropy

Minimizing Negative log-likelihood is equivalent to minimizing cross-entropy loss.

**Negative log-likelihood (NLL):**

$$y_i \log \hat{p}_i + (1 - y_i) \log(1 - \hat{p}_i)$$

**Cross Entropy:**

$$H(p, q) = - \sum_x p(x) \log q(x)$$

# Gradient Descent

**Goal:** minimize a differentiable function

$$\min_{\mathbf{x} \in \mathbb{R}^d} f(\mathbf{x}).$$

**Key intuition (1-D):**  $f(x): \mathbb{R} \rightarrow \mathbb{R}$

- The derivative  $f'(x)$  is the *slope* of  $f$  at  $x$ .
- If  $f'(x) < 0$  (downward slope), moving *right* decreases  $f$ .
- If  $f'(x) > 0$  (upward slope), moving *left* decreases  $f$ .

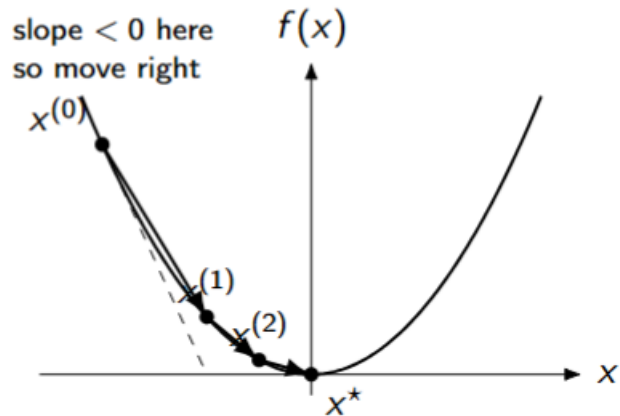
**Example:**  $f(x) = x^2 \Rightarrow f'(x) = 2x$

- If  $x < 0$ , then  $f'(x) = 2x < 0 \Rightarrow$  step to the right.
- If  $x > 0$ , then  $f'(x) = 2x > 0 \Rightarrow$  step to the left.
- Minimum at  $x^* = 0$ .

Update rule (1-D):  $x^{(t+1)} = x^{(t)} - \eta f'(x^{(t)})$ ,  $\eta_t > 0$  (step size / learning rate).

**Gradient descent (vector form):**  $f(\mathbf{x}): \mathbb{R}^d \rightarrow \mathbb{R}$

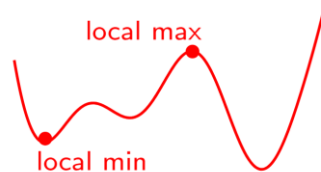
$$\mathbf{x}^{(t+1)} = \mathbf{x}^{(t)} - \eta_t \nabla f(\mathbf{x}^{(t)}).$$



# Gradient Descent (Step size and Initialization)

Two knobs you choose:

- **Initial point**  $\mathbf{x}^{(0)}$  (where you start) influences convergence point (multiple local minima in non-convex functions)
- **Step size**  $\eta_t$ :
  - Too large  $\Rightarrow$  may overshoot / diverge / oscillate.
  - Too small  $\Rightarrow$  very slow progress.
  - Often: small constant  $\eta$ , or decreasing  $\eta_t$  (e.g.,  $\eta_t \propto 1/t$ )



**Descent idea (local):** for small enough  $\eta_t$ ,

$$f(\mathbf{x}^{(t+1)}) \leq f(\mathbf{x}^{(t)}) \quad (\text{typically holds for smooth functions with a suitable step size}).$$

**Stochastic Gradient Descent (SGD)** (more details in future lectures):

$$\mathbf{x}^{(t+1)} = \mathbf{x}^{(t)} - \eta_t \widehat{\nabla} f(\mathbf{x}^{(t)}),$$

where  $\widehat{\nabla} f$  is a noisy/mini-batch estimate of the true gradient (cheaper per iteration; adds jitter that can help exploration).

Thank You