

Indian Institute of Information Technology, Allahabad



Introduction to Machine Learning

Logistic Regression and Linear Classification

By

Dr. Shiv Ram Dubey

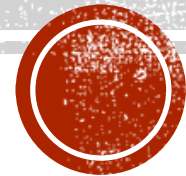
Associate Professor

Computer Vision and Biometrics Lab

Department of Information Technology

Indian Institute of Information Technology, Allahabad

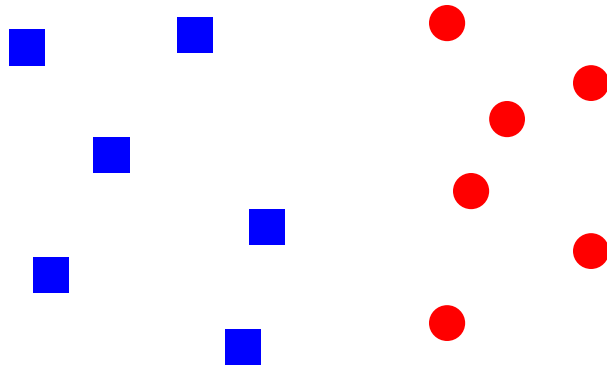
Web: <https://profile.iita.ac.in/srdubey/>



DISCLAIMER

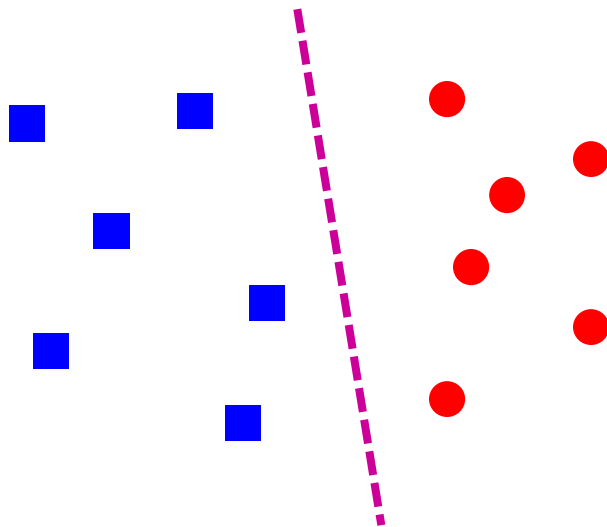
The content (text, image, and graphics) used in this slide are adopted from many sources for academic purposes. Broadly, the sources have been given due credit appropriately. However, there is a chance of missing out some original primary sources. The authors of this material do not claim any copyright of such material.

Linear Classifiers – 2 Class Problem



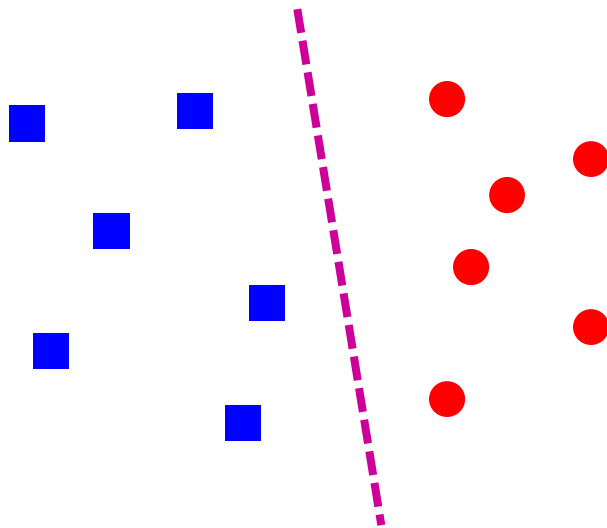
- Find a *linear function* to separate the classes:

Linear Classifiers – 2 Class Problem



- Find a *linear function* to separate the classes:

Linear Classifiers – 2 Class Problem



- Find a *linear function* to separate the classes:

$$f(\mathbf{x}) = \text{sign}(\mathbf{w} \cdot \mathbf{x} + b)$$

Loss Minimization View of ML

- **Three design decisions**

- **Model family:** What are the candidate models f ? (E.g., linear functions)
- **Loss function:** How to define “approximating”? (E.g., MSE loss)
- **Optimizer:** How do we optimize the loss? (E.g., gradient descent)

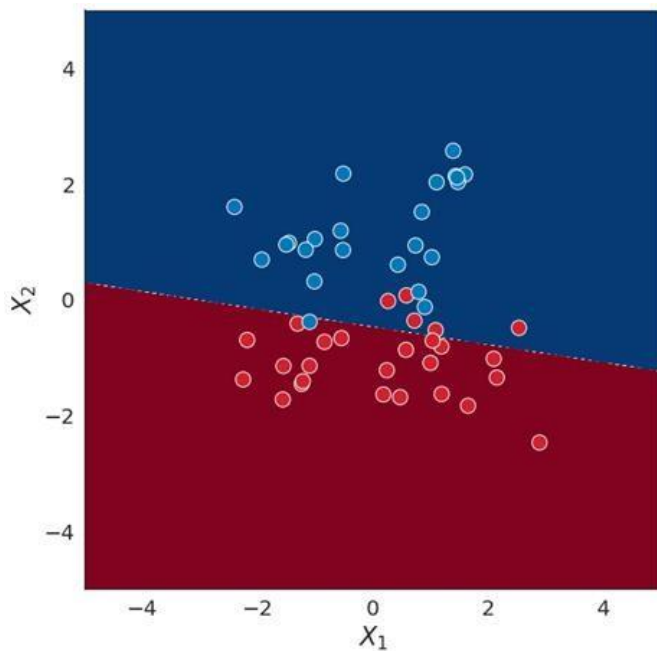
- How do we adapt to classification?

Linear Functions for (Binary) Classification

- **Input:** Dataset $Z = \{ (x_1, y_1), (x_2, y_2), \dots, (x_n, y_n) \}$

- **Classification:**

- Labels $y_i \in \{0, 1\}$
- Predict $y_i \approx 1(\mathbf{w}^T \mathbf{x}_i \geq 0)$
- $1(C)$ equals 1 if C is true and 0 if C is false
- How to learn $\mathbf{w}$? **Need a loss function!**

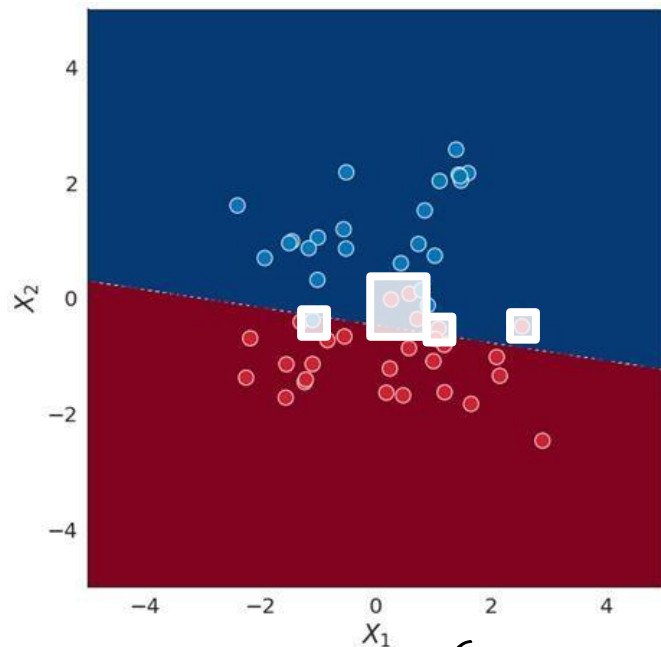


Loss Functions for Linear Classifiers

- (In)accuracy:

$$L(\mathbf{w}; \mathbf{Z}) = \frac{1}{n} \sum_{i=1}^n 1(y_i \neq f_{\mathbf{w}}(x_i))$$

- Computationally intractable
- Often, but not always the “true” loss (e.g., imbalanced data)



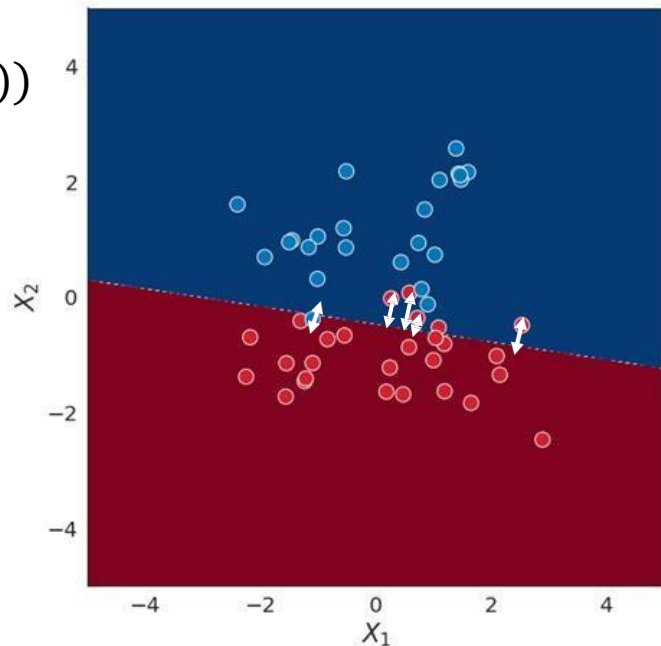
$$L(\mathbf{w}; \mathbf{Z}) = \frac{6}{50}$$

Loss Functions for Linear Classifiers

- **Distance:**

$$L(\mathbf{w}; \mathbf{Z}) = \frac{1}{n} \sum_{i=1}^n \text{dist}(\mathbf{x}_i, f_{\mathbf{w}}) \cdot 1(y_i \neq f_{\mathbf{w}}(\mathbf{x}_i))$$

- If $L(\mathbf{w}; \mathbf{Z}) = 0$, then 100% accuracy
- Variant of this loss results in SVM
- We consider a more general strategy



$$L(\mathbf{w}; \mathbf{Z}) = 1.2$$

What is Logistic Regression?

- Logistic regression is a statistical **classification** model that predicts the **probability** of a **binary** outcome (e.g., yes/no) by applying a **logistic (sigmoid) function** to a **linear** combination of input features.
- ...a frequent interview question

Maximum Likelihood View of ML

- **Two design decisions**

- **Likelihood:** Probability $p_{\mathbf{w}}(\mathbf{y} \mid \mathbf{x})$ of data $(\mathbf{x}, \mathbf{y})$ given parameters $\mathbf{w}$
- **Optimizer:** How do we optimize the Negative Log-Likelihood (NLL)? (E.g., gradient descent)

- **Corresponding Loss Minimization View:**

- **Model family:** Most likely label $f_{\mathbf{w}}(\mathbf{x}) = \arg \max_{\mathbf{y}} p_{\mathbf{w}}(\mathbf{y} \mid \mathbf{x})$
- **Loss function:** Negative log likelihood (NLL) $\ell(\mathbf{w}; \mathbf{Z}) = -\sum_{i=1}^n \log p_{\mathbf{w}}(\mathbf{y}_i \mid \mathbf{x}_i)$

- Very powerful framework for designing cutting edge ML algorithms
 - Write down the “right” likelihood, form tractable approximation if needed

Recall: Linear Regression under MLE View

- What about the model family?

$$\begin{aligned} f_w(x) &= \arg \max_y p_w(y | x) \\ &= \arg \max_y \frac{1}{\sqrt{2\pi}} \cdot e^{-\frac{(w^T x - y)^2}{2}} \\ &= w^T x \end{aligned}$$

- Recovers linear functions!

What About Classification?

Compare to linear regression:

$$p_w(y | x_i) \propto e^{-\frac{(w^T x_i - y)^2}{2}}$$

- Consider the following choice:

$$p_w(Y = 0 | x_i) \propto e^{-\frac{w^T x_i}{2}} \text{ and } p_w(Y = 1 | x_i) \propto e^{\frac{w^T x_i}{2}}$$

- Then, we have

$$p_w(Y = 1 | x_i) = \text{??????}$$

What About Classification?

Compare to linear regression:

$$p_w(y | x_i) \propto e^{-\frac{(w^T x_i - y)^2}{2}}$$

- Consider the following choice:

$$p_w(Y = 0 | x_i) \propto e^{-\frac{w^T x_i}{2}} \text{ and } p_w(Y = 1 | x_i) \propto e^{\frac{w^T x_i}{2}}$$

- Then, we have

$$p_w(Y = 1 | x_i) = \frac{e^{\frac{w^T x_i}{2}}}{e^{\frac{w^T x_i}{2}} + e^{-\frac{w^T x_i}{2}}} = \frac{1}{1 + e^{-w^T x_i}}$$

Sigmoid function

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

What About Classification?

Compare to linear regression:

$$p_w(y | x_i) \propto e^{-\frac{(w^T x_i - y)^2}{2}}$$

- Consider the following choice:

$$p_w(Y = 0 | x_i) \propto e^{-\frac{w^T x_i}{2}} \text{ and } p_w(Y = 1 | x_i) \propto e^{\frac{w^T x_i}{2}}$$

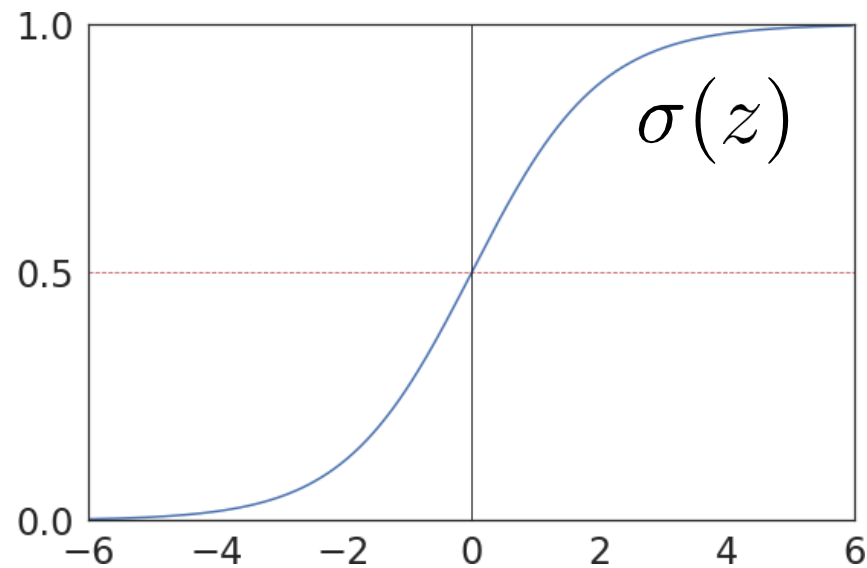
- Then, we have

$$p_w(Y = 1 | x_i) = \frac{e^{\frac{w^T x_i}{2}}}{e^{\frac{w^T x_i}{2}} + e^{-\frac{w^T x_i}{2}}} = \frac{1}{1 + e^{-w^T x_i}}$$

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

- Furthermore, $p_w(Y = 0 | x_i) = 1 - \sigma(w^T x_i)$

Logistic/Sigmoid Function



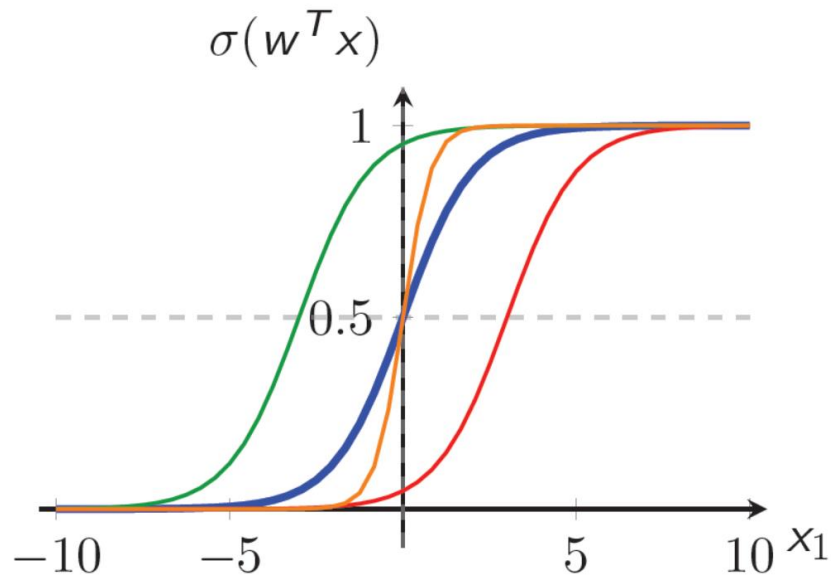
$$p_w(Y = 1 \mid x_i) = \sigma(w^T x_i)$$

Logistic/Sigmoid Function

Simple example:

$w = [w_0, w_1]$ and $x = [1, x_1]$

- w_0 controls the offset
- w_1 controls the steepness



— $w_0 = 0, w_1 = 1$	— $w_0 = -3, w_1 = 1$
— $w_0 = 3, w_1 = 1$	— $w_0 = 0, w_1 = 2.5$

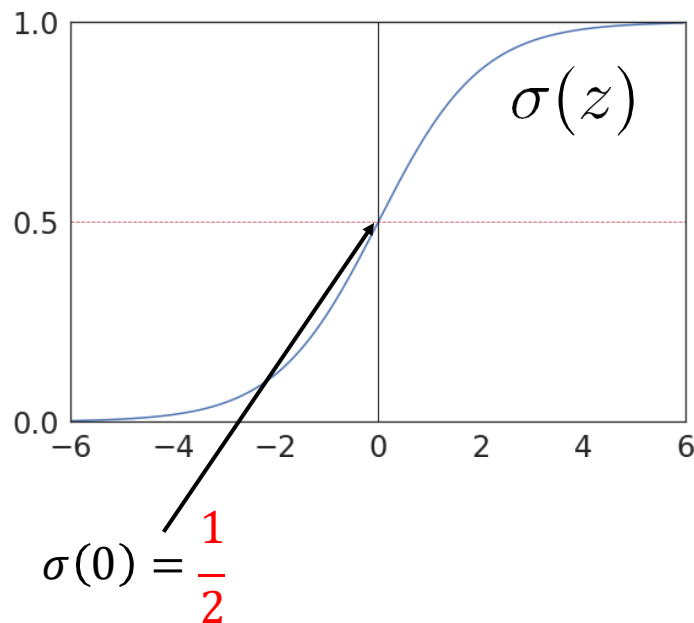
Logistic Regression Model Family

$$\begin{aligned} f_{\mathbf{w}}(\mathbf{x}) &= \arg \max_{\mathbf{y}} p_{\mathbf{w}}(\mathbf{y} \mid \mathbf{x}) \\ &= \arg \max_{\mathbf{y}} \begin{cases} \sigma(\mathbf{w}^T \mathbf{x}) & \text{if } \mathbf{y} = 1 \\ 1 - \sigma(\mathbf{w}^T \mathbf{x}) & \text{if } \mathbf{y} = 0 \end{cases} \\ &= \begin{cases} 1 & \text{if } \sigma(\mathbf{w}^T \mathbf{x}) \geq \frac{1}{2} \\ 0 & \text{otherwise} \end{cases} \end{aligned}$$

Logistic Regression Model Family

$$\begin{aligned} f_{\mathbf{w}}(\mathbf{x}) &= \arg \max_y p_{\mathbf{w}}(y | \mathbf{x}) \\ &= \arg \max_y \begin{cases} \sigma(\mathbf{w}^T \mathbf{x}) & \text{if } y = 1 \\ 1 - \sigma(\mathbf{w}^T \mathbf{x}) & \text{if } y = 0 \end{cases} \\ &= \begin{cases} 1 & \text{if } \sigma(\mathbf{w}^T \mathbf{x}) \geq \frac{1}{2} \\ 0 & \text{otherwise} \end{cases} \\ &= \begin{cases} 1 & \text{if } \mathbf{w}^T \mathbf{x} \geq 0 \\ 0 & \text{otherwise} \end{cases} \\ &= 1(\mathbf{w}^T \mathbf{x} \geq 0) \end{aligned}$$

- Recovers linear classifiers!



Logistic Regression Algorithm

- Then, we have the following Negative Log-Likelihood (NLL) loss:

$$\begin{aligned}\ell(\mathbf{w}; Z) &= -\sum_{i=1}^n \log p_{\mathbf{w}}(y_i | x_i) \\ &= -\sum_{i=1}^n 1(y_i = 1) \cdot \log(\sigma(\mathbf{w}^T x_i)) + 1(y_i = 0) \cdot \log(1 - \sigma(\mathbf{w}^T x_i)) \\ &= -\sum_{i=1}^n y_i \cdot \log(\sigma(\mathbf{w}^T x_i)) + (1 - y_i) \cdot \log(1 - \sigma(\mathbf{w}^T x_i))\end{aligned}$$

Binary Cross-Entropy (BCE) Loss

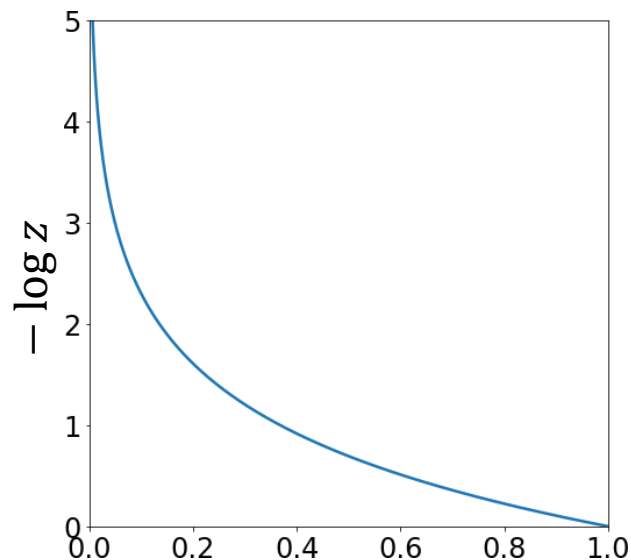
- Logistic regression minimizes this loss:

$$\hat{\mathbf{w}}(Z) = \arg \min_{\mathbf{w}} \ell(\mathbf{w}; Z)$$

Intuition on the Objective

- Loss for example i is

$$\begin{cases} -\log(\sigma(\mathbf{w}^\top \mathbf{x}_i)) & \text{if } y_i = 1 \\ -\log(1 - \sigma(\mathbf{w}^\top \mathbf{x}_i)) & \text{if } y_i = 0 \end{cases}$$



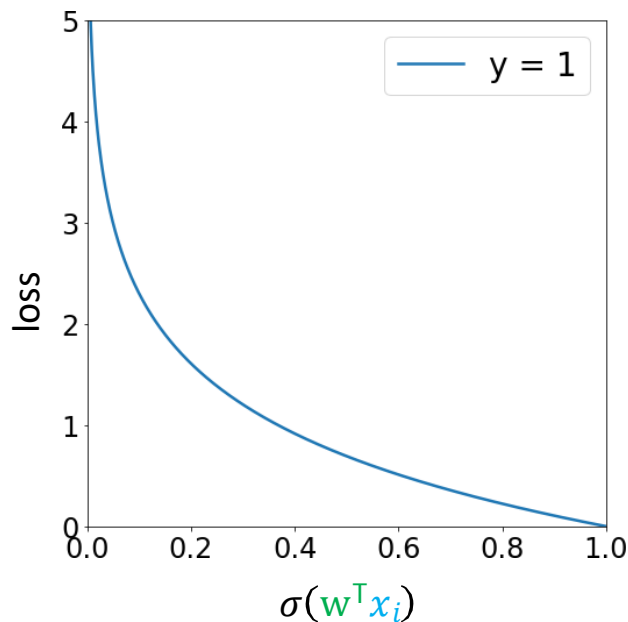
Intuition on the Objective

- Loss for example i is

$$\begin{cases} -\log(\sigma(\mathbf{w}^\top \mathbf{x}_i)) & \text{if } y_i = 1 \\ -\log(1 - \sigma(\mathbf{w}^\top \mathbf{x}_i)) & \text{if } y_i = 0 \end{cases}$$

- If $y_i = 1$:

- If $\sigma(\mathbf{w}^\top \mathbf{x}_i) = 1$, then loss = 0
- As $\sigma(\mathbf{w}^\top \mathbf{x}_i) \rightarrow 0$, loss $\rightarrow \infty$



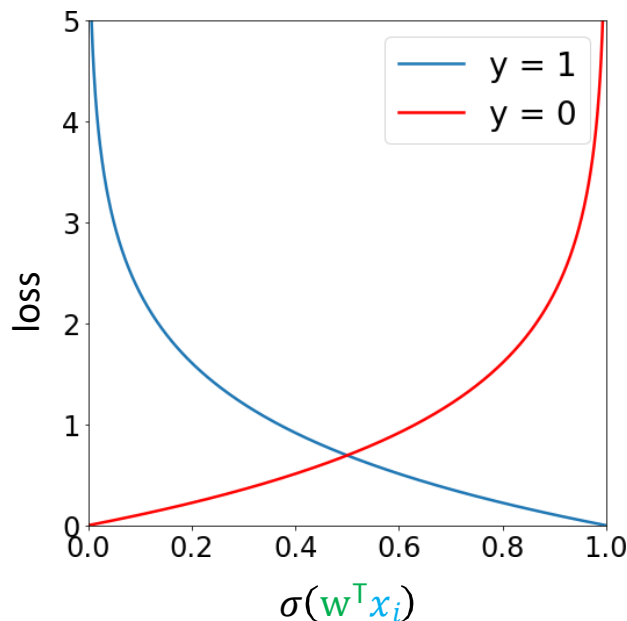
$$-y_i \cdot \boxed{\log(\sigma(\mathbf{w}^\top \mathbf{x}_i))} - (1 - y_i) \cdot \log(1 - \sigma(\mathbf{w}^\top \mathbf{x}_i))$$

Intuition on the Objective

- Loss for example i is

$$\begin{cases} -\log(\sigma(\mathbf{w}^\top \mathbf{x}_i)) & \text{if } y_i = 1 \\ -\log(1 - \sigma(\mathbf{w}^\top \mathbf{x}_i)) & \text{if } y_i = 0 \end{cases}$$

- If $y_i = 1$:
 - If $\sigma(\mathbf{w}^\top \mathbf{x}_i) = 1$, then loss = 0
 - As $\sigma(\mathbf{w}^\top \mathbf{x}_i) \rightarrow 0$, loss $\rightarrow \infty$
- If $y_i = 0$:
 - If $\sigma(\mathbf{w}^\top \mathbf{x}_i) = 0$, then loss = 0
 - As $\sigma(\mathbf{w}^\top \mathbf{x}_i) \rightarrow 1$, loss $\rightarrow \infty$



$$-y_i \cdot \log(\sigma(\mathbf{w}^\top \mathbf{x}_i)) - (1 - y_i) \cdot \log(1 - \sigma(\mathbf{w}^\top \mathbf{x}_i))$$

Optimization for Logistic Regression

- To optimize the loss, we need its gradient:

$$\begin{aligned}\nabla_{\mathbf{w}} \ell(\mathbf{w}; \mathbf{Z}) &= - \sum_{i=1}^n y_i \cdot \nabla_{\mathbf{w}} \log(\sigma(\mathbf{w}^\top \mathbf{x}_i)) + (1 - y_i) \cdot \nabla_{\mathbf{w}} \log(1 - \sigma(\mathbf{w}^\top \mathbf{x}_i)) \\ &= - \sum_{i=1}^n y_i \cdot \frac{\nabla_{\mathbf{w}} \sigma(\mathbf{w}^\top \mathbf{x}_i)}{\sigma(\mathbf{w}^\top \mathbf{x}_i)} - (1 - y_i) \cdot \frac{\nabla_{\mathbf{w}} \sigma(\mathbf{w}^\top \mathbf{x}_i)}{1 - \sigma(\mathbf{w}^\top \mathbf{x}_i)} \\ \sigma'(z) &= \sigma(z)(1 - \sigma(z)) \xrightarrow{\hspace{1cm}} \\ &= - \sum_{i=1}^n y_i \cdot \frac{\sigma(\mathbf{w}^\top \mathbf{x}_i)(1 - \sigma(\mathbf{w}^\top \mathbf{x}_i)) \cdot \mathbf{x}_i}{\sigma(\mathbf{w}^\top \mathbf{x}_i)} - (1 - y_i) \cdot \frac{\sigma(\mathbf{w}^\top \mathbf{x}_i)(1 - \sigma(\mathbf{w}^\top \mathbf{x}_i)) \cdot \mathbf{x}_i}{1 - \sigma(\mathbf{w}^\top \mathbf{x}_i)} \\ &= - \sum_{i=1}^n y_i \cdot (1 - \sigma(\mathbf{w}^\top \mathbf{x}_i)) \cdot \mathbf{x}_i - (1 - y_i) \cdot \sigma(\mathbf{w}^\top \mathbf{x}_i) \cdot \mathbf{x}_i \\ &= - \sum_{i=1}^n (y_i - \sigma(\mathbf{w}^\top \mathbf{x}_i)) \cdot \mathbf{x}_i\end{aligned}$$

Optimization for Logistic Regression

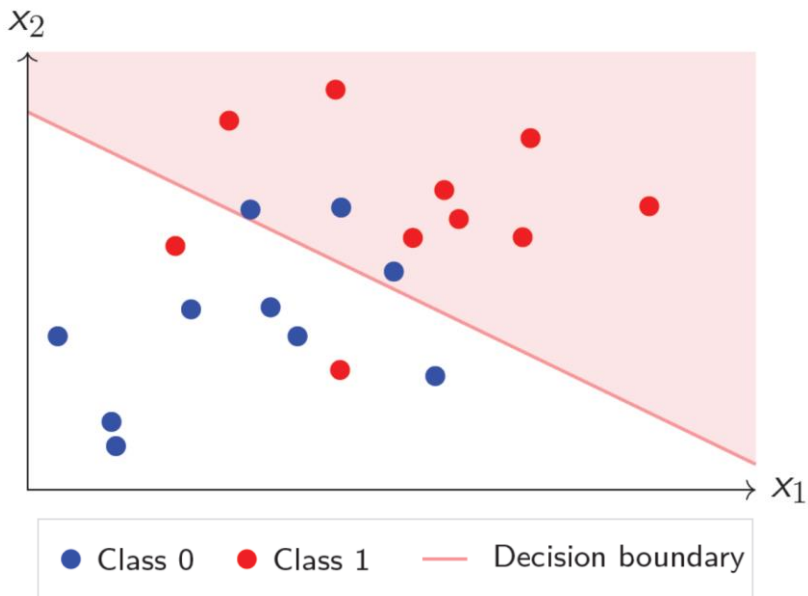
- Gradient of loss:

$$\nabla_{\mathbf{w}} \ell(\mathbf{w}; \mathbf{Z}) = - \sum_{i=1}^n (\sigma(\mathbf{w}^T \mathbf{x}_i) - y_i) \cdot \mathbf{x}_i$$

- Surprisingly similar to the gradient for linear regression!
 - Only difference is the σ
- Gradient descent works as before
 - No closed-form solution for $\hat{\mathbf{w}}(\mathbf{Z})$

Logistic Regression is a Linear Classifier

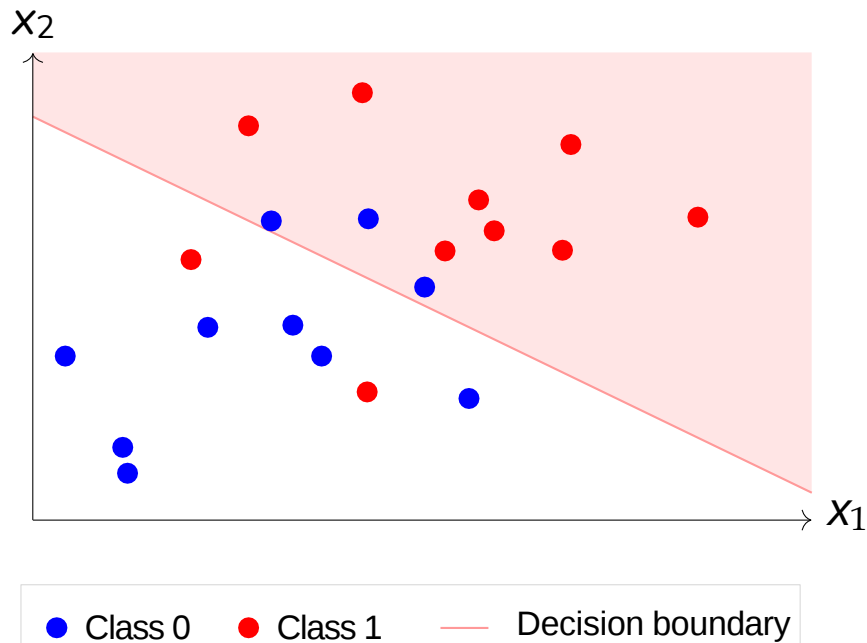
- A note on naming: it's called logistic *regression* because it's regressing the continuous value $P(Y|X) \in [0,1]$.
- The *classification* part only happens when we apply the threshold, i.e. predict Class 1 if $P_w(Y|X) \geq 0.5$.



Linear Features: Create Simple Decision Boundary

- Simple regression models
→ smooth predictors
- Simple classifier models
→ smooth decision boundaries

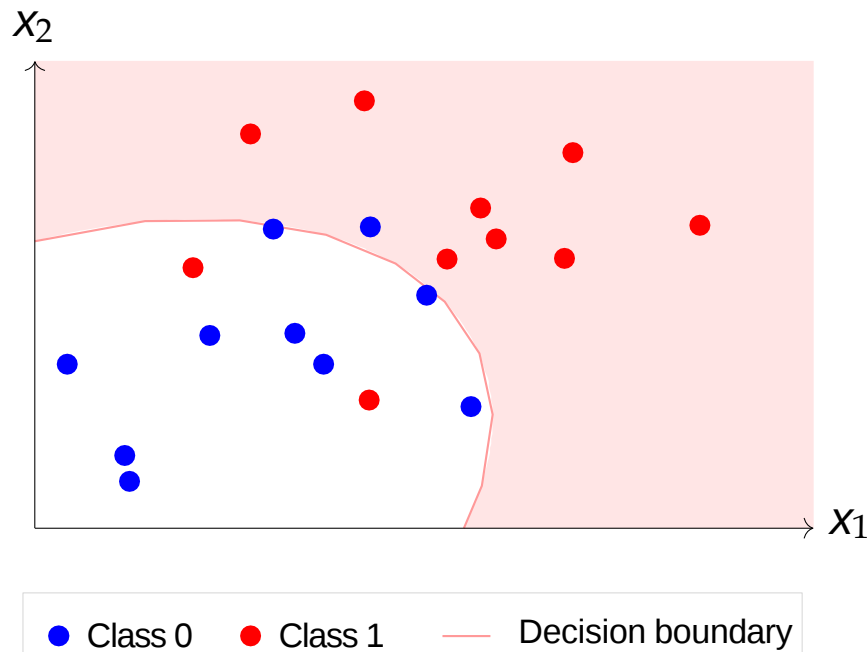
Feature	Value	Coeff w
$h_0(x)$	1	0.23
$h_1(x)$	x_1	1.12
$h_2(x)$	x_2	-1.07



Quadratic Features: Creates More Complex Decision Boundary

- Adding more features gives more complex models
- Decision boundary becomes more complex

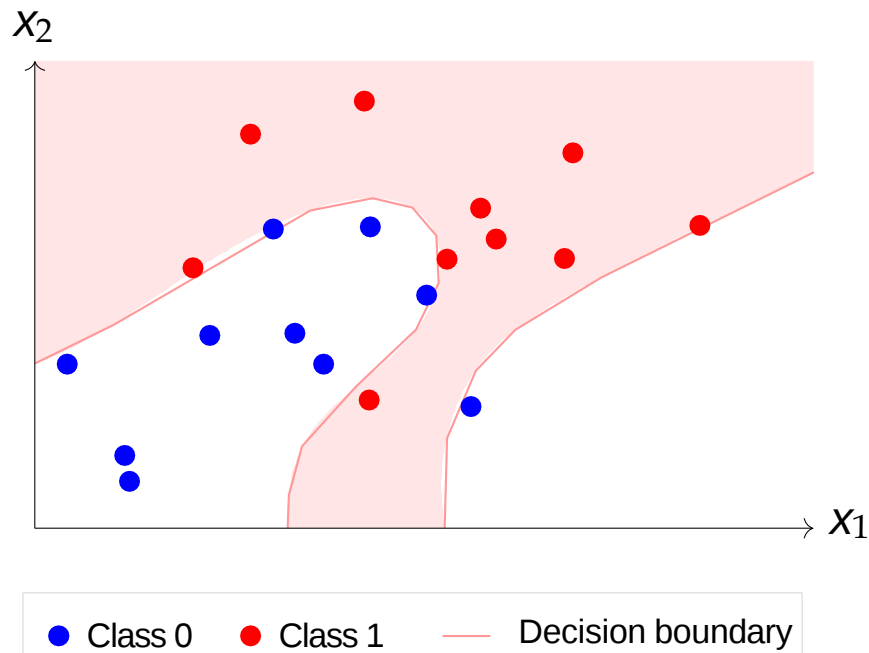
Feature	Value	Coeff
$h_0(x)$	1	1.68
$h_1(x)$	x_1	1.39
$h_2(x)$	x_2	-0.59
$h_3(x)$	x_1^2	-0.17
$h_4(x)$	x_2^2	-0.96
$h_5(x)$	x_1x_2	0.46



Higher-degree Polynomial Features: Can Lead to Overfitting

- Too much model complexity can lead to overfitting

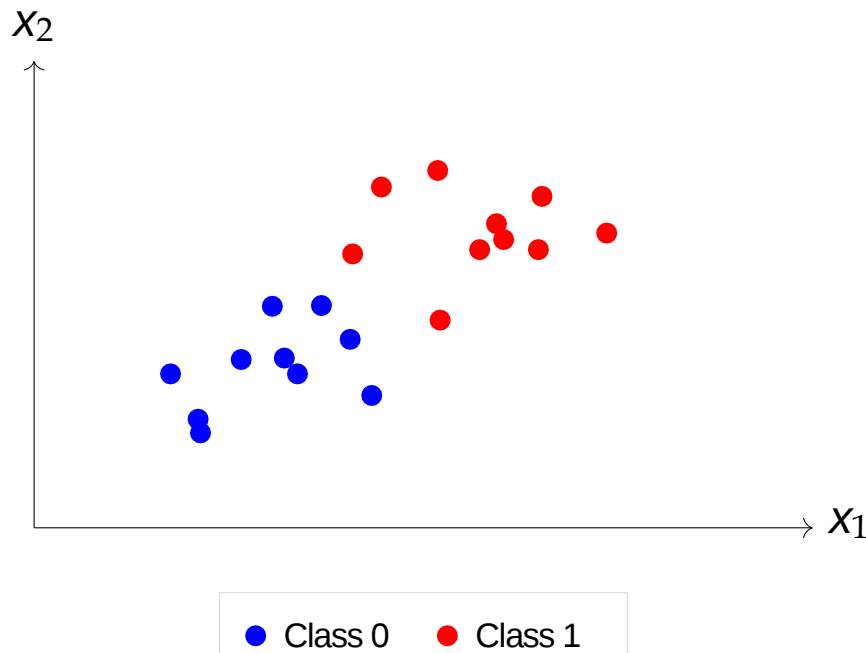
Feature	Value	Coeff
$h_0(x)$	1	21.6
$h_1(x)$	x_1	5.3
$h_2(x)$	x_2	-42.7
$h_3(x)$	x_1^2	-15.9
$h_4(x)$	x_2^2	-48.6
$\vdots$	$\vdots$	$\vdots$
$h_{k-1}(x)$	x_1^6	0.8
$h_k(x)$	x_2^6	-8.6



Even Linear Features Can Overfit

When would this badly overfit?

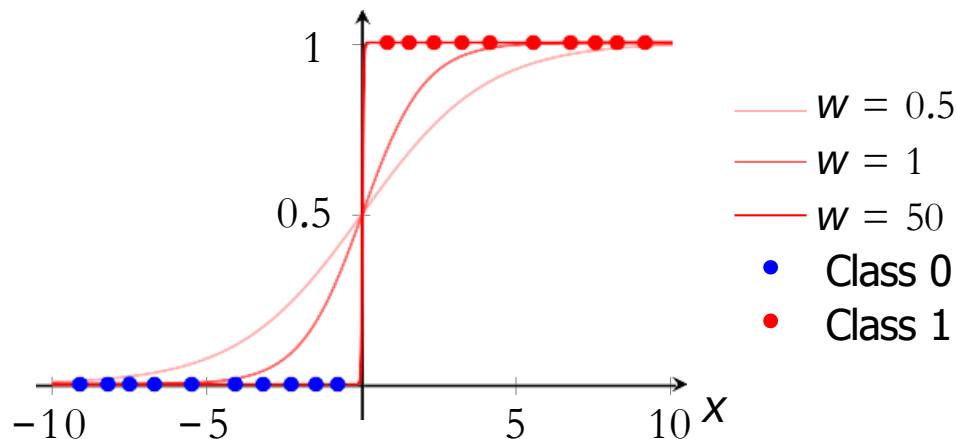
- When is the loss small?
- More likely with non-linear features. Why?



Linearly Separable Data will Overfit

When data is linearly separable, $\|w\| \rightarrow \infty$

$$P(Y = 1) = \sigma(wx)$$



Regularized Logistic Regression

We can add L_1 or L_2 regularization to the loss, e.g.:

$$\ell(\mathbf{w}; \mathbf{Z}) = - \sum_{i=1}^n y_i \cdot \log(\sigma(\mathbf{w}^\top \mathbf{x}_i)) + (1 - y_i) \cdot \log(1 - \sigma(\mathbf{w}^\top \mathbf{x}_i)) + \lambda \cdot \|\mathbf{w}\|_2^2$$

- Regularization ensures that $\mathbf{w}$ does not become too large
 - Prevents overconfidence
- Regularization can also be **necessary**
 - Without regularization (i.e., $\lambda = 0$) and data is linearly separable, then gradient descent diverges (i.e., $\mathbf{w} \rightarrow \pm\infty$)

Bias Term

- Most of the time, we want to add a bias term. Otherwise, the decision boundary is forced to pass through the origin (**not true in general**)

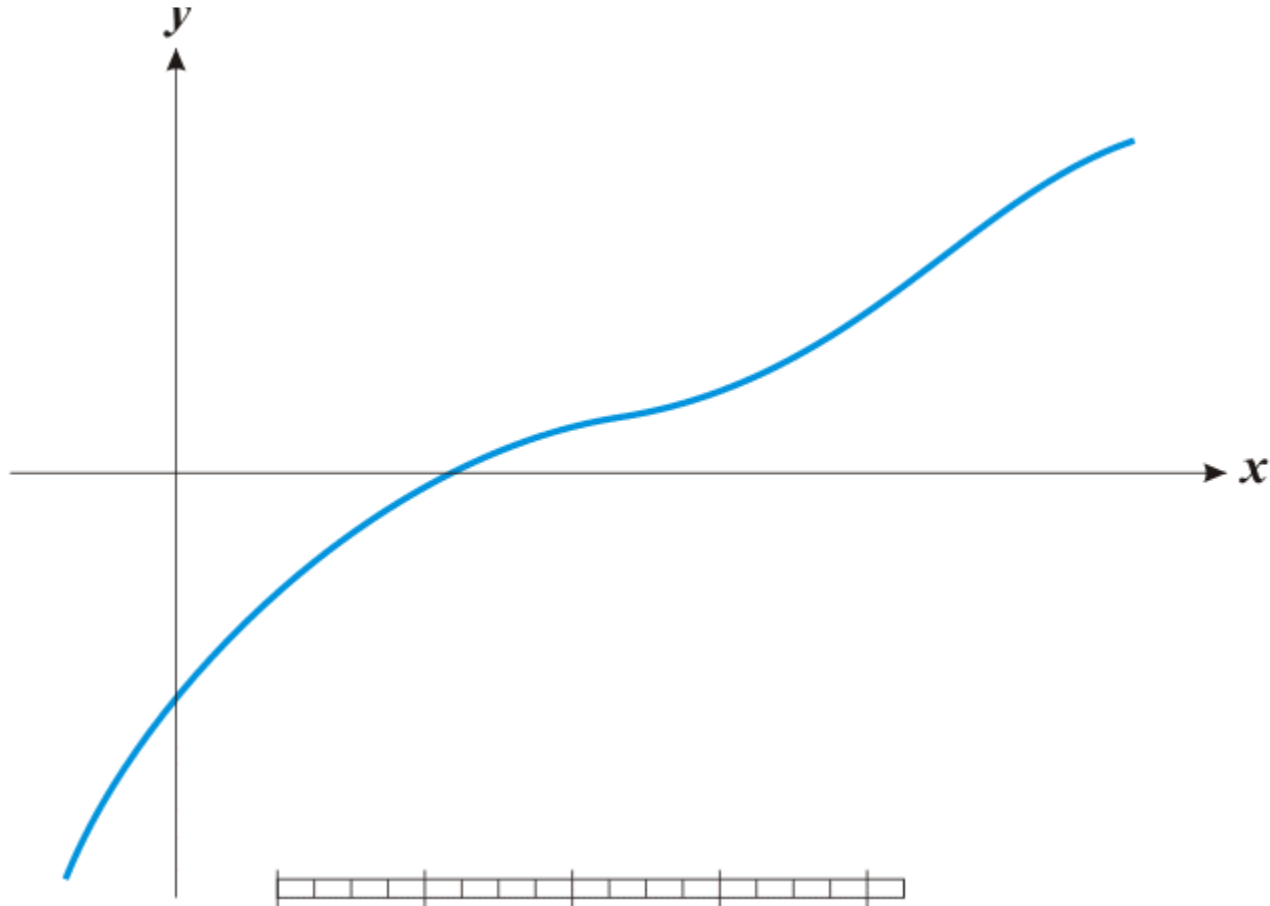
$$\ell(\mathbf{w}, \mathbf{b}; \mathbf{Z}) = - \sum_{i=1}^n y_i \cdot \log(\sigma(\mathbf{w}^\top \mathbf{x}_i + \mathbf{b})) + (1 - y_i) \cdot \log(1 - \sigma(\mathbf{w}^\top \mathbf{x}_i + \mathbf{b})) + \lambda \cdot \|\mathbf{w}\|_2^2$$

Newton's Method

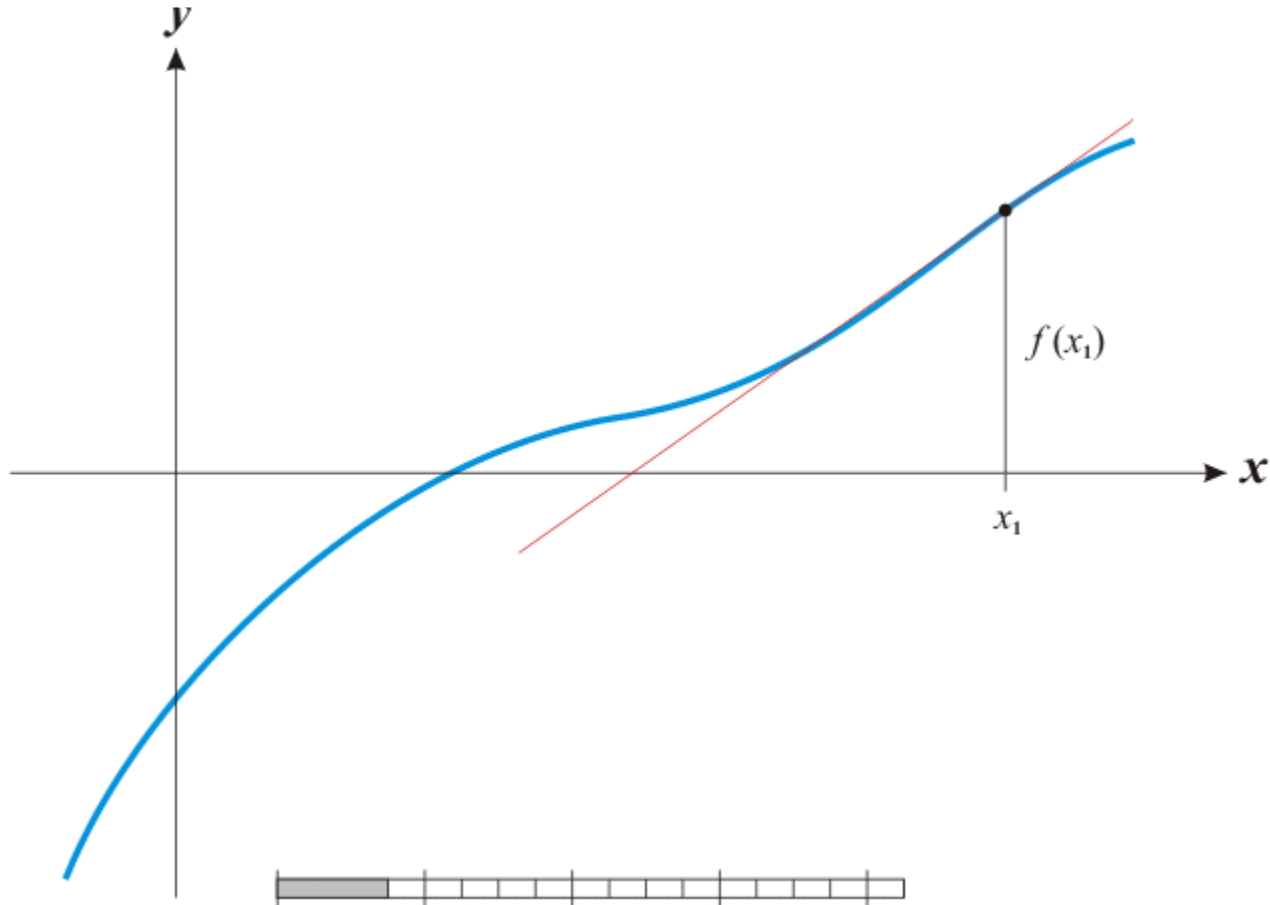
- **Newton's Method** is an iterative equation solver.
- It is an algorithm to find the roots of a polynomial function.
- In the simple, one-variable case, Newton's Method is implemented as follows:

1. Find the tangent line to $f(x)$ at point (x_n, y_n)
 - $y = f'(x_n)(x - x_n) + f(x_n)$
2. Find the x -intercept of the tangent line, x_{n+1}
 - $0 = f'(x_n)(x_{n+1} - x_n) + f(x_n)$
 - $-f(x_n) = f'(x_n)(x_{n+1} - x_n)$
 - $x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}$
3. Find the y value at the x -intercept.
 - $y_{n+1} = f(x_{n+1})$
4. If $y_{n+1} - y_n \approx 0$:
 - **return** y_{n+1} because we've converged!
5. Else update point (x_n, y_n) , and iterate
 - $x = x_{n+1}, y = y_{n+1}$, **goto** (1).

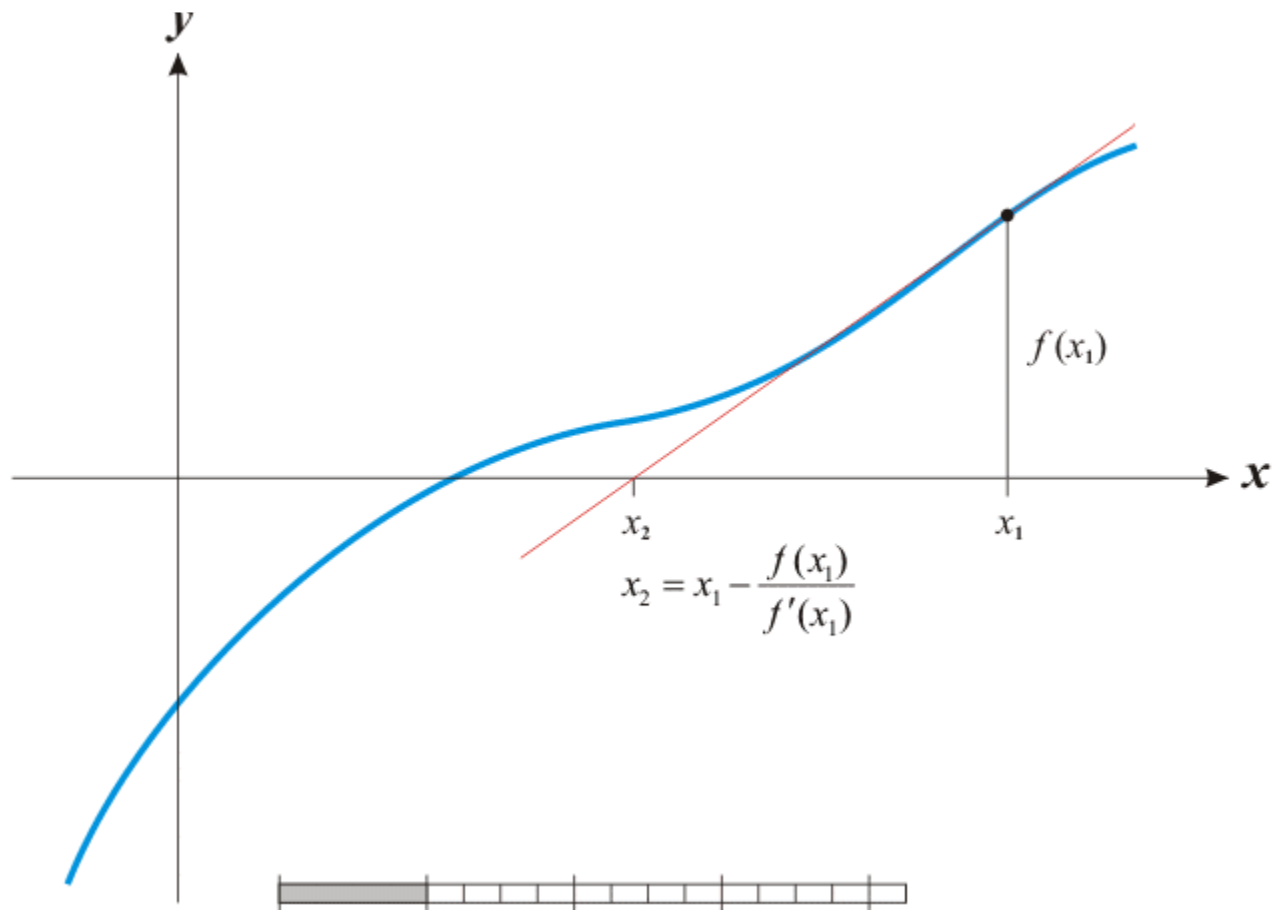
Newton's Method



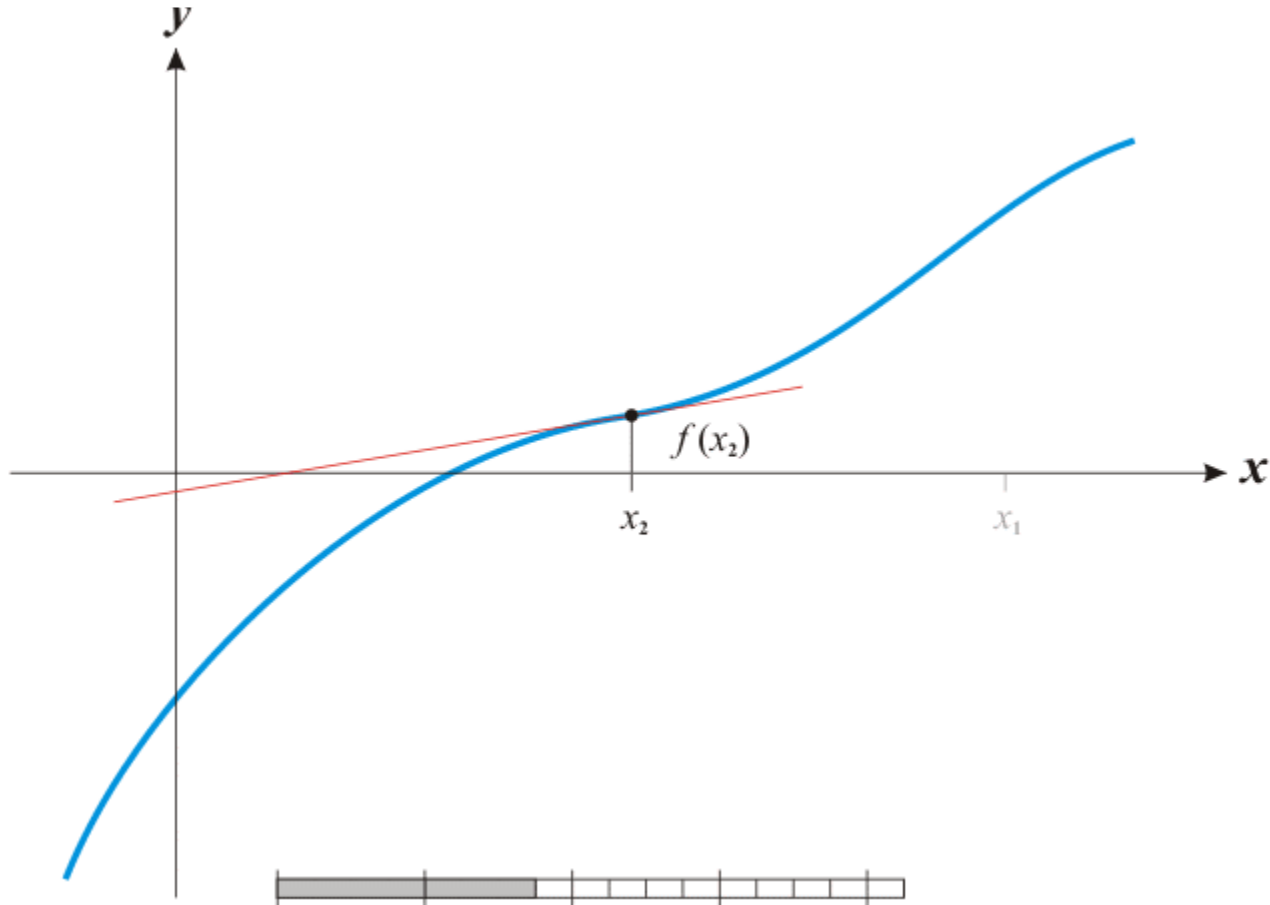
Newton's Method



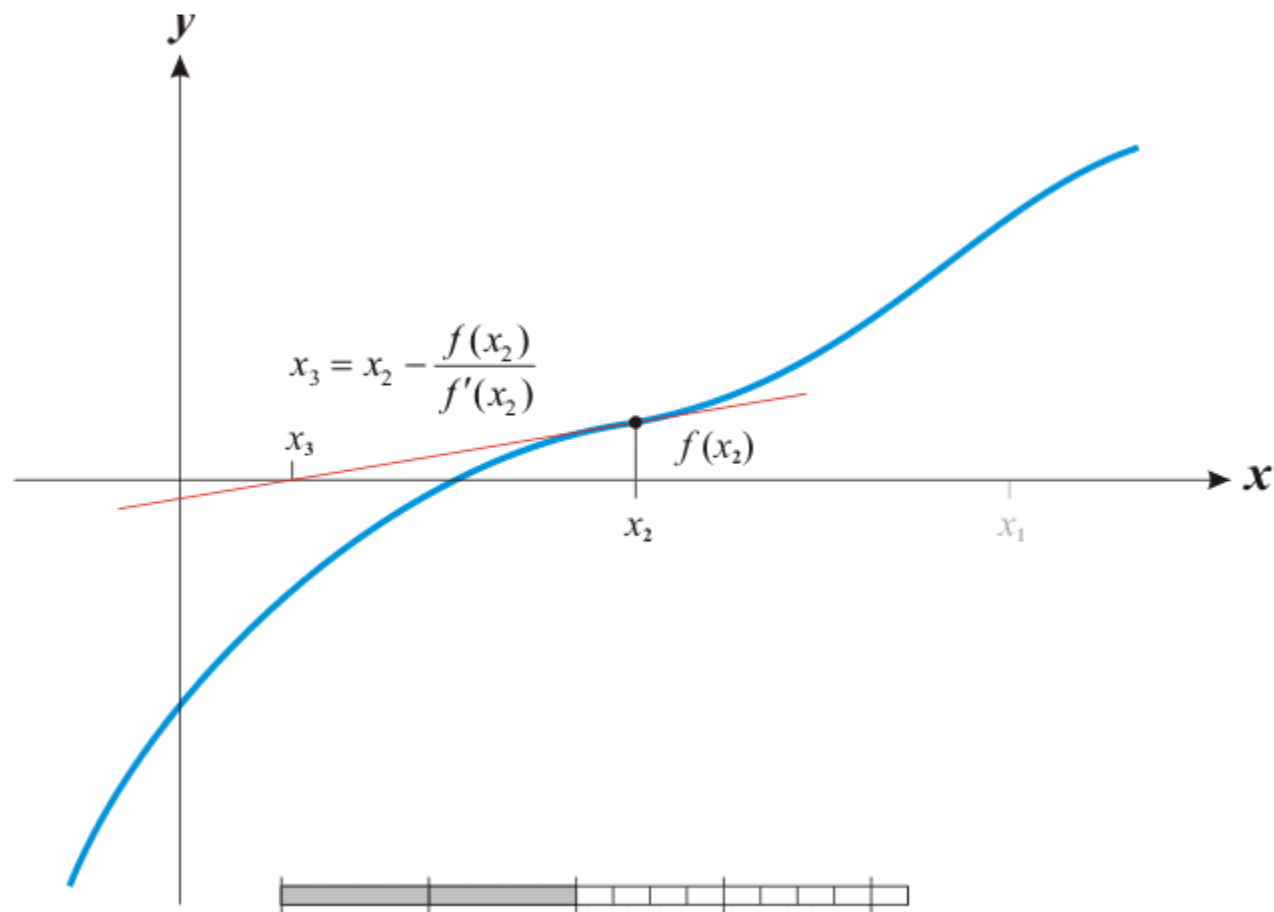
Newton's Method



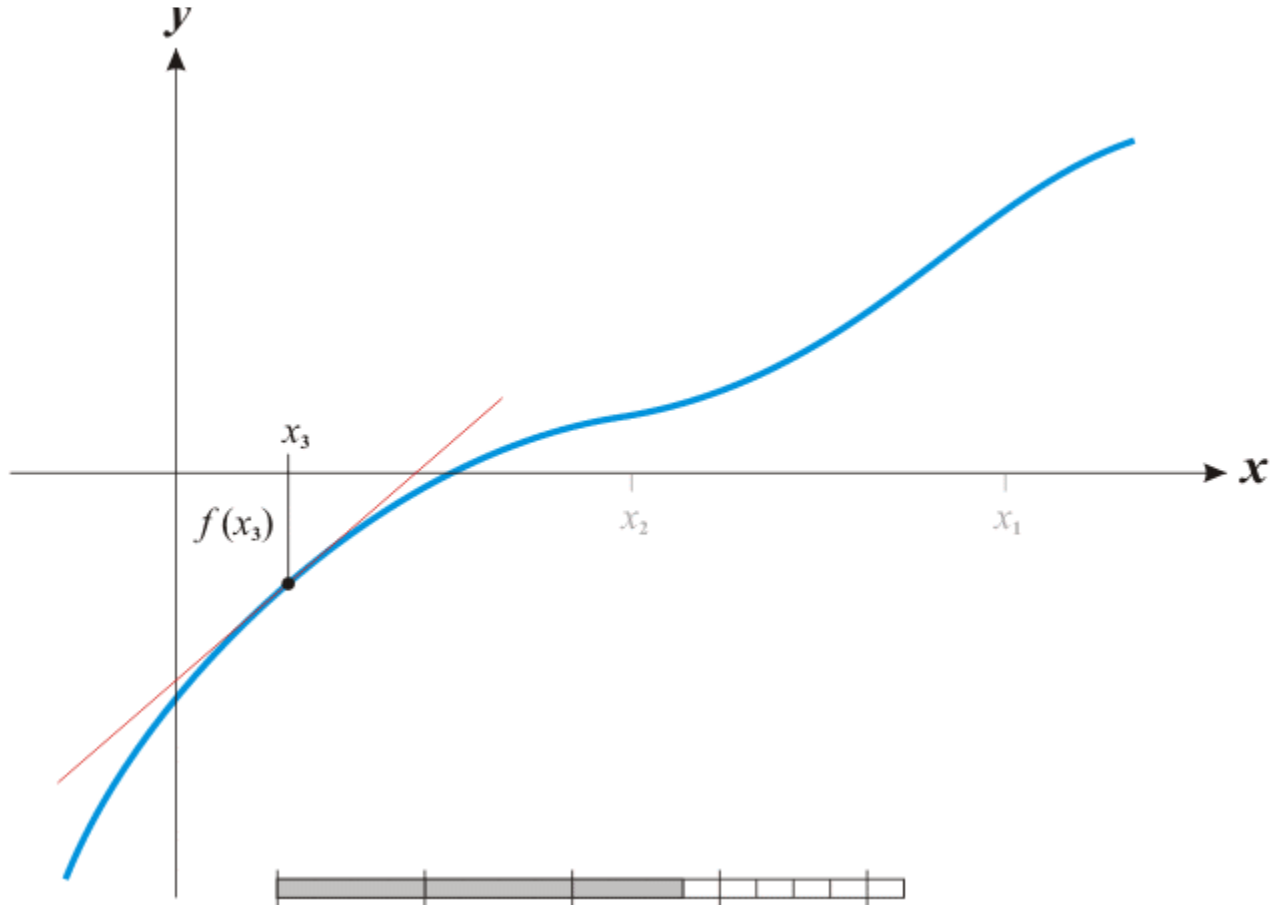
Newton's Method



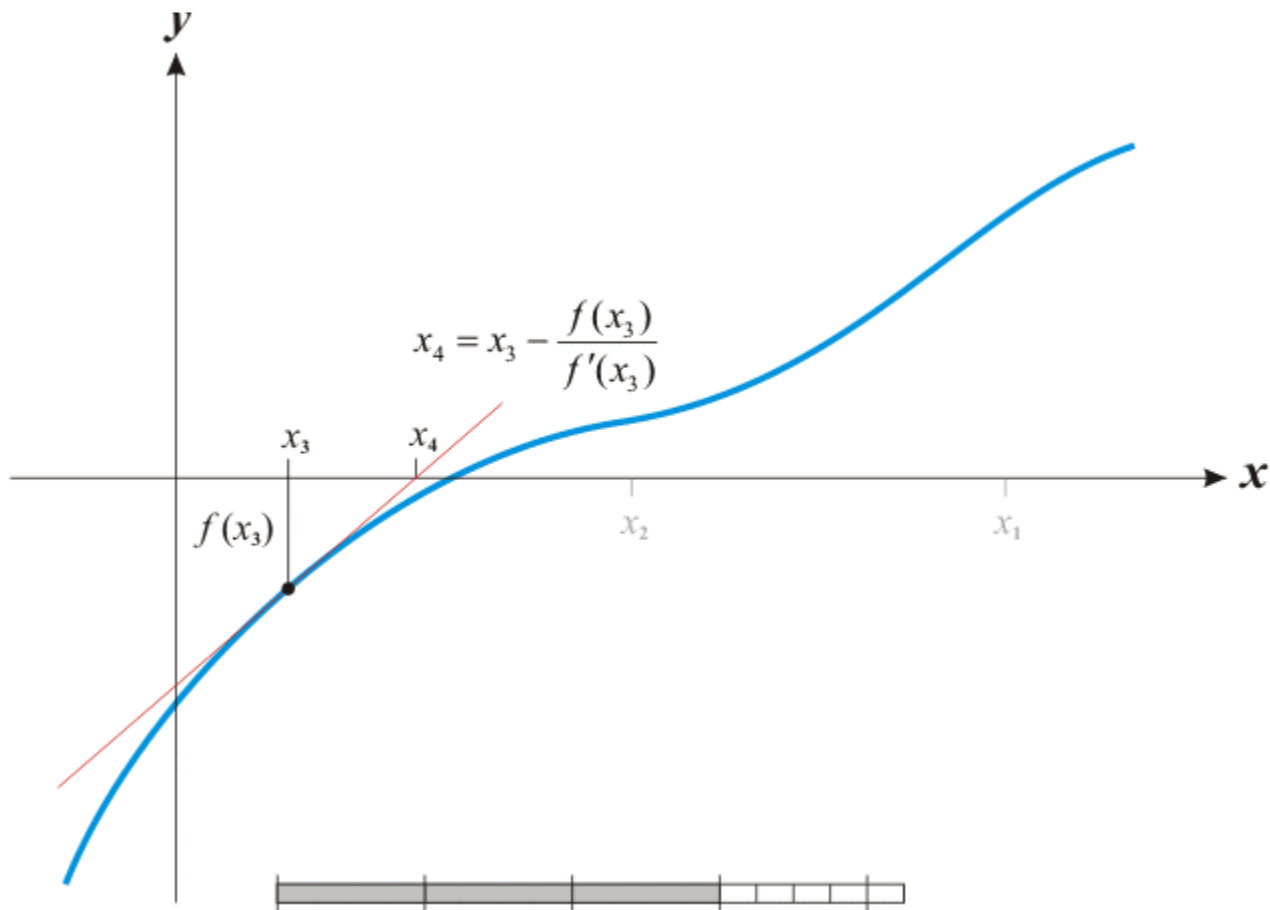
Newton's Method



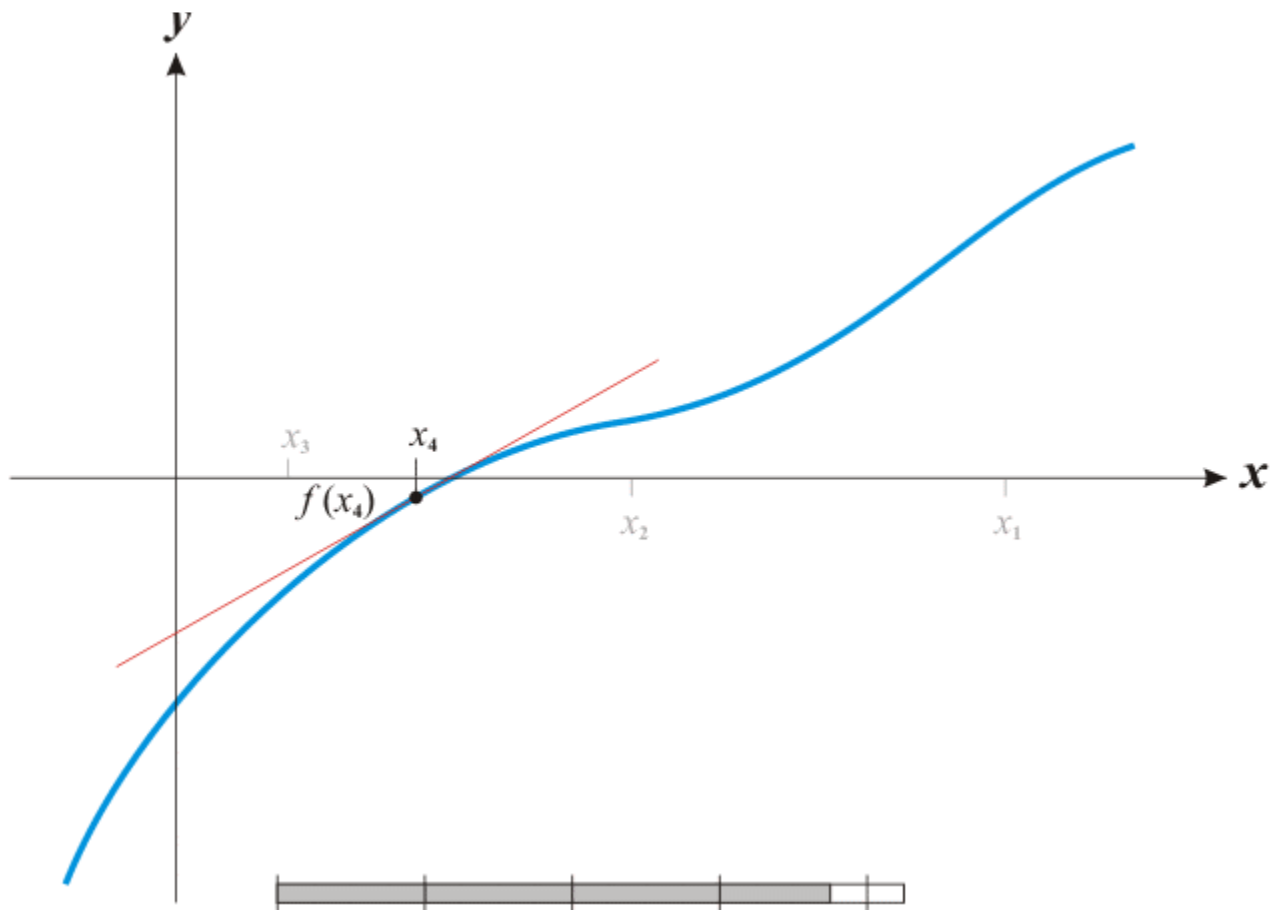
Newton's Method



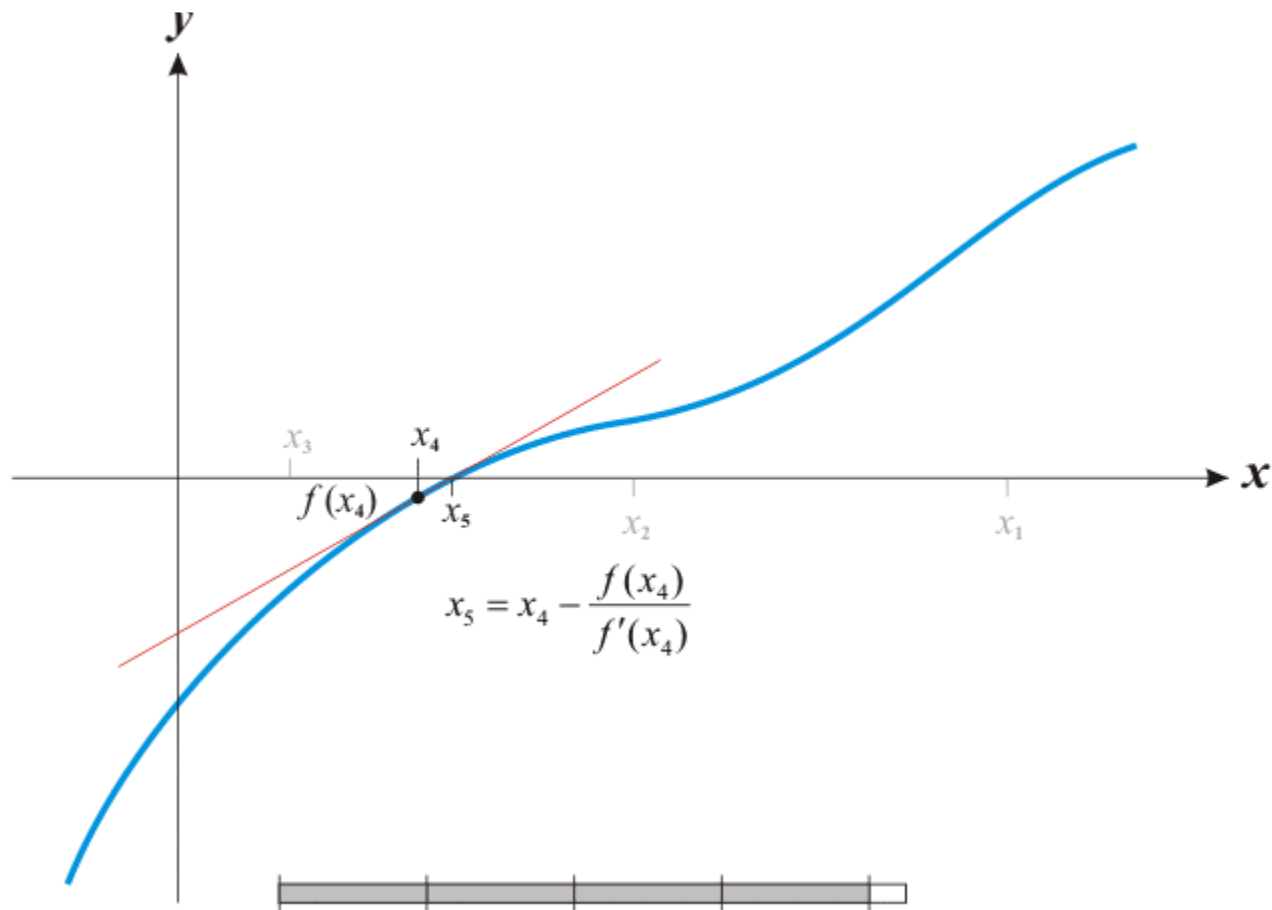
Newton's Method



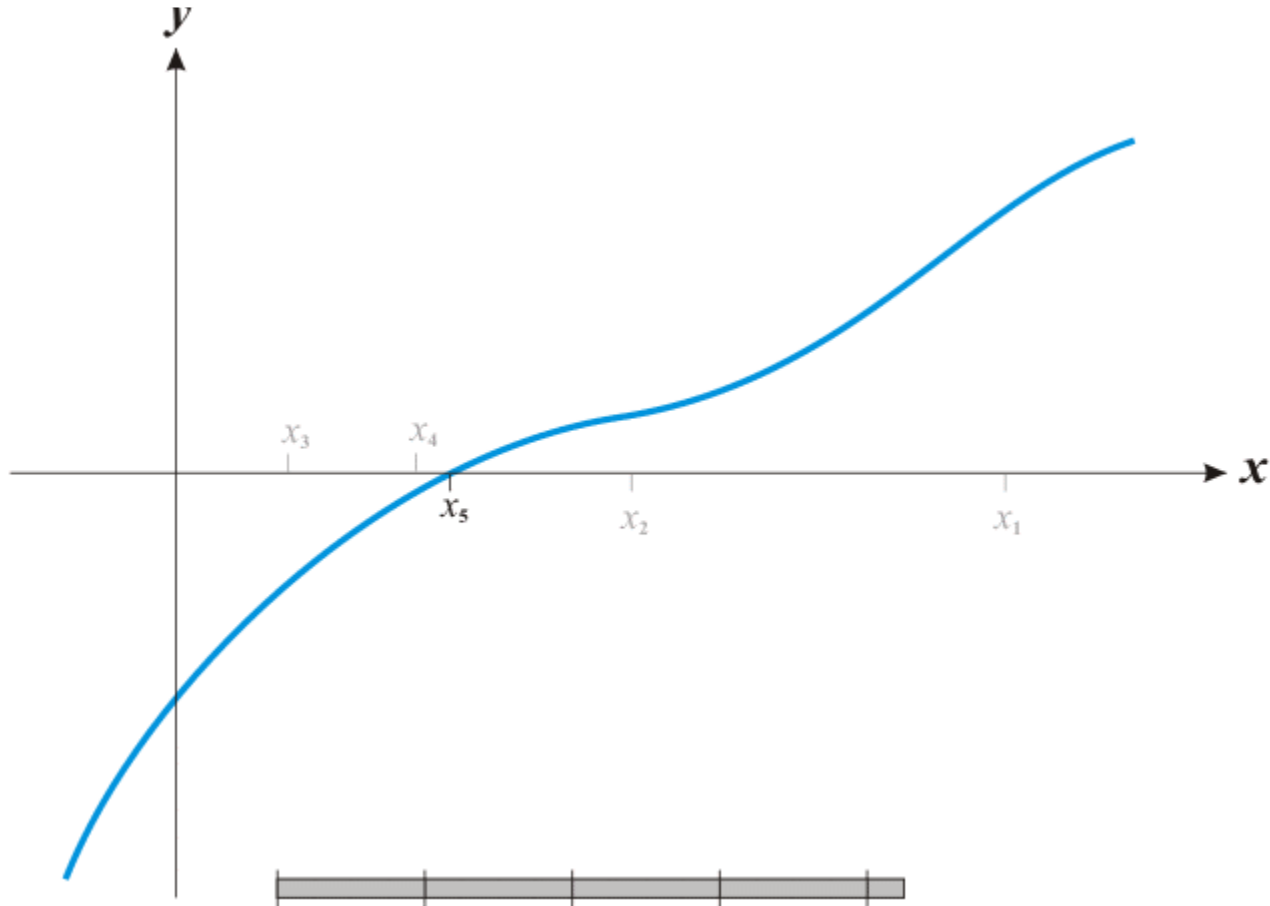
Newton's Method



Newton's Method



Newton's Method



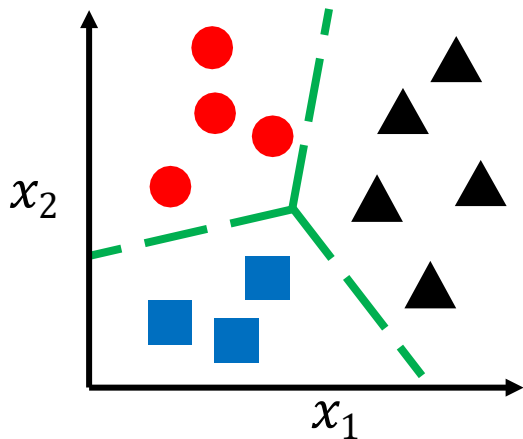
Newton's Method in N-dimensions

- Recall that in n-dimensions, we replace single-variable derivatives with a vector of partial derivatives called the **gradient**.
- Thus, our update rule, in its multivariate form, becomes our vector of parameters $\hat{x}$, minus the scalar $f(\hat{x})$ multiplied by the inverse of the gradient vector:

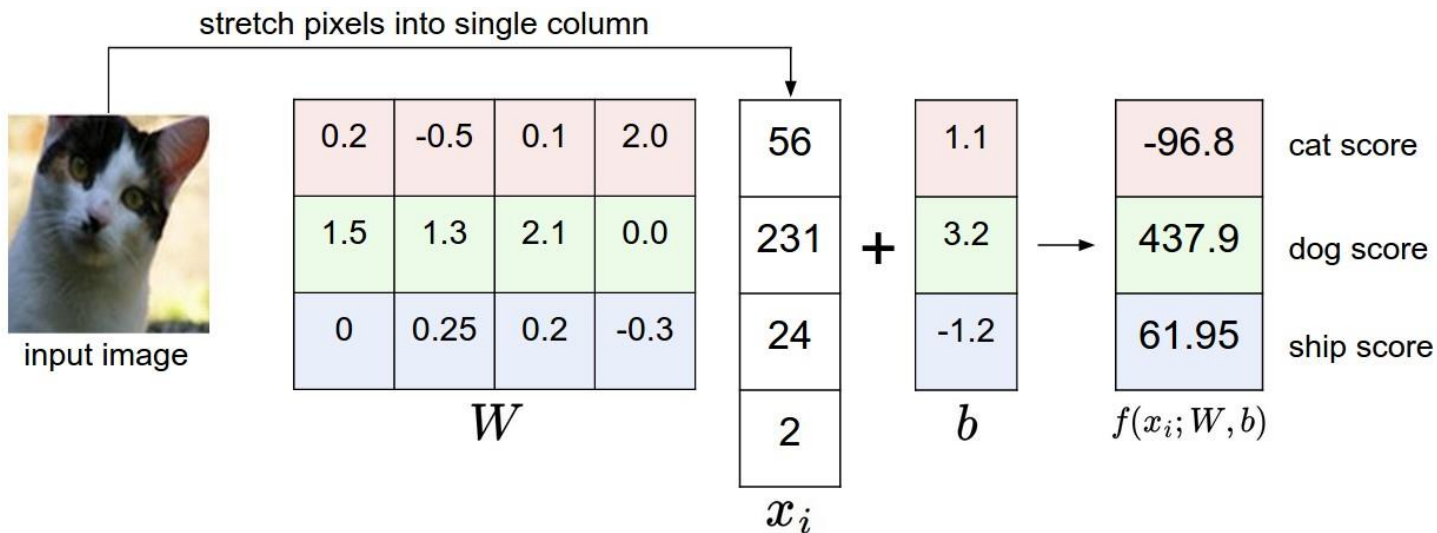
$$\hat{x}_{n+1} = \hat{x}_n - f(\hat{x}_n) * \nabla f(\hat{x}_n)^{-1}$$

Multi-Class Classification

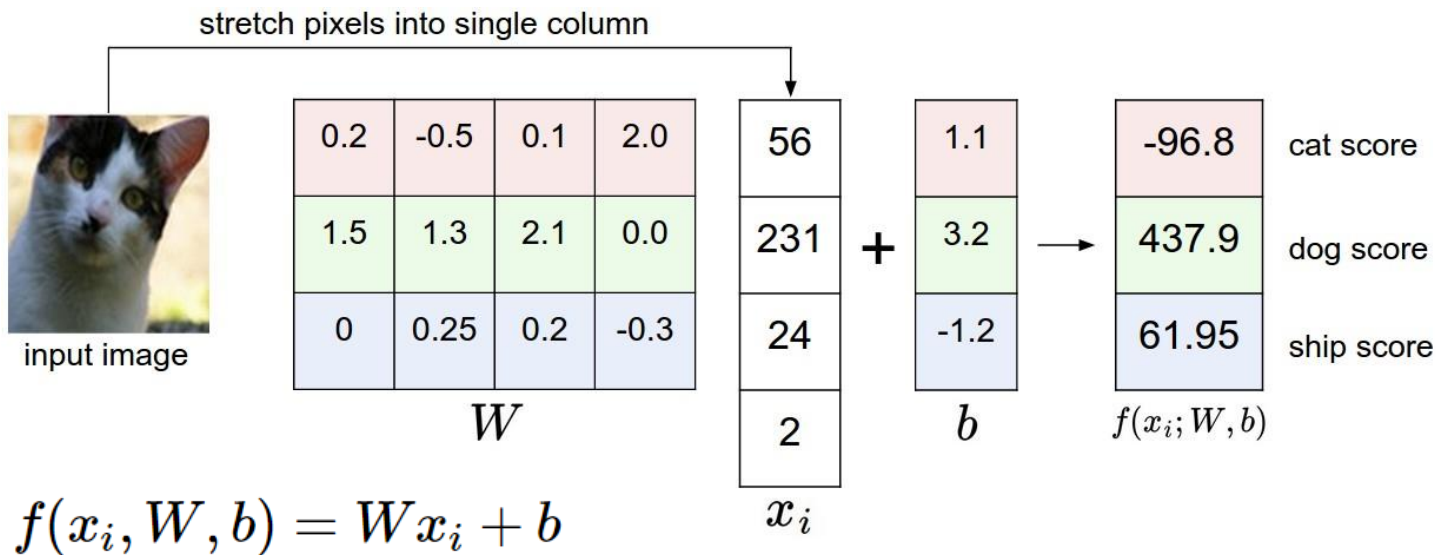
- What about more than two classes?
 - **Disease diagnosis:** healthy, cold, flu, pneumonia
 - **Object classification:** desk, chair, monitor, bookcase
 - In general, consider a finite space of labels $\mathcal{Y}$



Linear Classifiers – More than 2 Class



Linear Classifiers – More than 2 Class



Linear Classifiers – More than 2 Class

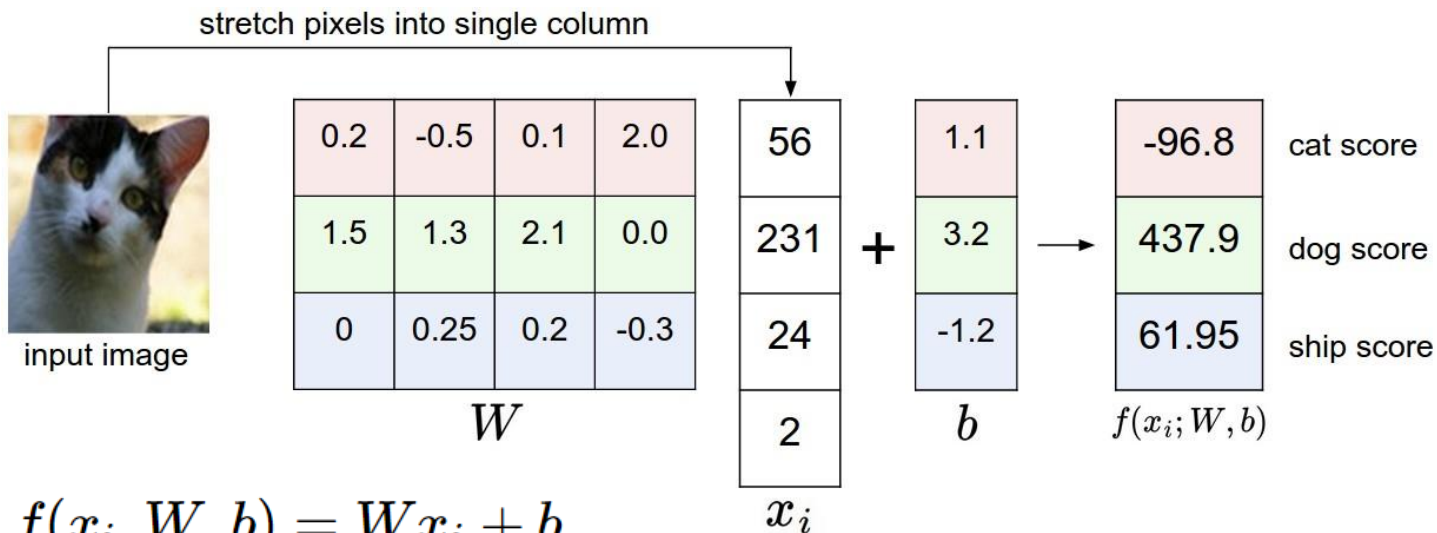
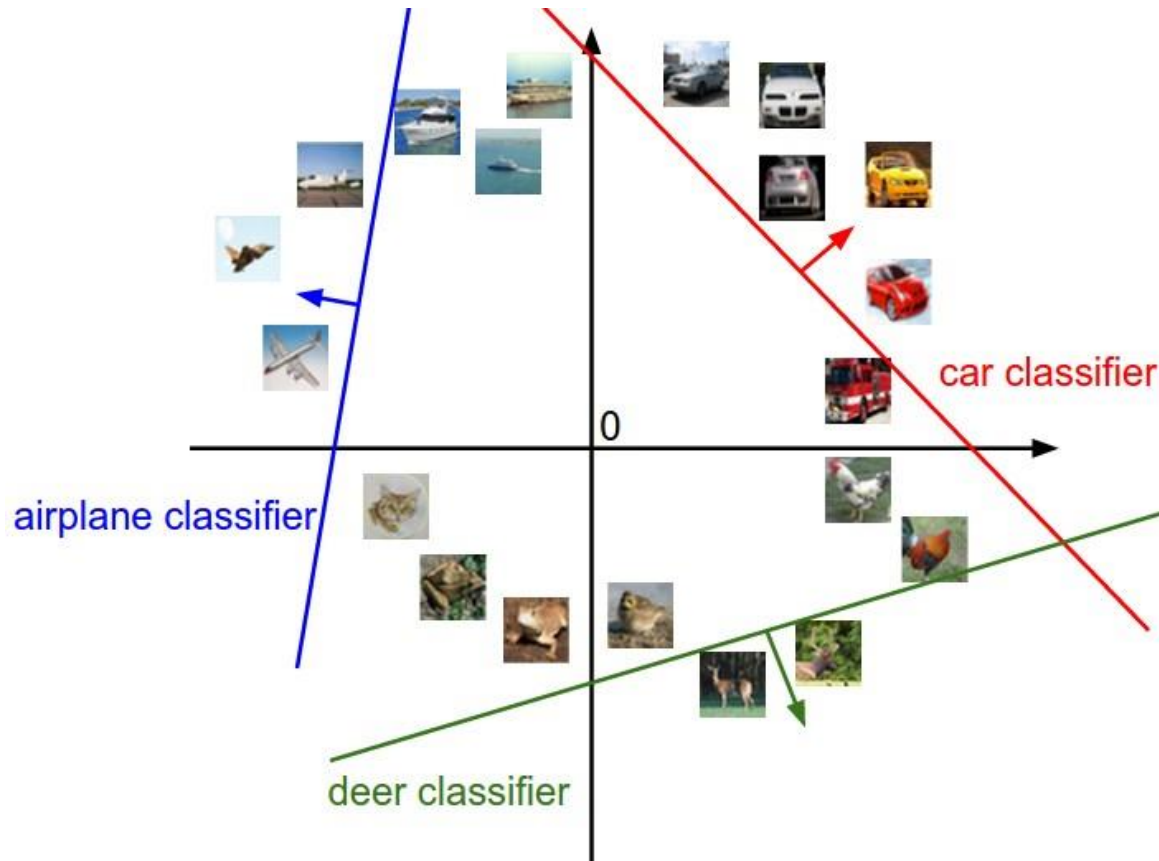


Image x_i has all of its pixels flattened out to a single column vector of shape $[D \times 1]$.

Matrix W (of size $[K \times D]$), and vector b (of size $[K \times 1]$) are the **parameters**.

K is the number of classes.

Analogy of Images as High-dimensional Points



Source: cs231n, <http://cs231n.github.io/linear-classify/>

Interpretation: Linear Classifiers as Template Matching



Example learned weights at the end of learning for CIFAR-10. Note that, for example, the ship template contains a lot of blue pixels as expected. This template will therefore give a high score once it is matched against images of ships on the ocean with an inner product.

Recall: Logistic Regression

- Consider the following choice:

$$p_w(Y = 0 \mid x_i) \propto e^{-\frac{w^\top x_i}{2}} \text{ and } p_w(Y = 1 \mid x_i) \propto e^{\frac{w^\top x_i}{2}}$$

- Then, we have

$$p_w(Y = 1 \mid x_i) = \frac{e^{\frac{w^\top x_i}{2}}}{e^{\frac{w^\top x_i}{2}} + e^{-\frac{w^\top x_i}{2}}} = \frac{1}{1 + e^{-w^\top x_i}}$$

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



- Furthermore, $p_w(Y = 0 \mid x_i) = 1 - \sigma(w^\top x_i)$

Multi-Class Logistic Regression

- **Strategy:** Include separate $\mathbf{w}_y$ for each label $y \in \mathcal{Y} = \{1, \dots, k\}$
- Let $p_w(y | x) \propto e^{\mathbf{w}_y^T x}$, i.e.

$$p_w(y | x) = \frac{e^{\mathbf{w}_y^T x}}{\sum_{y' \in \mathcal{Y}} e^{\mathbf{w}_{y'}^T x}}$$

- We define $\text{softmax}(z_1, \dots, z_k) = \left[\frac{e^{z_1}}{\sum_{i=1}^k e^{z_i}} \quad \dots \quad \frac{e^{z_k}}{\sum_{i=1}^k e^{z_i}} \right]$
- Thus, sometimes called **softmax regression**

Multi-Class Logistic Regression

- Objective function: Negative Log Likelihood

$$\ell(\mathbf{w}; \mathbf{Z}) = -\sum_{i=1}^n \log \frac{e^{\mathbf{w}_{y_i}^\top \mathbf{x}_i}}{\sum_{y' \in \mathcal{Y}} e^{\mathbf{w}_{y'}^\top \mathbf{x}_i}} = -\sum_{i=1}^n \left[\mathbf{w}_{y_i}^\top \mathbf{x}_i - \log \sum_{y' \in \mathcal{Y}} e^{\mathbf{w}_{y'}^\top \mathbf{x}_i} \right]$$

Multi-Class Logistic Regression

- **Define One-hot vector:** $\hat{y}_{ik} \in \{0, 1\}, \sum_k \hat{y}_{ik} = 1, \hat{y}_{iy_i} = 1$
- Then loss will become:

$$\ell(w; Z) = - \sum_i \sum_k \hat{y}_{ik} \log p_{ik}$$

$$\text{where: } p_{ik} = \frac{e^{w_k^T x_i}}{\sum_{k' \in \mathcal{Y}} e^{w_{k'}^T x_i}}$$

Multi-Class Logistic Regression

- **Define One-hot vector:** $\hat{y}_{ik} \in \{0, 1\}, \sum_k \hat{y}_{ik} = 1, \hat{y}_{iy_i} = 1$
- Then loss will become:

$$\ell(\mathbf{w}; \mathbf{Z}) = - \sum_i \sum_k \hat{y}_{ik} \log p_{ik}$$

$$\text{where: } p_{ik} = \frac{e^{\mathbf{w}_k^\top \mathbf{x}_i}}{\sum_{k' \in \mathcal{Y}} e^{\mathbf{w}_{k'}^\top \mathbf{x}_i}}$$

Let $K=2$, then it will fall back to the form of BCE!

- Then loss will become:

$$\ell(\mathbf{w}; \mathbf{Z}) = - \sum_i y_i \log p_i + (1 - y_i) \log(1 - p_i)$$

where:

$$p_i = \frac{e^{\mathbf{w}_0^\top \mathbf{x}_i}}{e^{\mathbf{w}_0^\top \mathbf{x}_i} + e^{\mathbf{w}_1^\top \mathbf{x}_i}} = \sigma((\mathbf{w}_0^\top - \mathbf{w}_1^\top) \mathbf{x}_i)$$

Multi-Class Logistic Regression

$$p_{ik} = \frac{e^{w_k^T x_i}}{\sum_{k' \in \mathcal{Y}} e^{w_{k'}^T x_i}}$$

- Gradient Derivation:

$$\ell(\mathbf{w}; \mathbf{Z}) = - \sum_i \sum_k \hat{y}_{ik} \log p_{ik} = - \sum_{i=1}^n \left[\sum_{k \in \mathcal{Y}} \left[\hat{y}_{ik} w_k^T x_i - \log \sum_{k' \in \mathcal{Y}} e^{w_{k'}^T x_i} \right] \right]$$

- We can derive gradients for both terms:

$$\nabla_{w_k} \ell(\mathbf{w}; \mathbf{Z}) = - \sum_{i=1}^n \hat{y}_{ik} x_i - \frac{e^{w_k^T x_i}}{\sum_{k' \in \mathcal{Y}} e^{w_{k'}^T x_i}} x_i = - \sum_{i=1}^n (p_{ik} - \hat{y}_{ik}) \cdot x_i$$

Multi-Class Logistic Regression

$$p_{ik} = \frac{e^{w_k^T x_i}}{\sum_{k' \in \mathcal{Y}} e^{w_{k'}^T x_i}}$$

- Gradient Derivation:

$$\ell(\mathbf{w}; \mathbf{Z}) = - \sum_i \sum_k \hat{y}_{ik} \log p_{ik} = - \sum_{i=1}^n \left[\sum_{k \in \mathcal{Y}} \left[\hat{y}_{ik} w_k^T x_i - \log \sum_{k' \in \mathcal{Y}} w_{k'}^T x_i \right] \right]$$

- We can derive gradients for both terms:

$$\nabla_{w_k} \ell(\mathbf{w}; \mathbf{Z}) = - \sum_{i=1}^n \hat{y}_{ik} x_i - \frac{e^{w_k^T x_i}}{\sum_{k' \in \mathcal{Y}} e^{w_{k'}^T x_i}} x_i = - \sum_{i=1}^n (p_{ik} - \hat{y}_{ik}) \cdot x_i$$

Diagram illustrating the gradient derivation for the second term:

- Probability**: Points to p_{ik} in the expression $(p_{ik} - \hat{y}_{ik}) \cdot x_i$.
- Label**: Points to $\hat{y}_{ik}$ in the expression $(p_{ik} - \hat{y}_{ik}) \cdot x_i$.
- Input**: Points to x_i in the expression $(p_{ik} - \hat{y}_{ik}) \cdot x_i$.

Multi-Class Logistic Regression

$$p_{ik} = \frac{e^{w_k^T x_i}}{\sum_{k' \in \mathcal{Y}} e^{w_{k'}^T x_i}}$$

- Gradient Derivation:

$$\ell(\mathbf{w}; \mathbf{Z}) = - \sum_i \sum_k \hat{y}_{ik} \log p_{ik} = - \sum_{i=1}^n \left[\sum_{k \in \mathcal{Y}} \left[\hat{y}_{ik} w_k^T x_i - \log \sum_{k' \in \mathcal{Y}} w_{k'}^T x_i \right] \right]$$

- We can derive gradients for both terms:

$$\nabla_{w_k} \ell(\mathbf{w}; \mathbf{Z}) = - \sum_{i=1}^n \hat{y}_{ik} x_i - \frac{e^{w_k^T x_i}}{\sum_{k' \in \mathcal{Y}} e^{w_{k'}^T x_i}} x_i = - \sum_{i=1}^n (p_{ik} - \hat{y}_{ik}) \cdot x_i$$

Diagram labels for the final equation:

- Probability: points to p_{ik}
- Label: points to $\hat{y}_{ik}$
- Input: points to x_i

- It has a very similar form as in binary case:

$$\nabla_{\mathbf{w}} \ell(\mathbf{w}; \mathbf{Z}) = - \sum_{i=1}^n (\sigma(\mathbf{w}^T x_i) - y_i) \cdot x_i$$

Multi-Class Logistic Regression

- **Aside:** This loss is also known as cross-entropy:

$$\ell(\mathbf{w}; \mathbf{Z}) = - \sum_i \sum_k q_{ik} \log p_{ik} = \sum_i H(q_i, p_i)$$

- **Label smoothing:** $\hat{\mathbf{y}}_{ik}$ can be any simplex vector (q), not necessarily one-hot
 - $[0, 0, \dots, 1, \dots, 0] \rightarrow [\epsilon/K, \epsilon/K, \dots, 1 - \epsilon + \epsilon/K, \dots, \epsilon/K]$

$$\nabla_{\mathbf{w}_k} \ell(\mathbf{w}; \mathbf{Z}) = - \sum_{i=1}^n (p_{ik} - q_{ik}) \cdot \mathbf{x}_i$$

Multi-Class Logistic Regression

- Model family

$$f_{\mathbf{w}}(\mathbf{x}) = \arg \max_y p_{\mathbf{w}}(\mathbf{y} \mid \mathbf{x}) = \arg \max_y \frac{e^{\mathbf{w}_y^T \mathbf{x}}}{\sum_{k' \in \mathcal{Y}} e^{\mathbf{w}_{k'}^T \mathbf{x}}} = \arg \max_y \mathbf{w}_y^T \mathbf{x}$$

- Optimization

- Gradient descent on Negative Log Likelihood

$$\nabla_{\mathbf{w}_k} \ell(\mathbf{w}; \mathbf{Z}) = - \sum_{i=1}^n (p_{ik} - \hat{y}_{ik}) \cdot \mathbf{x}_i$$

- Simultaneously update all parameters $\{\mathbf{w}_y\}_{y \in \mathcal{Y}}$

Linear Classifiers

“Parametric” classifier: model defined by a small number of parameters (w , b)

Pros:

- Very fast at test time

Cons:

- Slow at training time: need to estimate the parameters
- Data may not be linearly separable

Classification Metrics

- While we minimize the Negative Log Likelihood, we often evaluate using **accuracy**
- However, even accuracy isn't necessarily the “right” metric
 - If 99% of labels are negative (i.e., $y_i = 0$), accuracy of $f_w(x) = 0$ is 99%!
 - For instance, very few patients test positive for most diseases
 - “Imbalanced data”
- What are alternative metrics for these settings?

Classification Metrics

- **Classify test examples as follows:**
 - **True Positive (TP):** Actually positive, predicted positive
 - **False Negative (FN):** Actually positive, predicted negative
 - **True Negative (TN):** Actually negative, predicted negative
 - **False Positive (FP):** Actually negative, predicted positive
- Many metrics expressed in terms of these; for example:

$$\text{accuracy} = \frac{TP + TN}{n} \quad \text{error} = 1 - \text{accuracy} = \frac{FP + FN}{n}$$

Confusion Matrix

		Predicted Class	
		Yes	No
Actual Class	Yes	TP	FN
	No	FP	TN

Confusion Matrix

		Predicted Class	
		Yes	No
Actual Class	Yes	3 TP	4 FN
	No	6 FP	37 TN

Accuracy = 0.8

Classification Metrics

- For imbalanced metrics, we roughly want to disentangle:
 - Accuracy on “positive examples”
 - Accuracy on “negative examples”
- Different definitions are possible (and lead to different meanings)!

Sensitivity & Specificity

- **Sensitivity:** What fraction of **actual positives** are **predicted positive**?
 - **Good sensitivity:** If you have the disease, the test correctly detects it
 - Also called **true positive rate**
- **Specificity:** What fraction of **actual negatives** are **predicted negative**?
 - **Good specificity:** If you do not have the disease, the test says so
 - Also called **true negative rate**
- Commonly used in medicine

Sensitivity & Specificity

		Predicted Class		
		Yes	No	
Actual Class	Yes	TP	FN	sensitivity = $\frac{TP}{TP + FN}$
	No	FP	TN	specificity = $\frac{TN}{TN + FP}$

Sensitivity & Specificity

		Predicted Class		
		Yes	No	
Actual Class	Yes	3 TP 4 FN		sensitivity = $\frac{TP}{TP + FN}$
	No	6 FP 37 TN		specificity = $\frac{TN}{TN + FP}$

Sensitivity & Specificity

		Predicted Class		
		Yes	No	
Actual Class	Yes	3 TP	4 FN	sensitivity = 3/7
	No	6 FP	37 TN	specificity = 37/43

Precision & Recall

- **Recall:** What fraction of **actual positives** are **predicted positive**?
 - **Good recall:** If you have the disease, the test correctly detects it
 - Also called the **true positive rate** (and sensitivity)
- **Precision:** What fraction of **predicted positives** are **actual positives**?
 - **Good precision:** If the test says you have the disease, then you have it
 - Also called **positive predictive value**
- Used in information retrieval, NLP

Precision & Recall

		Predicted Class	
		Yes	No
Actual Class	Yes	TP	FN
	No	FP	TN

$$\text{recall} = \frac{TP}{TP + FN}$$

$$\text{precision} = \frac{TP}{TP + FP}$$

Precision & Recall

		Predicted Class	
		Yes	No
Actual Class	Yes	3 TP	4 FN
	No	6 FP	37 TN

$$\text{recall} = \frac{TP}{TP + FN}$$

$$\text{precision} = \frac{TP}{TP + FP}$$

Precision & Recall

		Predicted Class		
		Yes	No	
Actual Class	Yes	3 TP	4 FN	recall = 3/7
	No	6 FP	37 TN	

precision = 3/9

Classification Metrics

- **How to obtain a single metric?**

- Combination, e.g., $F_1 \text{ Score} = \frac{2 \cdot \text{precision} \cdot \text{recall}}{\text{precision} + \text{recall}}$ is the harmonic mean

- **How to choose the “right” metric?**

- No generally correct answer
- Depends on the goals for the specific problem/domain

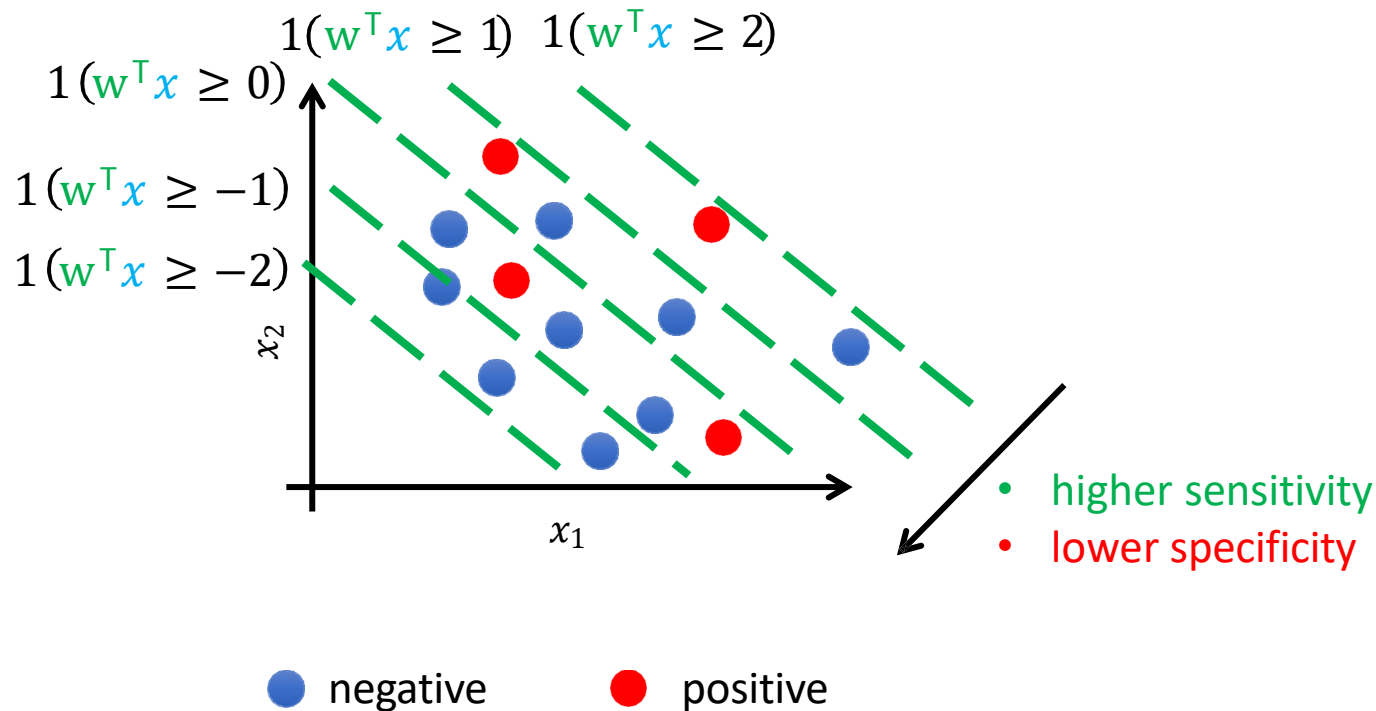
Optimizing Prediction Threshold

- Consider **hyperparameter** τ for the threshold:

$$f_w(x) = 1(w^T x \geq 0)$$

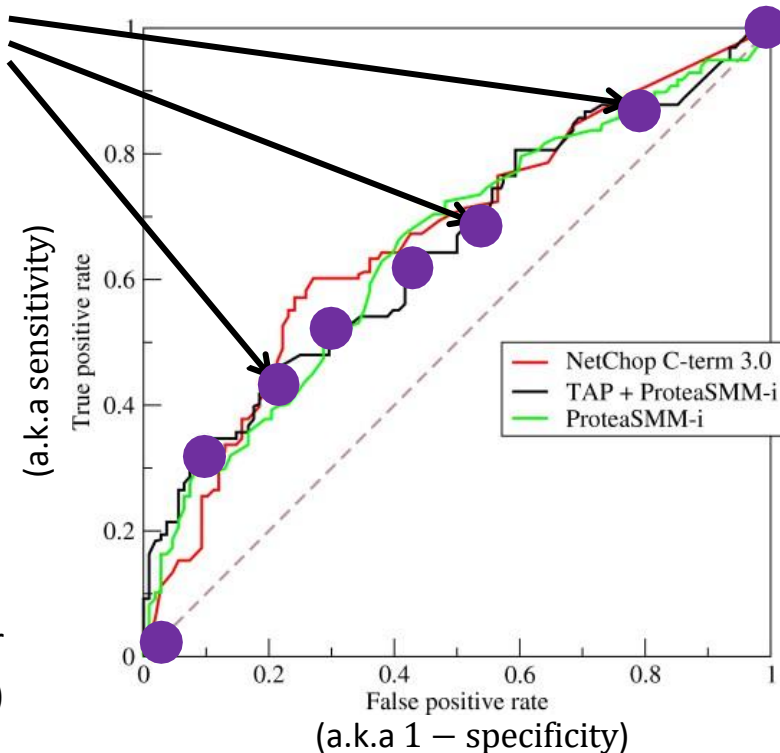

$$f_w(x) = 1(w^T x \geq \tau)$$

Optimizing Prediction Threshold



Visualization: ROC Curve

Each point on this curve corresponds to a choice of τ

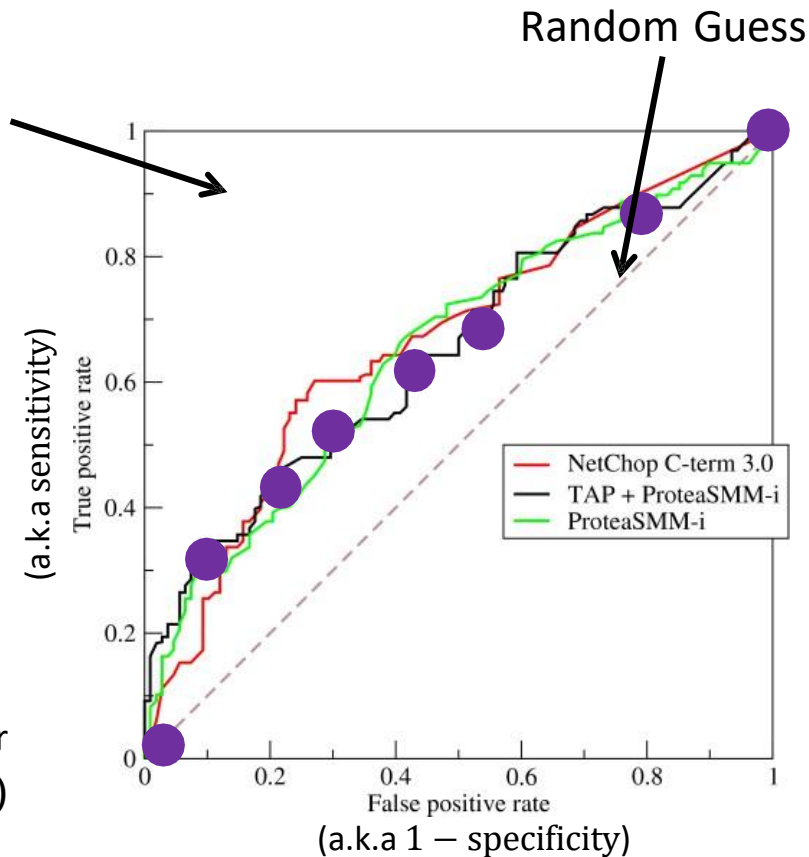


AUC-ROC: Area under ROC curve is another metric people consider when evaluating $\hat{w}(z)$

ROC =
Receiver
Operating
Characteristic

Visualization: ROC Curve

Better Area (High
TPR, low FPR)



ROC =
Receiver
Operating
Characteristic

AUC-ROC: Area under
ROC curve is another
metric people consider
when evaluating $\hat{w}(z)$

Acknowledgement

Most of the slides are based on the ML course of:

- CS231n, Stanford University
- Matt Barnes and Matt Golub, University of Washington
- Mingmin Zhao and Jiatao Gu, University of Pennsylvania

THANK YOU