

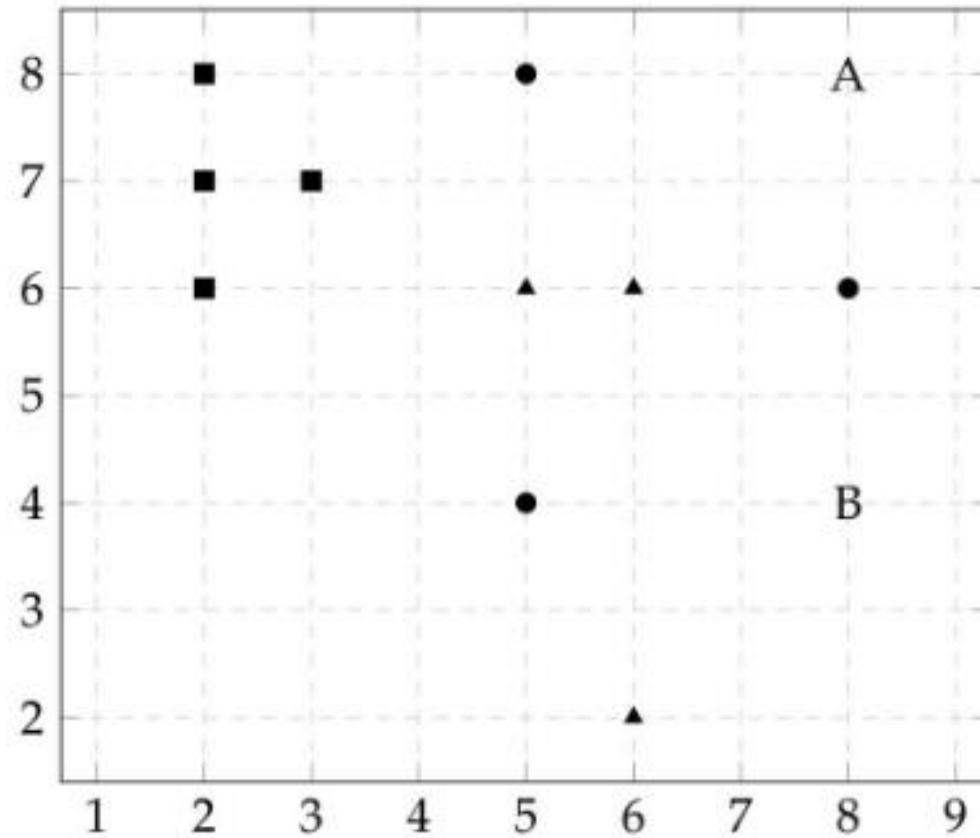
CSCI 567 Discussion

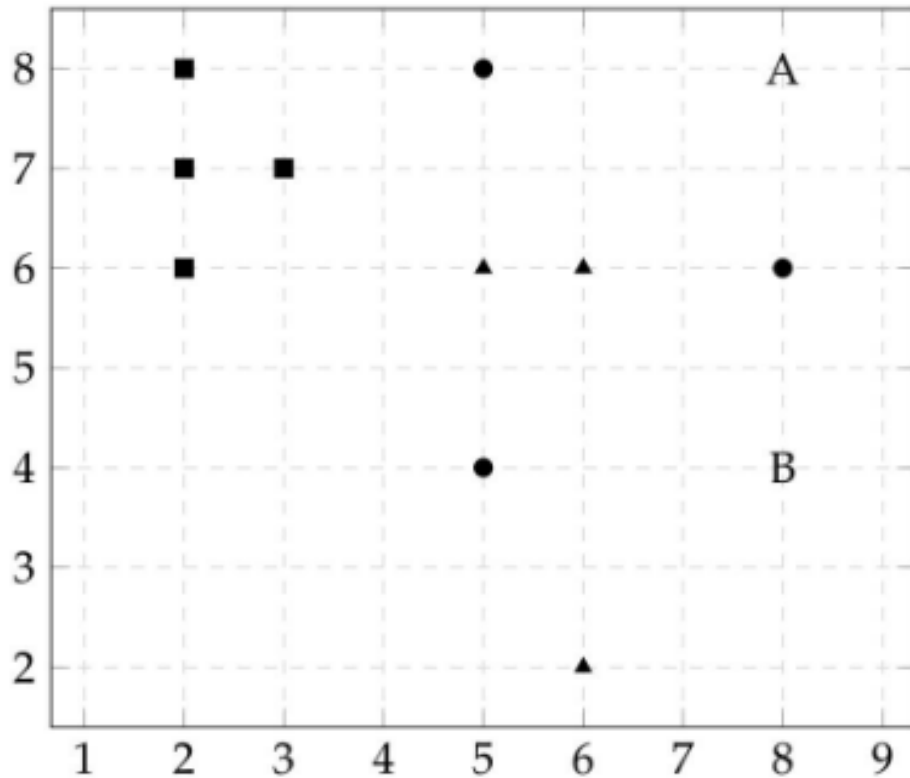
Week 4 - HW1 Review

Problem 1 Nearest Neighbor Classification

(10 points)

For the data given below, squares, triangles, and circles are three different classes in the training set, and A and B are two test points with an unknown class. We denote the total number of training points as N (which equals 10) and consider K-nearest-neighbor (KNN) classifier with L1 distance.





- Three classes:

■ Squares ▲ Triangles ● Circles

- Two test points:

A and B

- K-nearest neighbor (**KNN**) with **L1** distance

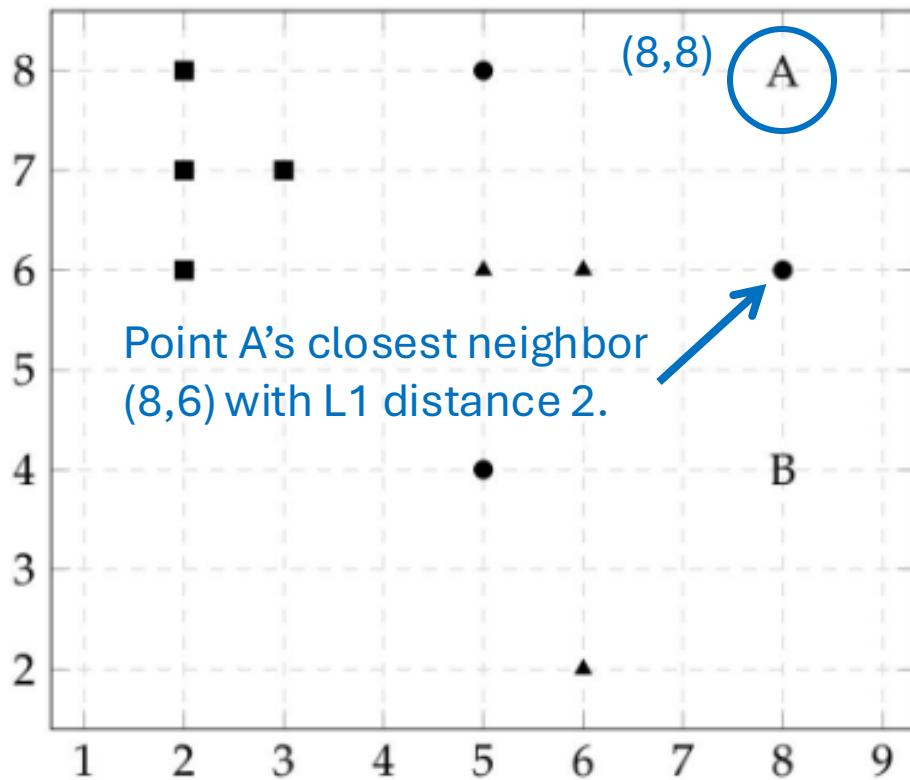
$$d((x_1, y_1), (x_2, y_2)) = |x_1 - x_2| + |y_1 - y_2|$$

1. What is the test point A classified as for $K = 1$? Explain briefly.

(2 points)

Recall KNN with $K=1$:

- We only look at the single closest neighbor.
- Whatever class that neighbor belongs to → the predicted class.



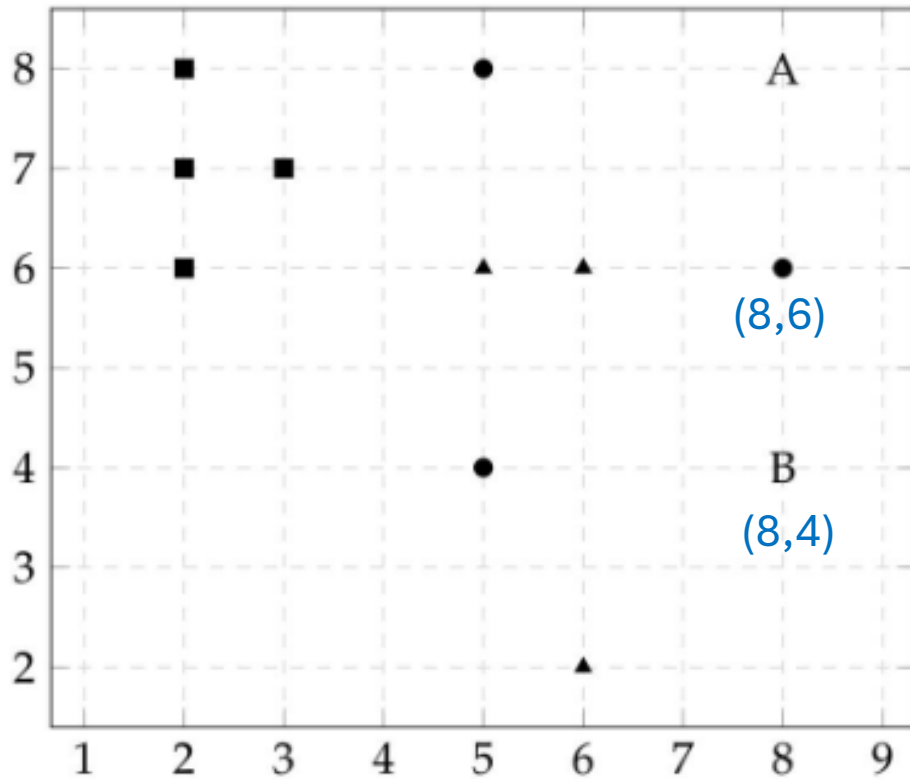
- Three classes:
 ■ Squares ▲ Triangles ● Circles
- Two test points:
 A and B
- K-nearest neighbor (**KNN**) with **L1** distance

$$d((x1, y1), (x2, y2)) = |x1 - x2| + |y1 - y2|$$

1. What is the test point A classified as for $K = 1$? Explain briefly. (2 points)

Answer:

- Circle.
- Explanation: The closest point to A is the circle at (8,6) (with L1 distance 2).



- Three classes:

■ Squares ▲ Triangles ● Circles

- Two test points:

A and B

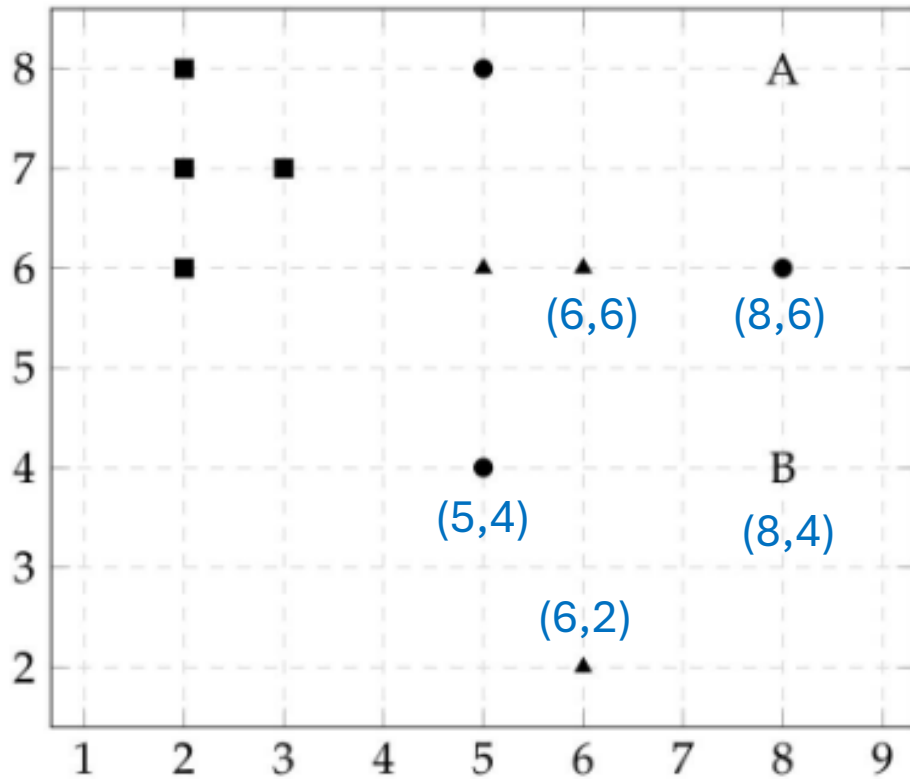
- K-nearest neighbor (**KNN**) with **L1** distance

$$d((x1, y1), (x2, y2)) = |x1 - x2| + |y1 - y2|$$

2. What is the smallest odd value of K for KNN to predict triangle for the test point B? Explain briefly. (4 points)

K = 1:

- Neighbors: ●
● at (8,6) with dist = 2.
- The prediction of B is ●, so $K = 1$ does not work.



- Three classes:

■ Squares ▲ Triangles ● Circles

- Two test points:

A and B

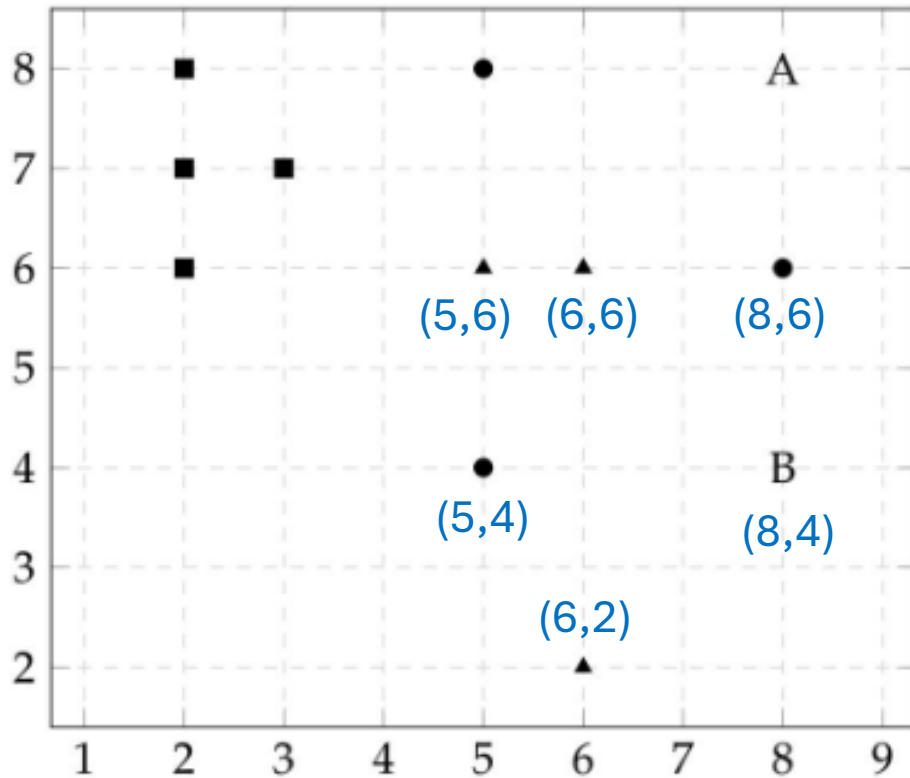
- K-nearest neighbor (**KNN**) with **L1** distance

$$d((x1, y1), (x2, y2)) = |x1 - x2| + |y1 - y2|$$

2. What is the smallest odd value of K for KNN to predict triangle for the test point B? Explain briefly. (4 points)

K = 3:

- Neighbors: ●, ●, ▲
 ● at (8,6) with dist = 2; ● at (5,4) with dist = 3; either ▲ at (6,2) with dist = 4 or ▲ at (6,6) with dist = 4.
- The prediction of B is ●, so $K = 3$ does not work.



- Three classes:

■ Squares ▲ Triangles ● Circles

- Two test points:

A and B

- K-nearest neighbor (**KNN**) with **L1** distance

$$d((x_1, y_1), (x_2, y_2)) = |x_1 - x_2| + |y_1 - y_2|$$

2. What is the smallest odd value of K for KNN to predict triangle for the test point B? Explain briefly. (4 points)

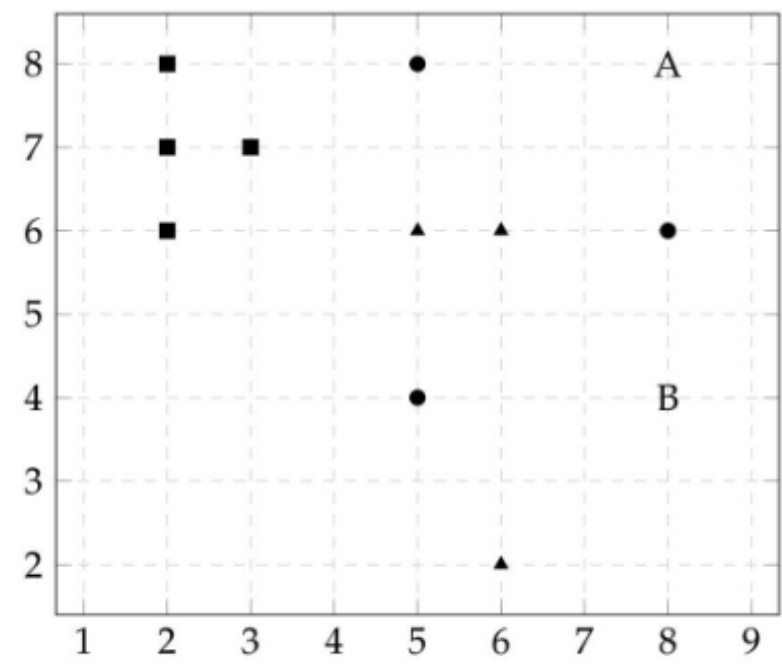
K = 5:

- Neighbors: ●, ●, ▲, ▲, ▲

● at (8,6) with dist = 2; ● at (5,4) with dist = 3; ▲ at (6,2) with dist = 4; ▲ at (6,6) with dist = 4; ▲ at (5,6) with dist = 5.

- The prediction of B is ▲, so $K = 5$ is the answer.

3. Suppose one performs leave-one-out validation (that is, N -fold cross validation) to choose the best hyper-parameter K . List all the points that are misclassified during the N runs when testing the hyper-parameter value $K = 1$, and report the averaged error rate for this hyper-parameter. (4 points)

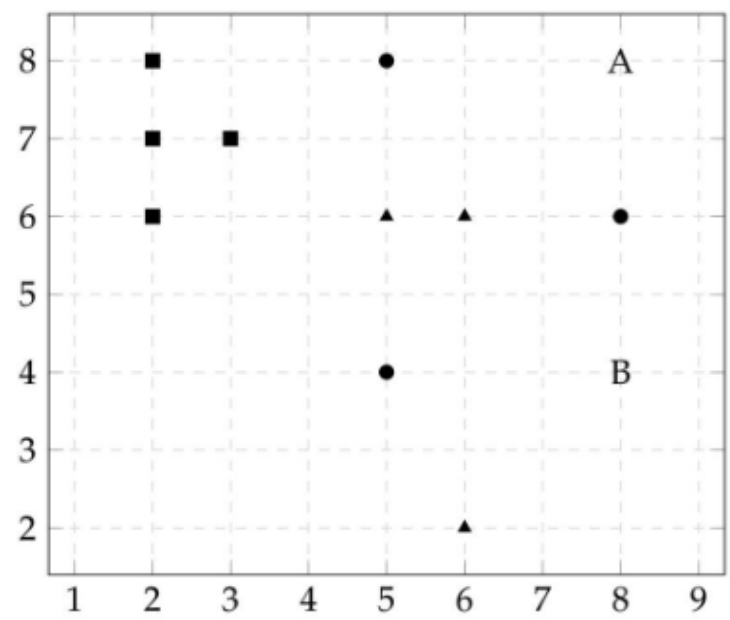


- Leave-one-out validation (N -fold cross validation, where $N = 10$):
- We have $N = 10$ training points.
 - For each run $i = 1 \dots N$:
 - **Leave out 1 point.**
 - **Classify** the point with KNN ($K = 1$) using the other 9 points.
 - **Check** if it is correct.



From: [https://en.wikipedia.org/wiki/Cross-validation_\(statistics\)#/media/File:LOOCV.gif](https://en.wikipedia.org/wiki/Cross-validation_(statistics)#/media/File:LOOCV.gif)

3. Suppose one performs leave-one-out validation (that is, N -fold cross validation) to choose the best hyper-parameter K . List all the points that are misclassified during the N runs when testing the hyper-parameter value $K = 1$, and report the averaged error rate for this hyper-parameter. (4 points)



Run	Leave Out	Closest Neighbor ($K = 1$)	Classification
1	■ at (2,8)	■ at (2,7)	■
2	● at (5,8)	▲ at (5,6)	▲
...	...	...	...
...	● at (8,6)	▲ at (6,6)	▲
...	● at (5,4)	▲ at (5,6)	▲
...	▲ at (6,2)	● at (5,4)	●
10	...	...	...

Answer:

The misclassified points: ▲ at (6,2) and all three ● (at (5,8), (8,6), and (5,4)).
The error rate: $4/10 = 0.4$

Problem 2 Linear Regression

(24 points)

2.1 (10 points) In the class, we discussed L2 regularized least square solution defined as

$$\mathbf{w}_* = \arg \min_{\mathbf{w} \in \mathbb{R}^D} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|_2^2 + \lambda \|\mathbf{w}\|_2^2 \quad (1)$$

where $\mathbf{X} \in \mathbb{R}^{N \times D}$ is the data matrix with each row corresponding to the feature of an example, $\mathbf{y} \in \mathbb{R}^N$ is a vector of all the outcomes, $\|\cdot\|_2$ stands for the L2 norm, and λ is the regularization coefficient. In this problem, we consider a different regularization method:

$$\mathbf{w}'_* = \arg \min_{\mathbf{w} \in \mathbb{R}^D} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|_2^2 + \mathbf{w}^T \mathbf{M} \mathbf{w} \quad (2)$$

where $\mathbf{M} \in \mathbb{R}^{D \times D}$ is a symmetric positive definite matrix.

$$w_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + \lambda \|w\|_2^2 \quad (1)$$

$$w'_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + w^T M w \quad (2)$$

Where,

$X \in \mathbb{R}^{N \times D}$ is the data matrix

$y \in \mathbb{R}^N$ is a vector of all the outcomes

λ is the regularization coefficient

$M \in \mathbb{R}^{D \times D}$ is a symmetric positive definite matrix.

1. Show that the new method is a generalization of the standard L2 regularization by picking a matrix M such that w'_* in Eq. (2) equals w_* in Eq. (1). (2 points)

For example,

$$\|w\|_2^2 = \sum_{j=1}^D w_j^2 = w^\top w$$

$$w = [2 \ 3 \ 4], w^T = \begin{bmatrix} 2 \\ 3 \\ 4 \end{bmatrix}$$

$$w^T w = [2 \ 3 \ 4] \begin{bmatrix} 2 \\ 3 \\ 4 \end{bmatrix} = 2^2 + 3^2 + 4^2 = \|w\|_2^2$$

$$w_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + \lambda \|w\|_2^2 \quad (1)$$

$$w'_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + w^T M w \quad (2)$$

Where,

$X \in \mathbb{R}^{N \times D}$ is the data matrix

$y \in \mathbb{R}^N$ is a vector of all the outcomes

λ is the regularization coefficient

$M \in \mathbb{R}^{D \times D}$ is a symmetric positive definite matrix.

1. Show that the new method is a generalization of the standard L2 regularization by picking a matrix M such that w'_* in Eq. (2) equals w_* in Eq. (1). (2 points)

Answer:

To make Eq. (2) equals to Eq. (1), we pick the matrix:

$$M = \lambda I$$

where I is the D by D identity matrix clearly makes $w^T M w$ equal to $\lambda \|w\|_2^2$.

$$w^T M w = w^T (\lambda I) w = \lambda w^T w = \lambda \|w\|_2^2 \quad \|w\|_2^2 = \sum_{j=1}^D w_j^2 = w^T w$$

Where,

$$w_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + \lambda \|w\|_2^2 \quad (1)$$

$$w'_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + w^T M w \quad (2)$$

$X \in \mathbb{R}^{N \times D}$ is the data matrix

$y \in \mathbb{R}^N$ is a vector of all the outcomes

λ is the regularization coefficient

$M \in \mathbb{R}^{D \times D}$ is a symmetric positive definite matrix.

2. Find the closed form of w'_* by writing down the gradient of $F(w) = \|Xw - y\|_2^2 + w^T M w$ and setting it to 0. (4 points)

$$F(w) = \|Xw - y\|_2^2 + w^T M w$$

$$\|Xw - y\|_2^2 = (Xw - y)^T (Xw - y)$$

$$= ((Xw)^T - y^T) (Xw - y).$$

using $(AB)^T = B^T A^T$

$$= (Xw)^T (Xw) - \underline{(Xw)^T y - y^T (Xw)} + y^T y$$

$$(Xw)^T y \text{ is a scalar: } (N \times D \ D \times 1)^T N \times 1 \rightarrow (N \times 1)^T N \times 1 \rightarrow 1 \times N \ N \times 1 \rightarrow 1 \times 1$$

$$\text{So we have, } (Xw)^T y = ((Xw)^T y)^T = y^T (Xw).$$

$$w_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + \lambda \|w\|_2^2 \quad (1)$$

Where,

$$w'_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + w^T M w \quad (2)$$

$X \in \mathbb{R}^{N \times D}$ is the data matrix

$y \in \mathbb{R}^N$ is a vector of all the outcomes

λ is the regularization coefficient

$M \in \mathbb{R}^{D \times D}$ is a symmetric positive definite matrix.

2. Find the closed form of w'_* by writing down the gradient of $F(w) = \|Xw - y\|_2^2 + w^T M w$ and setting it to 0. (4 points)

$$F(w) = \|Xw - y\|_2^2 + w^T M w$$

$$\|Xw - y\|_2^2 = (Xw - y)^T (Xw - y)$$

$$= ((Xw)^T - y^T) (Xw - y).$$

using $(AB)^T = B^T A^T$

$$= (Xw)^T (Xw) - (Xw)^T y - y^T (Xw) + y^T y$$

$$= (Xw)^T (Xw) - 2y^T (Xw) + y^T y$$

using $(Xw)^T y = y^T (Xw)$

$$= w^T X^T X w - 2y^T X w + y^T y$$

using $(AB)^T = B^T A^T$

$$w_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + \lambda \|w\|_2^2 \quad (1)$$

Where,

$$w'_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + w^T M w \quad (2)$$

$X \in \mathbb{R}^{N \times D}$ is the data matrix

$y \in \mathbb{R}^N$ is a vector of all the outcomes

λ is the regularization coefficient

$M \in \mathbb{R}^{D \times D}$ is a symmetric positive definite matrix.

2. Find the closed form of w'_* by writing down the gradient of $F(w) = \|Xw - y\|_2^2 + w^T M w$ and setting it to 0. (4 points)

$$F(w) = \|Xw - y\|_2^2 + w^T M w$$

$$F(w) = w^T X^T X w - 2 y^T X w + y^T y + w^T M w.$$

The gradient:

$$\nabla_w [w^T X^T X w] = 2X^T X w$$

$$\nabla_w [-2 y^T X w] = -2 X^T y$$

$$\nabla_w [y^T y] = 0$$

$$\nabla_w [w^T M w] = 2Mw$$

$$\frac{\partial \mathbf{x}^T \mathbf{A} \mathbf{x}}{\partial \mathbf{x}} = 2\mathbf{A} \mathbf{x}$$

$$\frac{\partial \mathbf{a}^T \mathbf{x}}{\partial \mathbf{x}} = \mathbf{a}$$

$$y^T X w = (X^T y)^T w$$

$\mathbf{A}$ is not a function of $\mathbf{x}$ and $\mathbf{A}$ is symmetric

$\mathbf{a}$ is not a function of $\mathbf{x}$

$$w_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + \lambda \|w\|_2^2 \quad (1)$$

Where,

$$w'_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + w^T M w \quad (2)$$

$X \in \mathbb{R}^{N \times D}$ is the data matrix

$y \in \mathbb{R}^N$ is a vector of all the outcomes

λ is the regularization coefficient

$M \in \mathbb{R}^{D \times D}$ is a symmetric positive definite matrix.

2. Find the closed form of w'_* by writing down the gradient of $F(w) = \|Xw - y\|_2^2 + w^T M w$ and setting it to 0. (4 points)

$$F(w) = \|Xw - y\|_2^2 + w^T M w$$

$$F(w) = w^T X^T X w - 2 y^T X w + y^T y + w^T M w.$$

The gradient:

$$\nabla F(w) = 2X^T X w - 2X^T y + 2Mw$$

$$2X^T(Xw - y) + 2Mw.$$

$$w_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + \lambda \|w\|_2^2 \quad (1)$$

Where,

$$w'_* = \arg \min_{w \in \mathbb{R}^D} \|Xw - y\|_2^2 + w^T M w \quad (2)$$

$X \in \mathbb{R}^{N \times D}$ is the data matrix

$y \in \mathbb{R}^N$ is a vector of all the outcomes

λ is the regularization coefficient

$M \in \mathbb{R}^{D \times D}$ is a symmetric positive definite matrix.

2. Find the closed form of w'_* by writing down the gradient of $F(w) = \|Xw - y\|_2^2 + w^T M w$ and setting it to 0. (4 points)

Set it to 0:

$$2X^T(Xw - y) + 2Mw.$$

$$2X^T X w - 2X^T y + 2Mw = 0$$

$$(X^T X + M)w = X^T y$$

$$Aw = b$$

Let $A := X^T X + M$ and $b := X^T y$

$$A^{-1}Aw = A^{-1}b$$

$$w = A^{-1}b \text{ , that is, } w'_* = (X^T X + M)^{-1} X^T y.$$

3. Recall the Newton method: $w^{(t+1)} \leftarrow w^{(t)} - H_t^{-1} \nabla F(w^{(t)})$ where $H_t = \nabla^2 F(w^{(t)})$. Show that no matter what the initialization $w^{(0)}$ is, Newton method always takes one step only to find the minimizer w'_* of $F(w) = \|Xw - y\|_2^2 + w^T M w$. (4 points)

$$w^{(t+1)} \leftarrow w^{(t)} - H_t^{-1} \nabla F(w^{(t)})$$

$$\nabla F(w) = 2X^T(Xw - y) + 2Mw$$

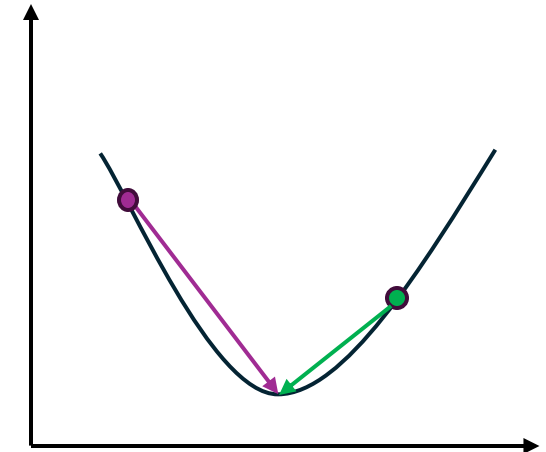
At **any point**, the Hessian is always $H := \nabla^2 F(w) = 2(X^T X + M)$.

$$H^{-1} = [2(X^T X + M)]^{-1} = \frac{1}{2}(X^T X + M)^{-1}$$

So applying Newton method to any initialization of $w^{(0)}$ gives,

$$\begin{aligned} w^{(1)} &= w^{(0)} - (X^T X + M)^{-1} (X^T (Xw^{(0)} - y) + Mw^{(0)}) \\ &= w^{(0)} - (X^T X + M)^{-1} ((X^T X + M)w^{(0)} - X^T y) \\ &= (X^T X + M)^{-1} X^T y = w'_*. \end{aligned}$$

Hence Newton method needs only **one step**, regardless of initialization.



Example: Newton method jumps directly to its minimizer in one step from any start.

2.2 (14 points) Assume we have a training set $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N) \in \mathbb{R}^D \times \mathbb{R}$, where each outcome y_n is generated by a probabilistic model $\mathbf{w}_*^T \mathbf{x}_n + \epsilon_n$ with ϵ_n being an independent Gaussian noise with zero-mean and variance σ^2 for some $\sigma > 0$. In other words, the probability density of any outcome $y \in \mathbb{R}$ given $\mathbf{x}_n$ is

$$\Pr(y \mid \mathbf{x}_n; \mathbf{w}_*, \sigma) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left(\frac{-(y - \mathbf{w}_*^T \mathbf{x}_n)^2}{2\sigma^2}\right).$$

1. Assume σ is fixed and given, find the maximum likelihood estimation for $\mathbf{w}_*$. In other words, first write down the joint density of the outcomes $y_1, \dots, y_N$ given $\mathbf{x}_1, \dots, \mathbf{x}_N$ as a function of the value of $\mathbf{w}_*$; then find the value of $\mathbf{w}_*$ that maximizes this density. You can assume $\mathbf{X}^T \mathbf{X}$ is invertible where $\mathbf{X}$ is the data matrix as used in Problem 2.1. (6 points)

The probability of seeing label $y_1, \dots, y_n$ given $\mathbf{x}_1, \dots, \mathbf{x}_n$ is that

$$\mathcal{P}(\mathbf{w}) = \prod_{n=1}^N \Pr(y_n \mid \mathbf{x}_n; \mathbf{w}, \sigma) = \prod_{n=1}^N \frac{1}{\sigma\sqrt{2\pi}} \exp\left(\frac{-(y_n - \mathbf{w}^T \mathbf{x}_n)^2}{2\sigma^2}\right).$$

(The joint density for a linear model $\mathbf{w}$)

Now, we need to find a $\mathbf{w}_*$ that maximizes $\mathcal{P}(\mathbf{w})$

1. Assume σ is fixed and given, find the maximum likelihood estimation for w_* . In other words, first write down the joint density of the outcomes $y_1, \dots, y_N$ given $x_1, \dots, x_N$ as a function of the value of w_* ; then find the value of w_* that maximizes this density. You can assume $X^T X$ is invertible where X is the data matrix as used in Problem 2.1. (6 points)

$$\mathcal{P}(w) = \prod_{n=1}^N \Pr(y_n | x_n; w, \sigma) = \prod_{n=1}^N \frac{1}{\sigma\sqrt{2\pi}} \exp\left(-\frac{(y_n - w^T x_n)^2}{2\sigma^2}\right).$$

Taking the negative log, this becomes

$$F(w) = -\ln P(w) = -\sum_{n=1}^N \ln \left[\frac{1}{\sigma\sqrt{2\pi}} \exp\left(-\frac{(y_n - w^T x_n)^2}{2\sigma^2}\right) \right]$$

$$\log_b(xy) = \log_b x + \log_b y$$

$$= -\sum_{n=1}^N \left[\ln\left(\frac{1}{\sigma\sqrt{2\pi}}\right) - \frac{(y_n - w^T x_n)^2}{2\sigma^2} \right]$$

$$\ln e^t = t$$

$$= -\sum_{n=1}^N \left[-\ln \sigma - \ln \sqrt{2\pi} - \frac{(y_n - w^T x_n)^2}{2\sigma^2} \right]$$

$$\log_b \sqrt[p]{x} = \frac{\log_b x}{p}$$

$$= \sum_{n=1}^N [\ln \sigma + \ln \sqrt{2\pi}] + \frac{1}{2\sigma^2} \sum_{n=1}^N (y_n - w^T x_n)^2$$

$$\log_b \frac{x}{y} = \log_b x - \log_b y$$

$$F(w) = N \ln \sqrt{2\pi} + N \ln \sigma + \frac{1}{2\sigma^2} \sum_{n=1}^N (y_n - w^T x_n)^2 = N \ln \sqrt{2\pi} + N \ln \sigma + \frac{1}{2\sigma^2} \|Xw - y\|_2^2$$

1. Assume σ is fixed and given, find the maximum likelihood estimation for w_* . In other words, first write down the joint density of the outcomes $y_1, \dots, y_N$ given $x_1, \dots, x_N$ as a function of the value of w_* ; then find the value of w_* that maximizes this density. You can assume $X^T X$ is invertible where X is the data matrix as used in Problem 2.1. (6 points)

$$\mathcal{P}(w) = \prod_{n=1}^N \Pr(y_n \mid x_n; w, \sigma) = \prod_{n=1}^N \frac{1}{\sigma \sqrt{2\pi}} \exp \left(\frac{-(y_n - w^T x_n)^2}{2\sigma^2} \right).$$

$$F(w) = N \ln \sqrt{2\pi} + N \ln \sigma + \frac{1}{2\sigma^2} \sum_{n=1}^N (y_n - w^T x_n)^2 = N \ln \sqrt{2\pi} + N \ln \sigma + \frac{1}{2\sigma^2} \|Xw - y\|_2^2$$

Maximizing $P(w)$ is the same as minimizing $F(w)$, the negative log of $P(w)$:

- σ is fixed and given, so it becomes minimizing the last part of $F(w)$: $\sum_{n=1}^N (y_n - w^T x_n)^2$

(the same objective as for least square regression)

Similarly,

$$\nabla_w \|y - Xw\|_2^2 = 2X^T (Xw - y) = 0$$

$$X^T X w = X^T y$$

$$w_* = (X^T X)^{-1} X^T y.$$

2. Now consider σ as a parameter of the probabilistic model too, that is, the model is specified by both w_* and σ . Find the maximum likelihood estimation for w_* and σ . (8 points)

Follow the same steps from the previous question, we can get $F(w, \sigma)$:

$$F(w, \sigma) = N \ln \sqrt{2\pi} + N \ln \sigma + \frac{1}{2\sigma^2} \|Xw - y\|_2^2.$$

MLE for w_* and σ is the minimizer of the function $F(w, \sigma)$.

Take the derivative over w and set it to 0 to find w_* :

$$\frac{\partial F(w, \sigma)}{\partial w} = \frac{1}{\sigma^2} X^\top (Xw - y) = \frac{1}{\sigma^2} (X^\top X w - X^\top y)$$

$$w_* = (X^\top X)^{-1} X^\top y.$$

(Does not depend on σ)

Take the derivative over σ :

$$\frac{\partial F(w, \sigma)}{\partial \sigma} = \frac{N}{\sigma} - \frac{\|Xw - y\|^2}{\sigma^3}.$$

$$N\sigma^2 - \|Xw - y\|^2 = 0 \quad \sigma > 0.$$

$$\sigma = \sqrt{\frac{\|Xw - y\|^2}{N}}$$

This depends on w , so we need to find w_* first.

2. Now consider σ as a parameter of the probabilistic model too, that is, the model is specified by both $\mathbf{w}_*$ and σ . Find the maximum likelihood estimation for $\mathbf{w}_*$ and σ . (8 points)

$$F(\mathbf{w}, \sigma) = N \ln \sqrt{2\pi} + N \ln \sigma + \frac{1}{2\sigma^2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|_2^2. \quad \sigma > 0.$$

We first fix σ and minimize over $\mathbf{w}$ (the same MLE from the previous question):

$$\mathbf{w}_* = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}.$$

For any $\sigma > 0$,

$$N \ln \sqrt{2\pi} + N \ln \sigma + \frac{1}{2\sigma^2} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|_2^2 \geq N \ln \sqrt{2\pi} + N \ln \sigma + \frac{1}{2\sigma^2} \|\mathbf{X}\mathbf{w}_* - \mathbf{y}\|_2^2$$

Then we plug in $\mathbf{w}_*$ and minimize $F(\mathbf{w}_*, \sigma)$ to find σ :

$$\begin{aligned} \frac{\partial F(\mathbf{w}_*, \sigma)}{\partial \sigma} &= \frac{N}{\sigma} - \frac{1}{\sigma^3} \|\mathbf{X}\mathbf{w}_* - \mathbf{y}\|_2^2 = 0. \\ \sigma &= \frac{1}{\sqrt{N}} \|\mathbf{X}\mathbf{w}_* - \mathbf{y}\|_2 = \frac{1}{\sqrt{N}} \|\mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} - \mathbf{y}\|_2. \end{aligned}$$

Problem 3 Linear Classifiers

(16 points)

In Lecture 3 we have seen the hinge loss $\ell(z) = \max\{0, 1 - z\}$, which is non-differentiable at $z = 1$. To avoid this issue, we can consider the square of hinge loss $\ell(z)^2$, which is differentiable everywhere. More specifically, given a binary dataset $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N) \in \mathbb{R}^D \times \{-1, 1\}$, we define the following new loss function for a linear model $\mathbf{w} \in \mathbb{R}^D$:

$$F(\mathbf{w}) = \frac{1}{N} \sum_{n=1}^N F_n(\mathbf{w}), \quad \text{where } F_n(\mathbf{w}) = \left(\max \left\{ 0, 1 - y_n \mathbf{w}^\top \mathbf{x}_n \right\} \right)^2. \quad (3)$$

1. For a fixed n , write down the gradient $\nabla F_n(\mathbf{w})$ (show your derivation), then fill in the missing details in the repeat-loop of the algorithm below which applies SGD to minimize F . (6 points)

When $1 - y_n \mathbf{w}^\top \mathbf{x}_n < 0$, $F_n(\mathbf{w}) = \left(\max \left\{ 0, 1 - y_n \mathbf{w}^\top \mathbf{x}_n \right\} \right)^2 = 0$ The gradient of it is 0.

When $1 - y_n \mathbf{w}^\top \mathbf{x}_n > 0$, $F_n(\mathbf{w}) = (1 - y_n \mathbf{w}^\top \mathbf{x}_n)^2$.

Take the derivative by using the chain rule:

$$\frac{d}{dz} z^2 = 2z \quad \nabla_{\mathbf{w}} F_n(\mathbf{w}) = 2(1 - y_n \mathbf{w}^\top \mathbf{x}_n) \cdot \nabla_{\mathbf{w}} (1 - y_n \mathbf{w}^\top \mathbf{x}_n) = 2(1 - y_n \mathbf{w}^\top \mathbf{x}_n)(-y_n \mathbf{x}_n)$$

$$\nabla F_n(\mathbf{w}) = -2y_n(1 - y_n \mathbf{w}^\top \mathbf{x}_n)\mathbf{x}_n.$$

Problem 3 Linear Classifiers

(16 points)

In Lecture 3 we have seen the hinge loss $\ell(z) = \max\{0, 1 - z\}$, which is non-differentiable at $z = 1$. To avoid this issue, we can consider the square of hinge loss $\ell(z)^2$, which is differentiable everywhere. More specifically, given a binary dataset $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N) \in \mathbb{R}^D \times \{-1, 1\}$, we define the following new loss function for a linear model $\mathbf{w} \in \mathbb{R}^D$:

$$F(\mathbf{w}) = \frac{1}{N} \sum_{n=1}^N F_n(\mathbf{w}), \quad \text{where } F_n(\mathbf{w}) = \left(\max\{0, 1 - y_n \mathbf{w}^T \mathbf{x}_n\} \right)^2. \quad (3)$$

1. For a fixed n , write down the gradient $\nabla F_n(\mathbf{w})$ (show your derivation), then fill in the missing details in the repeat-loop of the algorithm below which applies SGD to minimize F . (6 points)

When $1 - y_n \mathbf{w}^T \mathbf{x}_n < 0$, $F_n(\mathbf{w}) = \left(\max\{0, 1 - y_n \mathbf{w}^T \mathbf{x}_n\} \right)^2 = 0$ The gradient of it is 0.

When $1 - y_n \mathbf{w}^T \mathbf{x}_n > 0$, $F_n(\mathbf{w}) = (1 - y_n \mathbf{w}^T \mathbf{x}_n)^2$, $\max\{0, 1 - y_n \mathbf{w}^T \mathbf{x}_n\} = \begin{cases} 0 & \text{if } 1 - y_n \mathbf{w}^T \mathbf{x}_n \leq 0, \\ 1 - y_n \mathbf{w}^T \mathbf{x}_n & \text{if } 1 - y_n \mathbf{w}^T \mathbf{x}_n > 0. \end{cases}$
 $\nabla F_n(\mathbf{w}) = -2y_n(1 - y_n \mathbf{w}^T \mathbf{x}_n)\mathbf{x}_n$,

which is also 0 when $y_n \mathbf{w}^T \mathbf{x}_n$ approaches 1

$$\nabla F_n(\mathbf{w}) = -2y_n \max\{0, 1 - y_n \mathbf{w}^T \mathbf{x}_n\} \mathbf{x}_n.$$

Problem 3 Linear Classifiers

(16 points)

In Lecture 3 we have seen the hinge loss $\ell(z) = \max\{0, 1 - z\}$, which is non-differentiable at $z = 1$. To avoid this issue, we can consider the square of hinge loss $\ell(z)^2$, which is differentiable everywhere. More specifically, given a binary dataset $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N) \in \mathbb{R}^D \times \{-1, 1\}$, we define the following new loss function for a linear model $\mathbf{w} \in \mathbb{R}^D$:

$$F(\mathbf{w}) = \frac{1}{N} \sum_{n=1}^N F_n(\mathbf{w}), \quad \text{where } F_n(\mathbf{w}) = \left(\max\{0, 1 - y_n \mathbf{w}^T \mathbf{x}_n\} \right)^2. \quad (3)$$

1. For a fixed n , write down the gradient $\nabla F_n(\mathbf{w})$ (show your derivation), then fill in the missing details in the repeat-loop of the algorithm below which applies SGD to minimize F . (6 points)

Algorithm 1: SGD for minimizing Eq. (3)

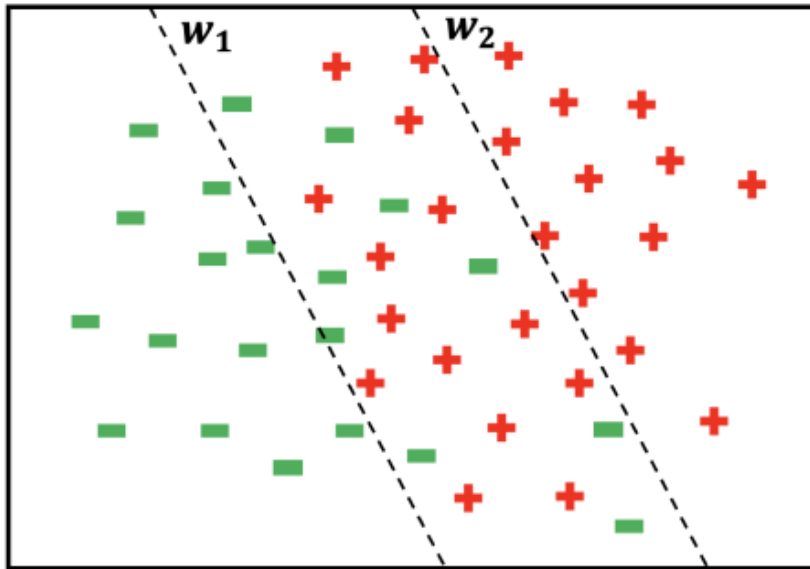
- 1 **Input:** A training set $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N) \in \mathbb{R}^D \times \{-1, 1\}$, learning rate $\eta > 0$
 - 2 **Initialization:** $\mathbf{w} = \mathbf{0}$
 - 3 **Repeat:**
 - 4 randomly pick an example $(\mathbf{x}_n, y_n)$
 - 5 update $\mathbf{w} \leftarrow \mathbf{w} + 2\eta y_n \max\{0, 1 - y_n \mathbf{w}^T \mathbf{x}_n\} \mathbf{x}_n$
-

$$\nabla F_n(\mathbf{w}) = -2y_n \max\{0, 1 - y_n \mathbf{w}^T \mathbf{x}_n\} \mathbf{x}_n.$$

2. Next, consider modifying $F_n(\mathbf{w})$ as

$$F_n(\mathbf{w}) = \begin{cases} (\max\{0, 1 - \mathbf{w}^\top \mathbf{x}_n\})^2, & \text{if } y_n = 1, \\ \underline{0.1} (\max\{0, 1 + \mathbf{w}^\top \mathbf{x}_n\})^2, & \text{else.} \end{cases} \quad (4)$$

- (a) Consider a binary classification dataset of points in two dimensions as shown in Figure 1, where the red, plus signs denote samples with label $+1$, and the green, minus signs denote samples with label -1 . When training a linear classifier with the modified loss in Eq. (4), which of w_1 or w_2 in Figure 1 do you think is more likely the resulting decision boundary? Explain briefly. (2 points)



- The loss function now does not penalize misclassifying negative points as heavily as it penalizes misclassifying positive points.

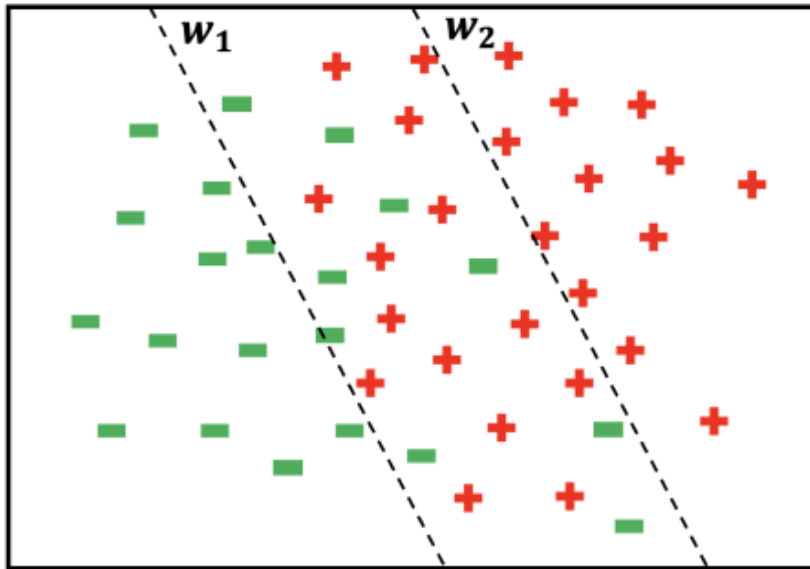
=> The model does not want to misclassify $+$ points.

Figure 1: A binary classification task

2. Next, consider modifying $F_n(\mathbf{w})$ as

$$F_n(\mathbf{w}) = \begin{cases} (\max \{0, 1 - \mathbf{w}^\top \mathbf{x}_n\})^2, & \text{if } y_n = 1, \\ \underline{0.1} (\max \{0, 1 + \mathbf{w}^\top \mathbf{x}_n\})^2, & \text{else.} \end{cases} \quad (4)$$

- (a) Consider a binary classification dataset of points in two dimensions as shown in Figure 1, where the red, plus signs denote samples with label $+1$, and the green, minus signs denote samples with label -1 . When training a linear classifier with the modified loss in Eq. (4), which of $\mathbf{w}_1$ or $\mathbf{w}_2$ in Figure 1 do you think is more likely the resulting decision boundary? Explain briefly. (2 points)



- The dataset is not linearly separable.
 - Any linear classifier will make mistakes.
- Therefore, $\mathbf{w}_1$ is more likely.

Figure 1: A binary classification task

2. Next, consider modifying $F_n(\mathbf{w})$ as

$$F_n(\mathbf{w}) = \begin{cases} (\max\{0, 1 - \mathbf{w}^\top \mathbf{x}_n\})^2, & \text{if } y_n = 1, \\ \underline{0.1} (\max\{0, 1 + \mathbf{w}^\top \mathbf{x}_n\})^2, & \text{else.} \end{cases} \quad (4)$$

(b) Based on your answer from the last question, give an example where one would want to modify the loss function in such a way. (2 points)

For example, in fraud detection (+1 represents a fraud), it is far more important to make sure that an actual fraud is not missed, and thus the loss function should give more weights to positive labels.

We accept any reasonable example where the one label represents a critical or high-stakes outcome and is far more important than the other label.

2. Next, consider modifying $F_n(w)$ as

$$F_n(w) = \begin{cases} (\max\{0, 1 - w^\top x_n\})^2, & \text{if } y_n = 1, \\ 0.1 (\max\{0, 1 + w^\top x_n\})^2, & \text{else.} \end{cases} \quad (4)$$

(c) Similarly to Question 3.1, write down the gradient of this modified loss F_n , then fill in the missing details in the repeat-loop of the algorithm below which applies SGD to minimize F . (6 points)

Similar steps as in Question 3.1,

When $y_n = 1$,

- If $1 - w^\top x_n \leq 0$, the gradient is 0.
- If $1 - w^\top x_n > 0$,

$$\nabla F_n(w) = 2(1 - w^\top x_n) \nabla(1 - w^\top x_n) = 2(1 - w^\top x_n)(-x_n) = -2 \max\{0, 1 - w^\top x_n\} x_n$$

Else,

- If $1 + w^\top x_n \leq 0$, the gradient is 0.
- If $1 + w^\top x_n > 0$,

$$\nabla F_n(w) = 0.1 \cdot 2(1 + w^\top x_n) \nabla(1 + w^\top x_n) = 0.2 \max\{0, 1 + w^\top x_n\} x_n.$$

2. Next, consider modifying $F_n(\mathbf{w})$ as

$$F_n(\mathbf{w}) = \begin{cases} (\max\{0, 1 - \mathbf{w}^\top \mathbf{x}_n\})^2, & \text{if } y_n = 1, \\ 0.1 (\max\{0, 1 + \mathbf{w}^\top \mathbf{x}_n\})^2, & \text{else.} \end{cases} \quad (4)$$

(c) Similarly to Question 3.1, write down the gradient of this modified loss F_n , then fill in the missing details in the repeat-loop of the algorithm below which applies SGD to minimize F . (6 points)

$$\nabla F_n(\mathbf{w}) = \begin{cases} -2 \max\{0, 1 - \mathbf{w}^\top \mathbf{x}_n\} \mathbf{x}_n, & \text{if } y_n = 1, \\ 0.2 \max\{0, 1 + \mathbf{w}^\top \mathbf{x}_n\} \mathbf{x}_n & \text{else.} \end{cases}$$

Algorithm 2: SGD for minimizing modified loss Eq. (4)

1 **Input:** A training set $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N) \in \mathbb{R}^D \times \{-1, 1\}$, learning rate $\eta > 0$

2 **Initialization:** $\mathbf{w} = \mathbf{0}$

3 **Repeat:**

4 randomly pick an example $(\mathbf{x}_n, y_n)$

5 update $\mathbf{w} \leftarrow \mathbf{w} - \begin{cases} -2\eta \max\{0, 1 - \mathbf{w}^\top \mathbf{x}_n\} \mathbf{x}_n, & \text{if } y_n = 1, \\ 0.2\eta \max\{0, 1 + \mathbf{w}^\top \mathbf{x}_n\} \mathbf{x}_n & \text{else.} \end{cases}$
