

DSC 140A Lecture

Lecture 1:

Feature : numerical

Feature vector

Dimensionality : # of features in a feature vector

Training set

$$X = \{(\vec{x}^{(1)}, y_1), \dots, (\vec{x}^{(n)}, y_n)\}$$

Train error

Test error $\Rightarrow$ generalization

Overfitting

Underfitting

Nearest Neighbors Predictor

useful for both classification & regression

Predict

Find the closest point to $\vec{z}$ in X :

$$i^* = \arg \min_{i \in \{1, \dots, n\}} \|\vec{x}^{(i)} - \vec{z}\|$$

Predicted label : y_i

• Euclidean Distance - 2-norm

$$\|\vec{p} - \vec{q}\| = \sqrt{\sum_{k=1}^d (p_k - q_k)^2} = \sqrt{(\vec{p} - \vec{q}) \cdot (\vec{p} - \vec{q})}$$

• Using Euclidean Distance treats all features equally

⇓ fix

Standardizing features (z-score).

$$(x_1, x_2)^T \Rightarrow (z_1, z_2)^T = \left(\frac{x_1 - \mu_1}{\sigma_1}, \frac{x_2 - \mu_2}{\sigma_2} \right)$$

• Decision Boundary

• KNN

• choose k nearest neighbors and vote

predict

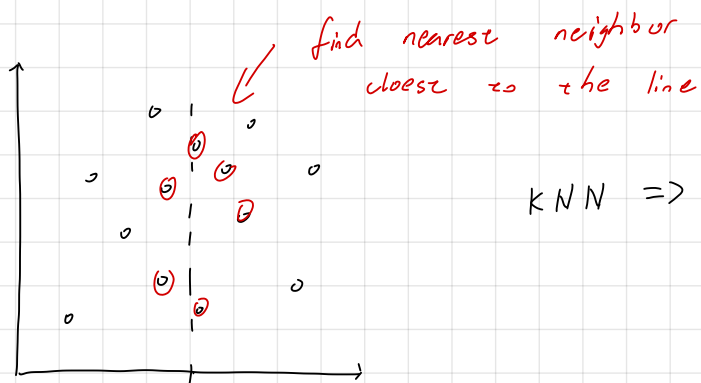
• find the most common labels in the k nearest neighbors points.

• k controls complexity

$k \uparrow \rightarrow$ underfitting $\rightarrow$ simple boundary

$k \downarrow \rightarrow$ overfitting $\rightarrow$ complex boundary

NN Regression



kNN $\Rightarrow$ average the neighbor

Empirical Risk Minimization

Prediction Function

$H(\vec{x}) \rightarrow y$

↑
takes in
a feature vector

↑
predicted label.

ex) $H(\text{experience}, \text{GPA}) = 50000 + 10000 \times \text{exp.} + 5000 \times \text{GPA}$

Assumption: the future will be like the past

↑
unseen

idea: use past data to "measure" how good a prediction function is.

Fit: formally quantify how well a prediction function fits the data?

⇓

by measuring error.

• loss function $\rightarrow$ measure difference between $H(\vec{x}^{(i)})$ & y_i

• Absolute loss: $\mathcal{L}_{\text{abs}}(H(\vec{x}^{(i)}), y_i) = |H(\vec{x}^{(i)}) - y_i|$

• Square loss: $\mathcal{L}_{\text{sq}}(H(\vec{x}^{(i)}), y_i) = (H(\vec{x}^{(i)}) - y_i)^2$

• Average loss $\Rightarrow$ Empirical Risk

$$\underline{R(H) = \frac{1}{n} \sum_{i=1}^n \mathcal{L}(H(\vec{x}^{(i)}), y_i)}$$

• depend on training set
& choice of loss function.

⇓

Minimizing Empirical Risk.

Restricting our search for prediction functions to a smaller set of functions.

↓
hypothesis class.

So, ERM

ex) linear function

1. hypothesis class
2. loss function
3. find H minimizing risk

Linear Function

$$H(x) = w_0 + w_1 x$$

Generally, $H(\vec{x}) = w_0 + w_1 x_1 + w_2 x_2 + \dots + w_d x_d.$

$w_0 \dots w_d$ are the parameters or weights.

- $w_0 \rightarrow$ bias, determine the prediction when all other are 0.
- $w_d \rightarrow$ how much prediction change
- parameter vector: $\vec{x} \in \mathbb{R}^d \Rightarrow \vec{w} \in \mathbb{R}^{d+1}$

$$\vec{w} = (w_0, w_1, \dots, w_d)^T$$

$$\text{So, } H(\vec{x}, \vec{w}) = (w_0, \dots, w_d)^T (1, x_1, \dots, x_d)^T$$

$$\underline{\text{Aug}(\vec{x}) = (1, x_1, \dots, x_d)^T = (1, \vec{x})^T}$$

$$\text{So, } H(\vec{x}, \vec{w}) = w_0 + w_1 x_1 + \dots + w_d x_d \\ = \vec{w} \cdot \text{Aug}(\vec{x}).$$

$$\text{Absolute loss: } R_{\text{abs}}(\vec{w}) = \frac{1}{n} \sum_{i=1}^n |H(\vec{x}^{(i)}) - y_i| \\ = \frac{1}{n} \sum_{i=1}^n |\vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) - y_i|$$

Risk Surface: imagine plotting $R_{\text{abs}}(\vec{w})$ for all values of $\vec{w}$

$\vec{w}$ that makes the surface lowest minimizes the risk

↓ How?

- Calculus $\rightarrow$ gradient $\frac{d}{d\vec{w}} R_{\text{abs}}(\vec{w})$
 set it equal to 0
 solve for $\vec{w}$

- However, R_{abs} is not differentiable

↓ try square loss

$$R_{\text{sq}}(\vec{w}) = \frac{1}{n} \sum_{i=1}^n (H(\vec{x}^{(i)}) - y_i)^2 \\ = \frac{1}{n} \sum_{i=1}^n (\vec{w}^{(i)} \cdot \text{Aug}(\vec{x}^{(i)}) - y_i)^2$$

(Mean Squared Error)
(MSE).

Least Square:

$$R_{\text{sq}} = \frac{1}{n} \sum_{i=1}^n (\vec{w}^{(i)} \cdot \text{Aug}(\vec{x}^{(i)}) - y_i)^2$$

$$= \frac{1}{n} \|X\vec{w} - \vec{y}\|^2 \quad \leftarrow \begin{array}{l} \text{let } p_i = \text{Aug}(\vec{x}^{(i)}) \cdot \vec{w} \\ \vec{p} = (p_1, \dots, p_n)^T \\ \vec{p} \in \mathbb{R}^n \end{array}$$

↑
design matrix

Design Matrix

$$X = \begin{pmatrix} \text{Aug}(\vec{x}^{(1)}) \\ \vdots \\ \text{Aug}(\vec{x}^{(n)}) \end{pmatrix}$$

$$= \begin{pmatrix} 1 & x_1^{(1)} & \dots & x_d^{(1)} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_1^{(n)} & \dots & x_d^{(n)} \end{pmatrix} \in \mathbb{R}^{n \times d+1}$$

$$\text{So } \vec{p} = X \vec{w}$$

$$K_{sq} = \frac{1}{n} \sum_{i=1}^n (p_i - y_i)^2$$

$$= \frac{1}{n} \|\vec{p} - \vec{y}\|^2$$

$$\Downarrow$$

$$K_{sq} = \frac{1}{n} \|X \vec{w} - \vec{y}\|^2$$

Find gradient, $\frac{d}{d\vec{w}} \left[\frac{1}{n} \|X \vec{w} - \vec{y}\|^2 \right]$

$\Downarrow$

$$\frac{dK_{sq}}{d\vec{w}} = \left(\frac{\partial K_{sq}}{\partial w_0}, \dots \right) \quad \text{partial derivative}$$

$$K_{sq} = (X \vec{w} - \vec{y}) \cdot (X \vec{w} - \vec{y})$$

$$= (X \vec{w} - \vec{y})^T (X \vec{w} - \vec{y})$$

$$= (\vec{w}^T X^T - \vec{y}^T) (X \vec{w} - \vec{y})$$

$$= \vec{w}^T X^T X \vec{w} - \vec{w}^T X^T \vec{y} - \vec{y}^T X \vec{w} + \vec{y}^T \vec{y}$$

$$\Downarrow$$

$$\vec{w}^T X^T \vec{y} = (\vec{y}^T X \vec{w})^T \leftarrow \text{scalar}$$

$$= \vec{y}^T X \vec{w}$$

$$= \vec{w}^T X^T X \vec{w} - 2 \vec{y}^T X \vec{w} + \vec{y}^T \vec{y}$$

Find gradient: $\frac{d}{d\vec{w}} \left[\vec{w}^T X^T X \vec{w} - 2 \vec{y}^T X \vec{w} + \vec{y}^T \vec{y} \right]$

$\Downarrow$

$$\frac{d}{d\vec{w}} \left((X \vec{w})^T (X \vec{w}) \right) \sim \frac{d}{d\vec{w}} (x^2 w^2) = 2 x^2 w$$

$$= 2 X^T X \vec{w}$$

$$\frac{d}{d\vec{w}} [Y^T X \vec{w}] = X^T \vec{y}$$

$$\frac{d}{d\vec{w}} [Y^T Y] = 0$$

$$S_0, = \frac{1}{n} (2 X^T X \vec{w} - 2 X^T \vec{y} + 0)$$

$$\text{set } \frac{dK_{sq}}{d\vec{w}} = 0$$

$$\frac{1}{n} (2 X^T X \vec{w} - 2 X^T \vec{y}) = 0$$

$$X^T X \vec{w} = X^T \vec{y} \quad (\text{normal equation})$$

$\Downarrow$

$$\vec{w}^* = (X^T X)^{-1} X^T \vec{y} \quad (\text{least square solution})$$

direct solution.

Gradient Descent

derivative $\rightarrow$ as slope of curve

$\hookrightarrow$ Positive: increasing

Negative: decreasing

Zero: flat

• Partial derivative.

for f as function of $\vec{z} = (z_1, z_2)^T$

• $\frac{\partial f}{\partial z_1}$: slope in z_1 direction

• $\frac{\partial f}{\partial z_2}$: slope in z_2 direction.

• Gradient

$$\frac{df}{d\vec{z}}(\vec{z}) = \begin{pmatrix} \frac{\partial f}{\partial z_1}(\vec{z}) \\ \vdots \\ \frac{\partial f}{\partial z_n}(\vec{z}) \end{pmatrix}$$

for $f: \mathbb{R}^d \rightarrow \mathbb{R}$

(the i th component is slope of f at $\vec{z}$ in i th direction).

• input vector $\in \mathbb{R}^d$

• output vector $\in \mathbb{R}^d$

• gradient approximation:

$$f(\vec{z} + \vec{\delta}) \approx f(\vec{z}) + \vec{\delta} \cdot \frac{df}{d\vec{z}}(\vec{z}).$$

↑
change in $\vec{z}$

↓
slope for $\vec{z}$

• gradient in the direction of steepest ascent.

⇓

negative

(steepest descent)

↓
why?

$$\text{change } \vec{\delta} \cdot \frac{df}{d\vec{z}}(\vec{z}) = \|\vec{\delta}\| \left\| \frac{df}{d\vec{z}} \right\| \cos \theta$$

↓

maximized
when $\theta = 0$

↓

direction of gradient

• Contour map

↓

gradient orthogonal to the contours

• Gradient descent (optimization).

⇓

find the minimum

• idea: iterate towards a minimum step by step

start at arbitrary location

move in direction of steepest descent. (negative gradient).

pick arbitrary point $\vec{z}^{(0)}$, learning rate $\eta > 0$

· compute gradient: $\frac{df}{d\vec{z}}(\vec{z}^{(t)})$ at $\vec{z}^{(t)}$

· update rule

$$\vec{z}^{(t+1)} = \vec{z}^{(t)} - \eta \times \frac{df}{d\vec{z}}(\vec{z}^{(t)}) \quad \text{for } \eta \text{ as learning rate}$$

· when converged, return $\vec{z}^{(t)}$

⇓
local minimum.

stop when $\|\frac{df}{d\vec{z}}(\vec{z})\|$ is small

or $\|\vec{z}^{(t+1)} - \vec{z}^{(t)}\|$ is small.

· Gradient descent with ERM:

$$\begin{aligned} \frac{dK}{d\vec{w}}(\vec{w}) &= \frac{d}{d\vec{w}} \left(\frac{1}{n} \sum_{i=1}^n \ell(\text{Aug}(\vec{x}^{(ii)}) \cdot \vec{w}, y_i) \right) \\ &= \frac{1}{n} \sum_{i=1}^n \frac{d\ell}{d\vec{w}}(\text{Aug}(\vec{x}^{(ii)}) \cdot \vec{w}, y_i) \end{aligned}$$

$d+1 \times 1$

$n \times d+1$

$d+1 \times 1$

for squared loss:

$$\frac{dK}{d\vec{w}}(\vec{w}) = \frac{2}{n} \sum_{i=1}^n (\text{Aug}(\vec{x}^{(ii)}) \cdot \vec{w} - y_i) \text{Aug}(\vec{x}^{(ii)})$$

So update rule:

$$\vec{w}^{(t+1)} = \vec{w}^{(t)} - \eta \times \frac{2}{n} \sum_{i=1}^n (\text{Aug}(\vec{x}^{(ii)}) \cdot \vec{w}^{(t)} - y_i) \text{Aug}(\vec{x}^{(ii)})$$

Benefit of gradient descent vs. normal equation

computation of gradient can be distributed, don't have to compute all at once;
memory efficient.

SGD

idea Use a random smaller subset of data B , to compute gradient.

$$\frac{d}{d\vec{w}} R(\vec{w}) \approx \frac{1}{|B|} \sum_{i \in B} \frac{d}{d\vec{w}} \ell(H(\vec{x}^{(i)}; \vec{w}), y_i).$$

• the smaller set B is called mini-batch

• it is called stochastic gradient.

↓

an approximation of the full gradient

bootstrapping

↓

• process same as gradient descent

except gradient is computed randomly (mini-batch)

Least Square gradient

$$\begin{aligned} \vec{g} &= \frac{2}{m} \sum_{i \in B} (\text{Aug}(\vec{x}^{(i)}) \cdot \vec{w} - y_i) \text{Aug}(\vec{x}^{(i)}) \\ &= \frac{2}{m} X_B^T (X_B \vec{w} - y_B) \end{aligned}$$

Trade off

• each step of GD, move in the "best" direction
but slowly

• each step of SGD, move in a "good" direction,
but faster.

useful for massive datasets

Absolute loss at some point

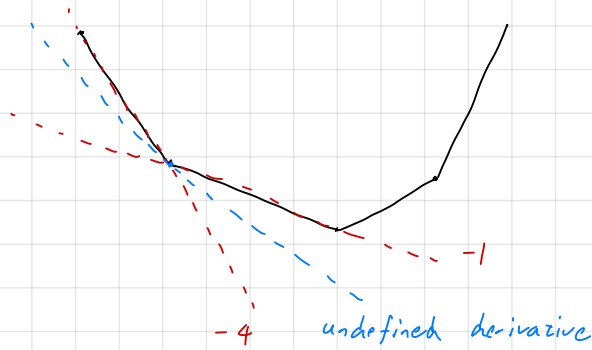
not differentiable $\checkmark \Rightarrow$ subgradient descent
(piecewise function)

↓↓

problem: reach a point where the gradient is not defined

⇓
a replacement

idea:



⇓
replace with subderivative

(any number between -4 and -1)

Any valid subderivative func.
defines a line lies below
the function.

⇓

$$f_s(z) = f(z_0) + s(z - z_0)$$

s is a subderivative if

$$f(z) \geq f_s(z) \text{ for all } z,$$

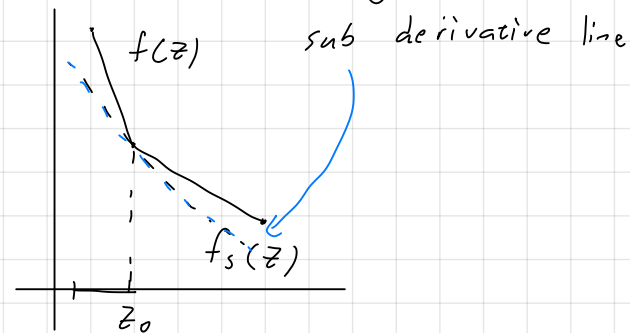
$$f(z) \geq f(z_0) + s(z - z_0).$$

⇓

subgradient:

→ in all dimension

valid subgradient if lies below or
at the function



$$f(\vec{z}) \geq f(\vec{z}^{(0)}) + \vec{s} \cdot (\vec{z} - \vec{z}^{(0)})$$

subgradient descent

problem: might not converge



b/c same derivative

⇓

decrease learning rate in
each iteration



⇓

choose a decreasing learning rate

schedule $\eta(t) > 0$

choose $\eta(t) = c/\sqrt{t}$, t is iteration, c is constant

Absolute loss subgradient descent

Same as GD except use subgradient instead of gradient.

⇓

sub gradient empirical risk

$$\text{subgrad } R(\vec{w}) = \frac{1}{n} \sum_{i=1}^n \text{subgrad } \ell(\vec{w}; \vec{x}^{(i)}, y_i)$$

$$\text{So, } \ell_{\text{abs}}(\vec{w} \cdot \text{Aug}(\vec{x}^{(i)}), y_i) = \begin{cases} \vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) - y_i & \text{if } \vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) > y_i \\ y_i - \vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) & \text{if } \vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) < y_i \\ 0 & \text{if } \vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) = y_i \end{cases}$$

⇓ subgradient

$$s(\vec{w}; \vec{x}^{(i)}, y_i) = \begin{cases} \text{Aug}(\vec{x}^{(i)}) & \text{if } \vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) > y_i \\ -\text{Aug}(\vec{x}^{(i)}) & \text{if } \vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) < y_i \\ \vec{0} & \text{if } \vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) = y_i \end{cases}$$

⇓ empirical risk

$$\frac{1}{n} \sum_{i=1}^n s(\vec{w}, \vec{x}^{(i)}, y_i)$$

zero vector as subgradient.

⇓
use this for GD.

$$\vec{w}^{(t+1)} = \vec{w}^{(t)} - \underbrace{\eta(t)}_{\uparrow} \times \frac{1}{n} \sum_{i=1}^n s(\cdot)$$

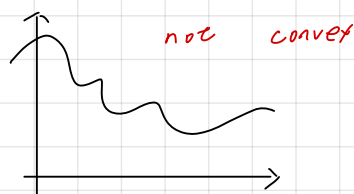
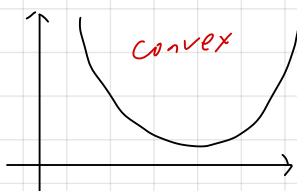
decreasing learning rate.
not fixed.

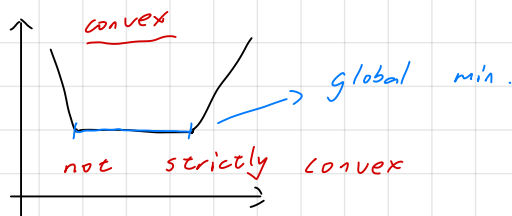


result

median regression, robust to outliers.

Convexity



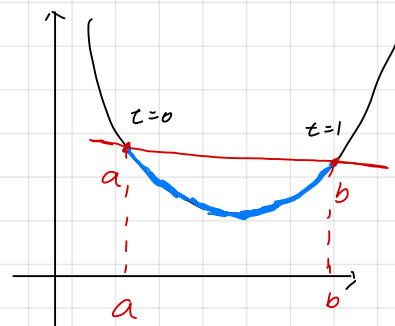


① Def: Function $f: \mathbb{R} \rightarrow \mathbb{R}$ is convex if for every choice of $a, b \in \mathbb{R}$ and $t \in [0, 1]$:

$$(1-t)f(a) + tf(b) \geq f((1-t)a + tb).$$

↓
the line connected
from a to b

↓
the function curve

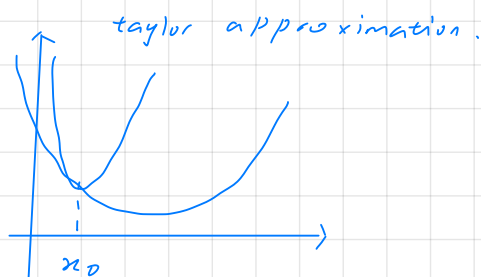


② if $\frac{d^2f}{dx^2}(x) \geq 0$ for all x , then f is convex.

↓
concave up

only if f is twice differentiable.

why?



"best" parabola at x_0

$$f_2(x) = f(x_0) + f'(x_0) \cdot (x - x_0) + \frac{1}{2} f''(x_0) \cdot (x - x_0)^2 \quad \leftarrow \text{Taylor expansion}$$

possibilities: upward, downward, flat.

convex if for every x_0 , parabola is upward-facing

⇓
 $f''(x_0) \geq 0.$

Convex Property

Suppose $f(x)$ & $g(x)$ are convex.

① $w_1 f(x) + w_2 g(x)$ is convex, provided $w_1, w_2 \geq 0$.

② $g(f(x))$ is convex, provided g is non-decreasing.

③ $\max \{f(x), g(x)\}$ is convex.

Theorem $f(x)$ convex, GD converges to global optimum of f .
given that step size is small enough.

Convexity for more dimension

Def: $f: \mathbb{R}^d \rightarrow \mathbb{R}$ is convex if for every choice of $\vec{a}, \vec{b} \in \mathbb{R}^d$,
 $t \in [0, 1]$.

$$(1-t)f(\vec{a}) + tf(\vec{b}) \geq f((1-t)\vec{a} + t\vec{b}).$$

Second Derivative in d -dimension:

there are d^2 of second derivatives.

$\Downarrow$ Hessian Matrix.

For $f: \mathbb{R}^d \rightarrow \mathbb{R}$,

$$H(\vec{x}) = \begin{pmatrix} \frac{\partial^2 f}{\partial x_1^2}(\vec{x}) & \dots & \frac{\partial^2 f}{\partial x_1 \partial x_d}(\vec{x}) \\ \frac{\partial^2 f}{\partial x_2 \partial x_1} & \dots & \frac{\partial^2 f}{\partial x_2 \partial x_d}(\vec{x}) \\ \vdots & \ddots & \vdots \\ \frac{\partial^2 f}{\partial x_d \partial x_1} & \dots & \frac{\partial^2 f}{\partial x_d^2}(\vec{x}) \end{pmatrix} \in \mathbb{R}^{d \times d}$$

$\Downarrow$

Def: convex if all eigenvalues of the Hessian matrix $H(\vec{x})$ are ≥ 0 .

simple case: matrix is diagonal $\Rightarrow$ diagonal entries are the eigenvalues.

Property:

① $w_1 f(\vec{x}) + w_2 g(\vec{x})$ is convex, for $w_1, w_2 \geq 0$.

② Let A be matrix, $\vec{x}, \vec{b}$ vector,

$g(\vec{x}) = f(A\vec{x} + \vec{b})$ is convex as a function of $\vec{x}$,
for $f(x)$ convex

ex) $f(\vec{w}) = (\vec{x} \cdot \vec{w} - \vec{y})^2$ convex

since $f(x) = x^2$ convex.

Linear Classifier

Problem: output of $H(\vec{x})$ is number

$$H(\vec{x}) = w_0 + w_1 x_1 + w_2 x_2.$$

$\Downarrow$

Encoding

$$\left. \begin{array}{l} y \in \{0, 1\} \begin{cases} y=0 \\ y=1 \end{cases} \\ y \in \{-1, 1\} \begin{cases} y=1 \\ y=-1 \end{cases} \end{array} \right\} \text{sign}(z) = \begin{cases} 1 & \text{if } z > 0 \\ -1 & \text{if } z < 0 \\ 0 & \text{if } z = 0 \end{cases}$$

$\Downarrow$
 $\text{sign}(H(\vec{x}, \vec{w}))$

linear classifier / linear decision function.

idea: a weighted vote

each term $w_i x_i$ votes on the label.

• $H(\vec{x})$ is a (hyper)plane.

• $H(\vec{x}) = 0$ is the decision boundary

if $\vec{x} \in \mathbb{R}^2$, a straight line

if $\vec{x} \in \mathbb{R}^d$, a $d-1$ dimensional (hyper)plane.

• Magnitude of $H(\vec{x})$ is proportional to the distance from the decision boundary.

• let $\vec{w}' = (w_1, \dots, w_d)$.

The decision boundary of $H(\vec{x}) = \text{Aug}(\vec{x}) \cdot \vec{w}$ is orthogonal to $\vec{w}'$.

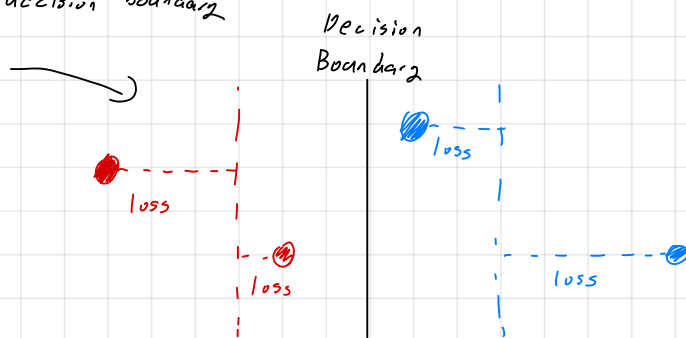
Least Square Classifier

train: normal equation $\vec{w}^* = (X^T X)^{-1} X^T \vec{y}$.

predict: $\text{sign}(\text{Aug}(\vec{x}^T) \cdot \vec{w}^*)$.

problem: penalizes points that are far from decision boundary even if they are correctly classified.

Outlier!



0-1 Loss

$$l_{0-1}(H(\vec{x}^{(i)}), y_i) = \begin{cases} 0 & \text{if } \text{sign}(H(\vec{x}^{(i)})) = y_i \\ 1 & \text{if } \text{sign}(H(\vec{x}^{(i)})) \neq y_i \end{cases}$$

$$R_{0-1}(H) = \frac{1}{n} \sum_{i=1}^n \begin{cases} 0 & \text{if } \dots = y_i \\ 1 & \text{if } \dots \neq y_i \end{cases}$$

Empirical risk

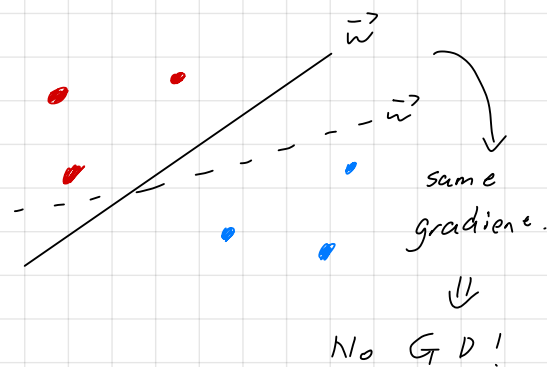


misclassification rate = $\frac{\# \text{ of correct prediction}}{n}$

minimize risk $\Rightarrow$ maximize accuracy

problem: gradient is $\vec{0}$ almost everywhere.

problem:



NP-Hard to optimize expected 0-1 loss.

Surrogate loss

$$l_{\text{perceptron}}(H(\vec{x}), y) = \begin{cases} 0 & \text{if } \text{sign}(H(\vec{x})) = y \\ |H(\vec{x})| & \text{if } \text{sign}(H(\vec{x})) \neq y \end{cases}$$

proportional to d from decision boundary.

$\text{sign}(H(\vec{x})) = y$

$\text{sign}(H(\vec{x})) \neq y$

perceptron loss



penalty grows with distance from decision boundary as point misclassified.

$$\text{Leron}(\text{Aug}(\vec{x}) \cdot \vec{w}, y) = \max(0, -y \text{Aug}(\vec{x}) \cdot \vec{w})$$

$\Downarrow$

Convex.

minimize empirical risk.

$\Downarrow$

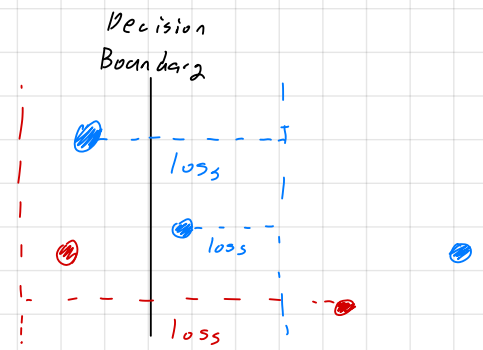
subgradient descent.

$\downarrow$

$$\text{if } -y \text{Aug}(\vec{x}) \cdot \vec{w} > 0, \text{ grad} = -y \text{Aug}(\vec{x})$$

$$-y \text{Aug}(\vec{x}) \cdot \vec{w} < 0, \text{ grad} = \vec{0}$$

$$-y \text{Aug}(\vec{x}) \cdot \vec{w} = 0, \text{ subgrad} = \vec{0}.$$



$$\text{subgrad Leron}(\text{Aug}(\vec{x}) \cdot \vec{w}, y).$$

$$= \begin{cases} \vec{0} & \text{if } \text{sign}(\text{Aug}(\vec{x}) \cdot \vec{w}) = y \\ -y \text{Aug}(\vec{x}) & \text{if } \text{sign}(\text{Aug}(\vec{x}) \cdot \vec{w}) \neq y. \end{cases}$$

train:

$$\vec{w}^{(t+1)} = \vec{w}^{(t)} - \eta(t) \times \frac{1}{n} \sum_{i=1}^n \begin{cases} \vec{0} \\ -y_i \text{Aug}(\vec{x}^{(i)}) \end{cases}$$

$$\text{sign}(\text{Aug}(\vec{x}^{(i)}) \cdot \vec{w}) = y_i$$

$$\text{sign}(\text{Aug}(\vec{x}^{(i)}) \cdot \vec{w}) \neq y_i$$

predict:

$$\text{sign}(\text{Aug}(\vec{x}) \cdot \vec{w}^*).$$

Maximum Margin Classifiers.

Problem with perceptron: no penalty to points correctly classified

$\Downarrow$

overfitting

Def: linearly separable — if there exists a linear classifier which perfectly classifies the data.

Def: Margin — smallest distance between decision boundary and a training point.

1)

Find classifier that maximize margin

2)

better generalization.

Constrain to ensure points sufficiently far from the boundary.

$$|H(\vec{x}^{(i)})| \geq 1.$$

3)

ensure that no point is within the "exclusion zone".

Problem: this only makes the prediction plane very steep.

4)

by making $w_0, \dots, w_n$ very large.

↓

smaller margin

Not what we want.

5)

idea: $\|\vec{w}\|$ as small as possible.

out of all $\vec{w}$ satisfying $y_i \vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) \geq 1$,

for all data points.

find that minimize $\|\vec{w}\|$.

$$\vec{w}^* = \underset{\vec{w}}{\text{argmin}} \|\vec{w}\| \quad \Leftarrow \text{How?}$$

convex quadratic optimization.

$$\forall i, y_i \vec{w} \cdot \text{Aug}(\vec{x}) \geq 1.$$

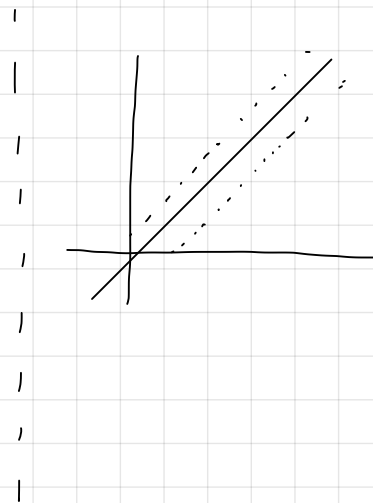
(Hard Support Vector Machine).

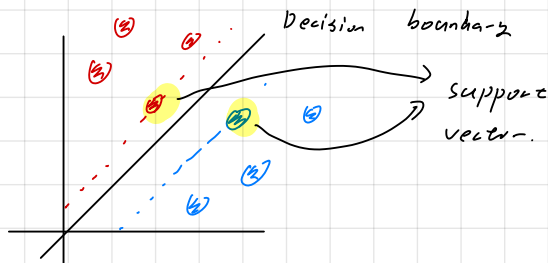
6)

Maximize margin.

• Support Vector - a training point $\vec{x}^{(i)}$, such that

$$y_i \vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) = 1.$$





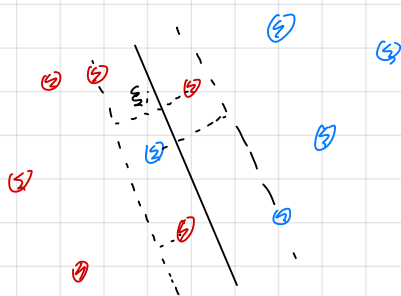
(The solution to Hard-SVM is always a linear combination of the support vector).

Let S be the set of support vectors.

$$\vec{w}^* = \sum_{i \in S} y_i a_i \text{Aug}(\vec{x}^{(i)}).$$

Problem: what if non-separability?

idea: allow some classifications to be ξ_i wrong but not too wrong.



misclassification

For $C \geq 0$.

$$\min_{\vec{w} \in \mathbb{R}^{d+1}, \vec{\xi} \in \mathbb{R}} \|\vec{w}\|^2 + C \sum_{i=1}^n \xi_i$$

Large C : don't give much slack, avoid
Small C : allow more slack at the cost of misclassification.

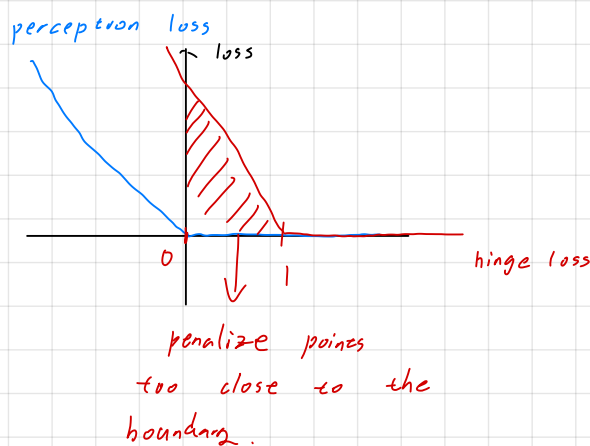
subject to $y_i \vec{w} \cdot \text{Aug}(\vec{x}^{(i)}) \geq 1 - \xi_i$ for all i , $\xi_i \geq 0$.

(Soft margin SVM), As $C \rightarrow \infty$, margin hardens.

Hinge loss

$$\ell_{\text{hinge}} = \begin{cases} 0 & y H(\vec{x}) \geq 1 \\ 1 - y H(\vec{x}) & y H(\vec{x}) < 1. \end{cases}$$

$$= \max\{0, 1 - y H(\vec{x})\}.$$



Thus, soft-SVM $\Leftrightarrow$ finding $\vec{w}$ that minimizes

$$R_{\text{svm}}(\vec{w}) = \|\vec{w}\|^2 + C \sum_{i=1}^n \max\{0, 1 - y_i \vec{w} \cdot \vec{x}^{(i)}\}.$$

SGD to train

a regularized risk

Feature Maps

Non-linear feature maps



careful about overfitting.

allow us to

- fit complex, non-linear patterns
- while still using linear models.

Feature Map: $\vec{\phi} : \mathbb{R}^d \rightarrow \mathbb{R}^k$ function that takes in a d-dimensional vector and outputs a k-dimensional feature vector.

i.e. creates new features from the old ones.

ex) $\vec{\phi} : \mathbb{R}^2 \rightarrow \mathbb{R}^5$:

$$\vec{\phi}(x_1, x_2) = (x_1, x_2, x_1^2, x_2^2, x_1 x_2)^T.$$

A basis function : $\phi_i : \mathbb{R}^d \rightarrow \mathbb{R}$.

↑
vector

↑

single new feature.

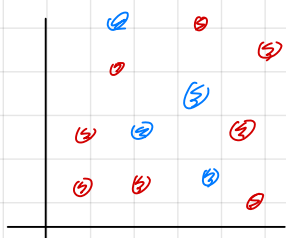
think of $\vec{\phi} : \mathbb{R}^d \rightarrow \mathbb{R}^k \Rightarrow \vec{\phi}(\vec{x}) = (\phi_1(\vec{x}), \dots, \phi_k(\vec{x}))$.

Feature map made of many basis function.

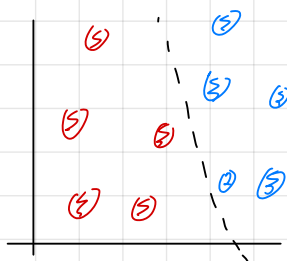
Why feature map?



turn non-linear patterns. $\Rightarrow$ linear patterns.



$\Rightarrow$



Input Space

Feature Space.

Procedure:

- Train:
1. Training set
 2. Feature map, new data set
 3. linear model

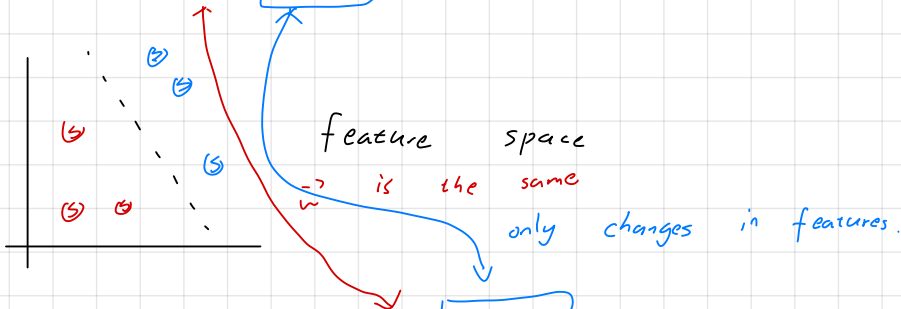
- Predict:
1. given new input $\vec{x}$.
 2. map to feature space $\vec{x} \rightarrow \vec{\phi}(\vec{x})$.
 3. Predict: $H(\vec{x}, \vec{w}) = \vec{w} * \text{Aug}(\vec{\phi}(\vec{x}))$.

Design matrix

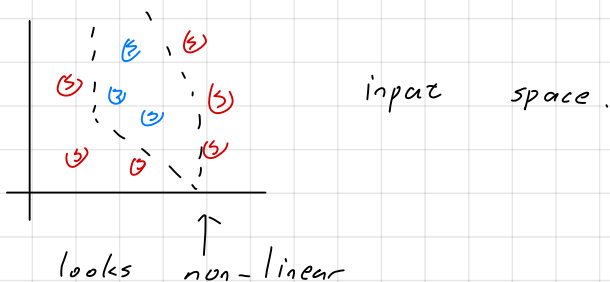
$$(n \times h) \Rightarrow (n \times k)$$

Prediction function:

$$H_{\phi}(\vec{z}) = \vec{w} \cdot \text{Aug}(\vec{z}) \quad \text{takes in vector } \vec{z} \text{ in feature space}$$



$$H(\vec{x}) = H_{\phi}(\vec{\phi}(\vec{x})) = \vec{w} \cdot \text{Aug}(\vec{\phi}(\vec{x})) \quad \text{take in vector } \vec{x} \text{ in input space}$$



ERM with Feature Maps

H_{ϕ} that minimizes risk in feature space

$\Leftrightarrow$

H that minimizes risk in input space.

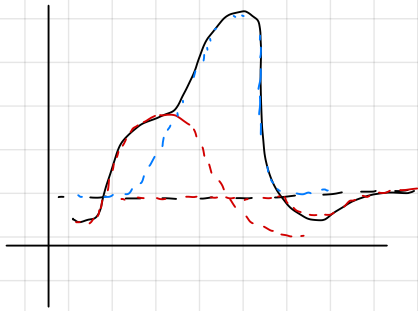
as long as H is a linear function of the parameter.

Thus, we can fit any function of

$$H(x) = w_1 \phi_1(x) + \dots + w_k \phi_k(x)$$

Gaussian Radial Basis Function (RBF)

idea: sum of bump



$$H(x) = w_1 \text{bump}_1(x) + w_2 \text{bump}_2(x) + w_3 \text{bump}_3(x).$$

Bump:

$$\phi_i(x) = \exp\left(-\frac{(x - \mu_i)^2}{\sigma_i^2}\right)$$

measure how close x_i to μ_i .

with center μ_i and width σ_i for each Gaussian basis function.

⇓

makes very general basis function.

by adjusting $\vec{w}$, we can fit complex pattern.

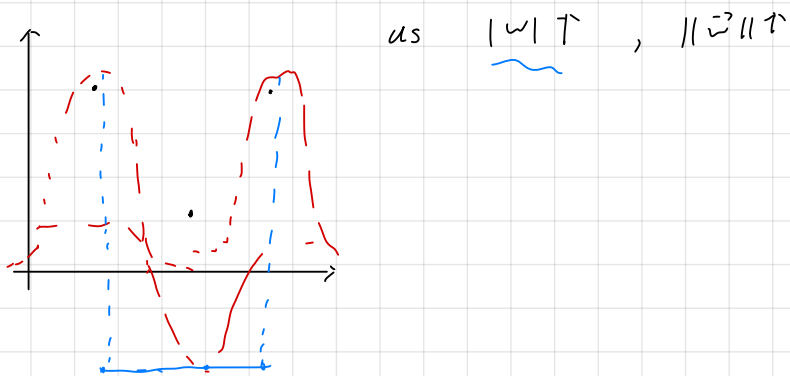
Height = w_i

Too many Gaussian function:

⇒ Overfitting.

⇓

Regularization:



$||\vec{\tilde{w}}||$ large as Gaussian function more complex / overfitting



measure the complexity of the model.



minimize $||\vec{\tilde{w}}||$ to prevent overfitting

Regularized least square:

idea: penalize large $||\vec{\tilde{w}}||$

$$\tilde{R}(\vec{\tilde{w}}) = \frac{1}{n} \sum_{i=1}^n (\vec{\tilde{w}} \cdot \phi(\vec{x}^{(i)}) - y_i)^2 + \underbrace{\lambda ||\vec{\tilde{w}}||^2}_{\text{regularizer.}}$$

hyperparameter that controls the trade-off

regularization term

Ridge Regression

$$\widehat{R}(\vec{\tilde{w}}) = \frac{1}{n} ||\Phi \vec{\tilde{w}} - \vec{y}||^2 + \lambda ||\vec{\tilde{w}}||^2.$$

Calculate $\frac{d\widehat{R}}{d\vec{\tilde{w}}}$ and set to 0:

solution:

$$\vec{\tilde{w}}^* = (\Phi^T \Phi + n\lambda \mathbf{I})^{-1} \Phi^T \vec{y}$$

↑
add small number λ to diagonal of $\Phi^T \Phi$.

⇓
improve condition number
helpful when multicollinearity exists.

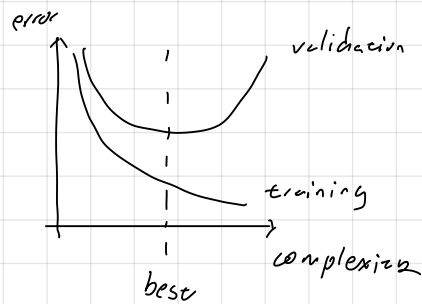
$\lambda \uparrow \Rightarrow$ more regularization.

$$\Downarrow \\ \|\vec{w}\|^2 \downarrow$$

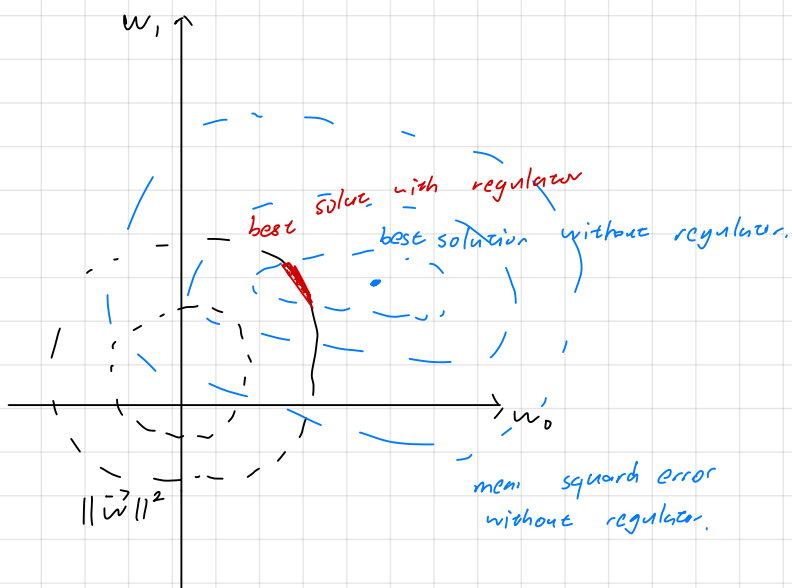
$\Downarrow$
simpler model / penalize complex model

How pick λ ?

Cross-validation.



Understanding: $\Rightarrow$ an increase in risk for a decrease in complexity.



Solution: least risk $\neq$

least $\|\vec{w}\|^2$.

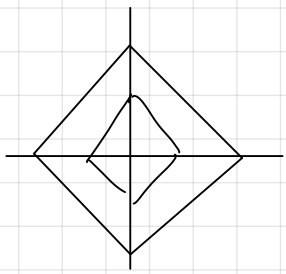
Lasso

p -norm: for $\forall p \in [0, \infty]$.

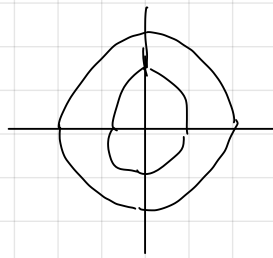
$$\|\vec{u}\|_p = \left(\sum_{i=1}^d |u_i|^p \right)^{1/p}.$$

$p=2 \Rightarrow$ euclidean distance.

$p=1 \Rightarrow$ manhattan distance



1-norm.



2-norm.

Use 1-norm $\Rightarrow$ Lasso.

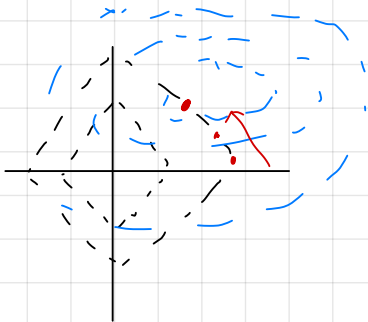
$$\tilde{R}(\vec{w}) = \frac{1}{n} \sum_{i=1}^n (\vec{w} \cdot \phi(\vec{x}^{(i)}) - y_i)^2 + \lambda \|\vec{w}\|_1$$

- No closed form solution.

- subgradient Descent?

- Sparse solution $\rightarrow$ feature importance.

• $\Downarrow$
why?



General regularization:

$$\tilde{R}(\vec{w}) = R(\vec{w}) + p(\vec{w}).$$

Regularized linear model:

Loss	Regularization	Name
square	2-norm	ridge regression.
square	1-norm	Lasso
square	1-norm + 2-norm	elastic net
hinge	2-norm	soft-SVM

Other regularization method: "Early Stopping". DNN.

High-Dimensional Feature Maps

linear prediction rule $\rightarrow$ straight linear line
(limitation!)



Basis functions $\vec{\phi}(\vec{x})$



maps data to a new space that is linearly separable.



$\Rightarrow$



$\Rightarrow$ train

$$H(\vec{x}) = w_0 + w_1 \phi_1(\vec{x}) + \dots$$

Data space

feature space

$\mathbb{R}^d$

$\vec{\phi}(\vec{x})$

$\mathbb{R}^k$

Decision boundary no longer a straight line.

Problem: How to choose $\vec{\phi}$?

Design a general feature map function that makes any dataset linearly-separable.



idea: very high-dimensional generic feature map.

(each additional feature makes the data easier to classify).

ex) Monomial: $\vec{\phi}(\vec{x}) = (1, x_1, x_2, x_1 x_2, x_1^2, x_2^2)^T$

$$(1, x_i, x_i x_j, x_i^2, \dots)^T$$

So, $\vec{x} \in \mathbb{R}^d$, then $\vec{\phi}(\vec{x}) \in \mathbb{R}^{1+2d+(\frac{d}{2})}$

$$(1, x_i, x_i x_j, x_i^2, x_i x_j x_k, x_i^3)^T$$

$$\text{or } \vec{\phi}(\vec{x}) \in \mathbb{R}^{1+3d+(\frac{d}{2})+(\frac{d}{3})}$$

problem: computationally costly for fitting a linear prediction rule.

Kernel Trick

idea: we can train many linear predictors as if we have mapped data to feature space, without actually doing so.

$$k(\vec{x}, \vec{x}') = \vec{\phi}(\vec{x}) \cdot \vec{\phi}(\vec{x}') \leftarrow \text{computationally costly}$$

↑
there is such function that does not require mapping to feature

↓
Kernel Function

$$\text{ex) } \vec{\phi}(\vec{x}) = (1, x_1^2, x_2^2, \sqrt{2}x_1, \sqrt{2}x_2, \sqrt{2}x_1x_2)^T$$

$$k(\vec{x}, \vec{x}') = (1 + \vec{x} \cdot \vec{x}')^2 \text{ is a kernel for this } \vec{\phi}.$$

polynomial kernel

$$k(\vec{x}, \vec{x}') = (1 + \vec{x} \cdot \vec{x}')^k \text{ is kernel for } k\text{-order monomial mappings.}$$

Idea: replacing all instances of $\vec{\phi}(\vec{x}) \cdot \vec{\phi}(\vec{x}')$ with $k(\vec{x}, \vec{x}')$

kernelize the algorithm

avoid explicitly mapping to feature space. (kernel trick)

- Only certain feature maps have efficiently-computed kernels.
- only certain learning algorithms can be kernelized.

Kernel Ridge Regression & Kernel SVM

ridge regression:

$$\arg \min_{\vec{w}} \frac{1}{n} \sum_{i=1}^n (\vec{\phi}(\vec{x}^{(i)}) \cdot \vec{w} - y_i)^2 + \lambda \|\vec{w}\|^2.$$

$\Uparrow$
 $\checkmark$

$$\arg \min_{\vec{w}} \frac{1}{n} \|\phi \vec{w} - \vec{y}\|^2 + \lambda \vec{w}^T \vec{w}$$

$\nearrow$

kernelize, replace all instances of $\vec{\phi}(\vec{x}) \cdot \vec{\phi}(\vec{x}')$ with $k(\vec{x}, \vec{x}')$

$\downarrow$

rewrite in dual form:

Fact: The solution $\vec{w}^*$ is a linear combination of $\vec{\phi}(\vec{x}^{(i)})$:

$$\vec{w}^* = \sum_{i=1}^n a_i \vec{\phi}(\vec{x}^{(i)})$$

why?

$$\frac{2}{n} \sum_{i=1}^n (\vec{\phi}(\vec{x}^{(i)}) \cdot \vec{w} - y_i) \vec{\phi}(\vec{x}^{(i)}) + 2\lambda \vec{w}, \text{ gradient of regularized risk}$$

$\downarrow$

Setting to zero, solving for $\vec{w}^*$:

$$\vec{w}^* = \sum_{i=1}^n \underbrace{\left(-\frac{1}{n\lambda} \vec{\phi}(\vec{x}^{(i)}) \cdot \vec{w}^* - y_i \right)}_{a_i} \vec{\phi}(\vec{x}^{(i)})$$

$\downarrow$

$$\vec{w}^* = \phi^T \vec{a}, \text{ where } \vec{a} = (a_1, \dots, a_n)^T$$

So, the problem:

$$\arg \min_{\vec{w}} \frac{1}{n} \|\phi \vec{w} - \vec{y}\|^2 + \lambda \vec{w}^T \vec{w}$$

$\Downarrow$

this can be kernelized.

$$\arg \min_{\vec{a}} \frac{1}{n} \|\phi \phi^T \vec{a} - \vec{y}\|^2 + \lambda \vec{a}^T \phi \phi^T \vec{a}$$

new optimization solution:

inside ϕ , we have

$$\phi = \begin{pmatrix} \vec{\phi}(\vec{x}^{(1)}) & \dots \\ \vdots & \ddots \\ \vec{\phi}(\vec{x}^{(n)}) & \dots \end{pmatrix}$$

$$\begin{aligned} (\phi \phi^T)_{ij} &= \vec{\phi}_i \cdot (\vec{\phi}^T)_j \\ &= \vec{\phi}_i \cdot \vec{\phi}_j \\ &= \vec{\phi}(\vec{x}^{(i)}) \cdot \vec{\phi}(\vec{x}^{(j)}) \\ &= k(\vec{x}^{(i)}, \vec{x}^{(j)}) \end{aligned}$$

And so, (i, j) entry of $\phi \phi^T$ is $\vec{\phi}(\vec{x}^{(i)}) \cdot \vec{\phi}(\vec{x}^{(j)}) = k(\vec{x}^{(i)}, \vec{x}^{(j)})$.

$$\phi \phi^T = \begin{pmatrix} k(\vec{x}^{(1)}, \vec{x}^{(1)}) & \dots & \dots \\ \vdots & \ddots & \vdots \\ \vdots & \vdots & k(\vec{x}^{(n)}, \vec{x}^{(n)}) \end{pmatrix} \quad \text{kernel matrix}$$

$$\text{So, we have } \arg \min_{\vec{a}} \frac{1}{n} \|K \vec{a} - \vec{y}\|^2 + \lambda \vec{a}^T K \vec{a}$$

$\Downarrow$

$$\vec{a}^* = (K + n \lambda I)^{-1} \vec{y}$$

The kernel ridge regression.

Making a prediction on new point:

$$H(\vec{x}) = \vec{w}^* \cdot \vec{\phi}(\vec{x}) = \sum_{i=1}^n a_i \vec{\phi}(\vec{x}^{(i)}) \cdot \vec{\phi}(\vec{x})$$

$$= \sum_{i=1}^n a_i k(\vec{x}^{(i)}, \vec{x})$$

A weighted sum of kernel evaluations.

Kernel Soft-SVM

$$\underset{\vec{a}}{\operatorname{argmin}} \left(\lambda \vec{a}^T K \vec{a} + \frac{1}{n} \sum_{i=1}^n \max \left\{ 0, 1 - y_i (K \vec{a})_i \right\} \right).$$

make new prediction:

$$H(\vec{x}) = \sum_{i \in S} q_i k(\vec{x}^{(i)}, \vec{x})$$

↑
the set of support vector.

Downside:

$n \times n$ kernel matrix can be very large.

Kernel Function

valid kernel: must complete the dot product w.r.t. some mapping, $\vec{\phi}(\vec{x})$.

⇓

$$k(\vec{x}, \vec{x}') = \vec{\phi}(\vec{x}) \cdot \vec{\phi}(\vec{x}')$$

New kernel constructed from other kernel functions.

- $k(\vec{x}, \vec{x}') = k_1 + k_2$
- $k(\vec{x}, \vec{x}') = k_1 \times k_2$
- $k(\vec{x}, \vec{x}') = k_3(\vec{\phi}(\vec{x}), \vec{\phi}(\vec{x}'))$
- $k(\vec{x}, \vec{x}') = f(\vec{x}) k_1 f(\vec{x}')$

Theorem:

A symmetric function k is a valid kernel

⇓

the kernel matrix, K , is positive semi-definite for any choice of data $\vec{x}^{(1)}, \dots, \vec{x}^{(n)}$.

RBF kernel Gaussian kernel

$$k(\vec{x}, \vec{x}') = e^{\frac{-\|\vec{x} - \vec{x}'\|^2}{2\sigma^2}} \leftarrow \text{distance between } \vec{x} \text{ and } \vec{x}'$$

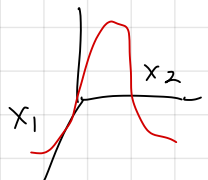
$$= e^{-\gamma \|\vec{x} - \vec{x}'\|^2}, \text{ where } \gamma = 1/(2\sigma^2)$$

Interpretation: RBF kernel measures similarity of $\vec{x}$ and $\vec{x}'$.

very similar: $k(\vec{x}, \vec{x}) \approx 1$

very different: $k(\vec{x}, \vec{x}) \approx 0$.

Parameter σ (or γ) controls the scale

ex) 

$$H(\vec{x}) = \sum_{i=1}^n a_i k(\vec{x}^{(i)}, \vec{x})$$

larger σ (smaller γ),
wider the Gaussian.

sum of bumps

$$\vec{x}, \vec{x}' \in \mathbb{R}^2$$

$$\vec{x} = (x_1, x_2)$$

KBF kernel

Prediction: $H(\vec{x}) = \sum_{i=1}^n a_i k(\vec{x}^{(i)}, \vec{x})$

• One parameter a_i learned for each training point $\vec{x}^{(i)}$.

• $k(\vec{x}^{(i)}, \vec{x}) \approx 0$ for any $\vec{x}^{(i)}$ far from $\vec{x}$.

• $H(\vec{x})$ is largely determined by the training points closest to $\vec{x}$.

↓
associated weights

Since $k(\vec{x}^{(i)}, \vec{x})$
determine the distance

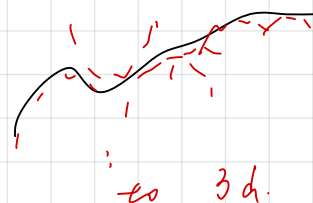
Understanding:

RBF function placed at each training point $\Rightarrow$ sum of bumps

$H(\vec{x})$ largely determined by training point closest to $\vec{x}$.

as a surface

↓
a generalization of KNN rule



RBF Kernel Map

The mapping $\vec{\phi}(\vec{x})$ with entries of the form: $\mathbb{R}^2 \rightarrow \mathbb{R}^\infty$

$$e^{-\|\vec{x}\|^2/2} x_i, \quad \frac{1}{\sqrt{2!}} e^{-\|\vec{x}\|^2/2} x_i x_j, \quad \frac{1}{\sqrt{3!}} e^{-\|\vec{x}\|^2/2} x_i x_j x_k, \dots$$

mapping to an infinite dimensional Hilbert space

Hyperparameter: $\gamma \rightarrow$ kernel width.

underfitting

$\gamma \downarrow \rightarrow$ STDT $\downarrow$

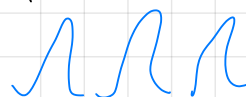
$\uparrow \uparrow$

$\gamma \downarrow \rightarrow$ more wider $\Rightarrow$ gaussian overlapping

$\gamma \uparrow \rightarrow$ more narrower $\Rightarrow$ gaussian separating

$\Rightarrow$ STD $\downarrow \rightarrow$

overfitting



Probabilistic Modeling

ex) Given a new flower with $X = x$ petals, predict its species, Y .
 $\uparrow$
 random variable

Assumption: When nature generates a new flower, it picks X & Y from some probability distribution.

• Joint distribution $\mathbb{P}(X=x, Y=y)$ gives us full info.

	$Y=0$	$Y=1$
$X=0$	0%	.
.	5%	.
.	10%	.
.	.	.
$X=6$	0%	.

$\Rightarrow$ Entries of the joint distribution table sum to 100%

$$\sum_{x \in \{0, \dots, 6\}} \sum_{y \in \{0, 1\}} \mathbb{P}(X=x, Y=y) = 1.$$

• Marginal distribution for X is found by summing over values of Y

$\Downarrow$

$$\mathbb{P}(X=x) = \sum_{y \in \{0, 1\}} \mathbb{P}(X=x, Y=y)$$

or

$$\mathbb{P}(Y=y) = \sum_{x \in \{0, \dots, 6\}} \mathbb{P}(X=x, Y=y)$$

Conditional Probability:

$$\mathbb{P}(Y=y | X=x) = \frac{\mathbb{P}(X=x, Y=y)}{\mathbb{P}(X=x)} \quad \text{or} \quad \mathbb{P}(X=x | Y=y) = \frac{\mathbb{P}(X=x, Y=y)}{\mathbb{P}(Y=y)}$$

	Y=0	Y=1		$\mathbb{P}(Y=y X=1)$
X=0	0%	0%	→	Y=0 100%
X=1	5%	0%		Y=1 0%
⋮				
X=4	5%	20%	→	$\mathbb{P}(Y=y X=4)$
⋮				

Y=0	20%	← $\frac{5\%}{25\%}$
Y=1	80%	← $\frac{20\%}{25\%}$

Bayes Theorem

$$\mathbb{P}(Y=y | X=x) = \frac{\mathbb{P}(X=x | Y=y) \mathbb{P}(Y=y)}{\mathbb{P}(X=x)}$$

⇓

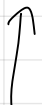
Bayes Decision Theorem

Predict the most likely label y given $X=x$

i.e. predict the y that maximizes $\mathbb{P}(Y=y | X=x)$

In Binary Probabilistic Classification:

- Predict 1 if $\mathbb{P}(Y=1 | X=x) > \mathbb{P}(Y=0 | X=x)$; otherwise predict 0.



Bayes Classification rule

• alternative form:

Predict 1 if $\mathbb{P}(X=x | Y=1) \mathbb{P}(Y=1) > \mathbb{P}(X=x | Y=0) \mathbb{P}(Y=0)$

otherwise predict 0.

Continuous Distribution

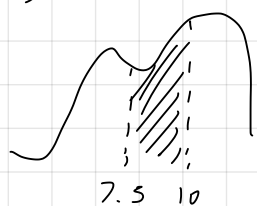
pdf: $p: \mathbb{R} \rightarrow \mathbb{R}^+$

$$\mathbb{P}(a < X < b) = \int_a^b p_X(x) dx$$

How likely it is to get a value of X in any interval $[a, b]$.

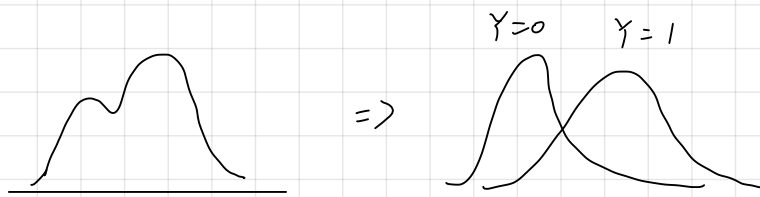
$\mathbb{P}(X = x) = 0$. $\Rightarrow$ prob. of continuous r.v. being exactly a certain value is zero.

ex)



$$\mathbb{P}(7.5 < X < 10) = \int_{7.5}^{10} p_X(x) dx$$

• Conditional density $p(x|Y=y)$, conditional on Y .

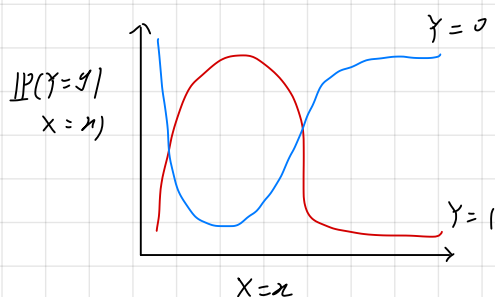


• Conditional on X:

$\Downarrow$

conditional distribution of Y given X :

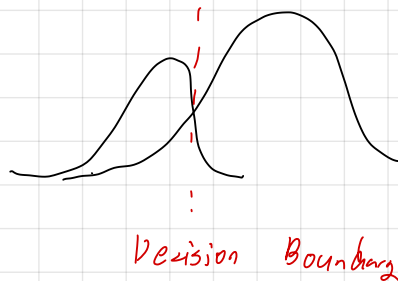
Discrete, b/c Y is discrete



• Joint Distribution

$$p(x, 0) = p(x | Y=0) \mathbb{P}(Y=0)$$

$$p(x, 1) = p(x | Y=1) \mathbb{P}(Y=1)$$



predict class 1 if $p(x | Y=1) \mathbb{P}(Y=1) > p(x | Y=0) \mathbb{P}(Y=0)$.

• Multivariate Distribution

random vector $\vec{X}$

$$\text{pdf: } p_X(\vec{X}) : \mathbb{R}^d \rightarrow \mathbb{R}^+$$

• Bayes Error

• Predict 0, true 1 ①

• Predict 1, true 0 ②

$$\mathbb{P}(\text{error}) = \mathbb{P}(\text{Case ①}) + \mathbb{P}(\text{Case ②}).$$

$$\begin{aligned} \text{Case ①: } \mathbb{P}(\text{case ①}) &= \mathbb{P}(Y \text{ is actually } 1, \text{ predict } 0) \\ &= \mathbb{P}(Y \text{ is actually } 1) \cdot \mathbb{P}(\text{predict } 0 | Y \text{ is actually } 1) \end{aligned}$$

$$\begin{aligned} \text{Case ② } \mathbb{P}(\text{case ②}) &= \mathbb{P}(Y \text{ is actually } 0, \text{ predict } 1) \\ &= \mathbb{P}(Y \text{ is actually } 0) \cdot \mathbb{P}(\text{predict } 1 | Y \text{ is actually } 0) \end{aligned}$$

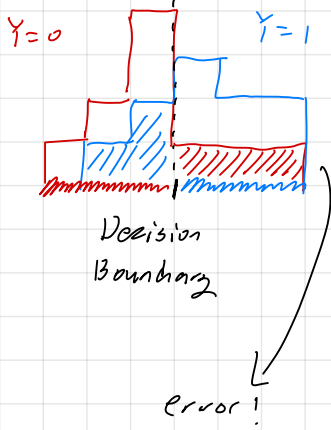
• Bayes classifier has the smallest possible error probability of any classifier.

⇓

the error is called Bayes error

• In most cases, the minimum possible error probability is > 0 .

Bayes Error if $\mathbb{P}(Y=0) = \mathbb{P}(Y=1) = 0.5$.



$$\sim P_0 = (x | Y=0)$$

$$\sim P_1 = (x | Y=1)$$

$$\underbrace{\mathbb{P}(Y \text{ actually } 1)}_{0.5} \cdot \underbrace{\mathbb{P}(\text{predict } 0 | Y=1)}_{0.3}$$

↓
area of P_1 under $Y=0$

$$\underbrace{\mathbb{P}(Y \text{ actually } 0)}_{0.3} \cdot \underbrace{\mathbb{P}(\text{predict } 1 | Y=0)}_{0.3}$$

↓
area of P_0 under $Y=1$

Bayes classifier !
is optimal

⇓

problem it requires knowing the joint distribution.

⇓

Estimate : do not know joint distribution.

we know the data.

estimate : $\mathbb{P}(X=x, Y=y) \approx \frac{\#(X=x \ \& \ Y=y)}{n}$
joint Prob.

X	Y
5	0
3	0
4	1
4	1

estimate marginal : $\mathbb{P}(Y=y) \approx \frac{\#(Y=y)}{n}$

estimate conditional : $\mathbb{P}(X=x | Y=y) \approx \frac{\#(X=x \ \& \ Y=y)}{\#(Y=y)}$

$$\mathbb{P}(Y=y | X=x) \approx \frac{\#(X=x \ \& \ Y=y)}{\#(X=x)}$$

Law of large number : as data size $n \rightarrow \infty$, these estimated prob.

converges to their true values

Using estimated prob. for Bayes classifier: (no longer optimal)

estimate $\mathbb{P}(Y=0 | X=x)$
 $\mathbb{P}(Y=1 | X=x)$ } $\rightarrow$ compare and predict based on Bayes decision rule.

Multivariate estimate:

$$\mathbb{P}(Y=y | X=x_1, X=x_2, \dots, X=x_n)$$

$$\mathbb{P}(X=x_1, \dots | Y=y)$$

Histogram Density Estimators

estimate density for continuous distribution

problem: most values never seen in data
 $\downarrow$

Histogram estimator

• Suppose $X_1, \dots, X_n$ came from density f .

• Divide domain into k bins: $[a_i, b_i)$.

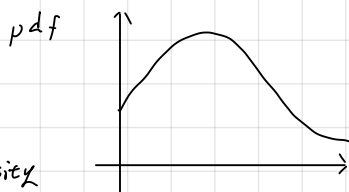
• within each bin, $\downarrow$ often equal-sized.

$$f(x) \text{ within bin } i \approx \frac{\# \text{ data points } \in [a_i, b_i)}{n \times \underbrace{(b_i - a_i)}_{\text{bin width}}}$$

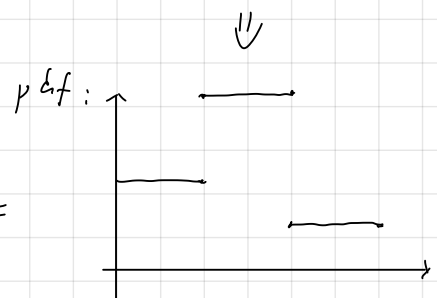
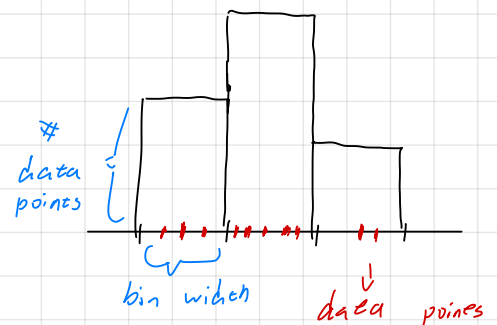
As we get more data

- decrease bin width
- increase number of bins

As n and $\# \text{ bins} \rightarrow \infty$,
estimator approaches true density
by law of large number.



histogram
integrals to 1



get smooth

estimate $p(X | Y=y)$

Conditional probability: restrict data where $Y=1/0$ and use histogram estimator

↓
by law of large number, approaches true density



estimate $\mathbb{P}(Y=y|X=x)$, discrete prob, conditioned on continuous variable

↓
Two Approaches:

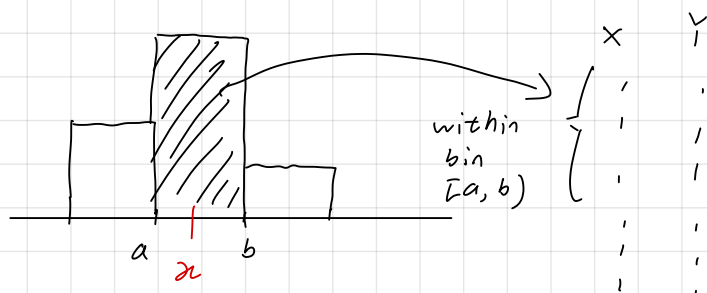
↓
problem: x may not be seen in data

① Count $\#(Y=1)$ and $\#(Y=0)$ within bin containing x

• find bin containing x

• estimate

$$\mathbb{P}(Y=y|X=x) \approx \frac{\#(Y=y \text{ within bin})}{\#(\text{points within bin})}$$



$\Rightarrow$ count $Y=y$, divide by $\#$ points

Bayes' rule:

predict $Y=y$ that occurs most within the bin.

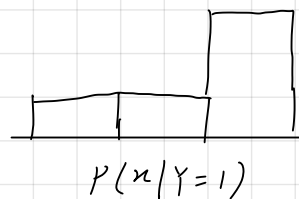
② Compute from Bayes' rule

• estimate $p(x|Y=y)$ $\mathbb{P}(Y=y)$ $p_X(x)$

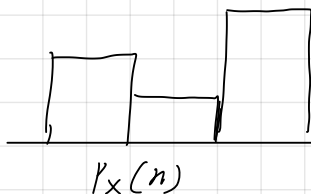
• Use Bayes' rule $\mathbb{P}(Y=y|X=x) = \frac{p(x|Y=y)\mathbb{P}(Y=y)}{p_X(x)}$

$\mathbb{P}(Y=y) \rightarrow$ count

$p(x|Y=y) \rightarrow$ restrict data, histogram estimator

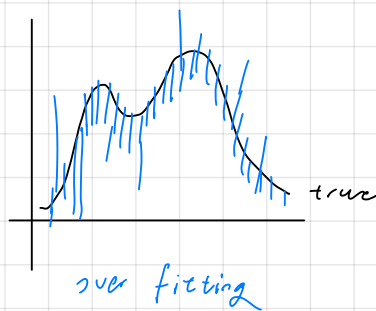
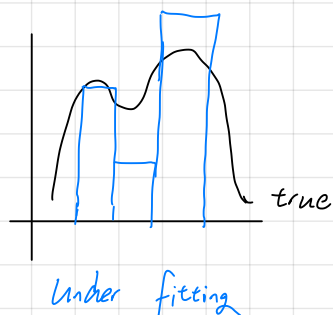


$p_X(x) \rightarrow$ histogram estimator



Both approaches the same if we use the same bins

Like KNN, # of bins might result over- or under-fitting.

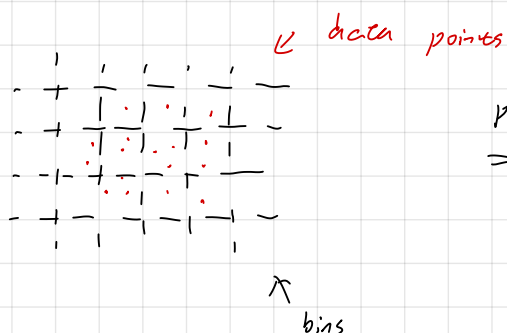


Multivariate Histogram Density Estimators

Estimate $p(\vec{x})$

- divide $\mathbb{R}^d$ into rectangular bins with regular side-length $l_1, \dots, l_d$
- estimate

$$f(\vec{x}) \text{ within bin} \approx \frac{\# \text{ data points} \in [a_i, b_i)}{n \times (l_1 \times \dots \times l_d)}$$



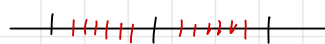
pdf
 $\Rightarrow$



as $n \rightarrow \infty$
bin

Curse of dimensionality

Not accurate b/c fewer points in each bins as we divide each dimensionality while estimating



6 points in each bin

$\Rightarrow$



one data point in each bin.

Parametric Density Estimation

Histogram needs massive data in high dimension.

↓

has no shape of true density / very flexible but yet "data hungry"

↓

idea: assume true density follows some distribution.

ex) Gaussians

$$p(x; \mu, \sigma) = \frac{1}{\sigma \sqrt{2\pi}} e^{\frac{-(x-\mu)^2}{2\sigma^2}}$$

CLT: sums of indep. random variable are Gaussian.

↓
parameters

↓ ↓
 μ : center σ : width

Parametric Distributions determined by finite number of parameters.

ex) $p(x; \mu, \sigma) = \frac{1}{x \sigma \sqrt{2\pi}} \exp\left(-\frac{(\ln x - \mu)^2}{2\sigma^2}\right)$ - log normal.

• product of indep. positive random number

• ex) length of comments

$$p(x; k, \theta) = \frac{1}{\Gamma(k) \theta^k} x^{k-1} e^{-x/\theta} \quad - \text{Gamma}$$

• ex) wait time, ...

$$p(x; \alpha, a) = \frac{\alpha x^\alpha}{x^{\alpha+1}} \quad - \text{Pareto}$$

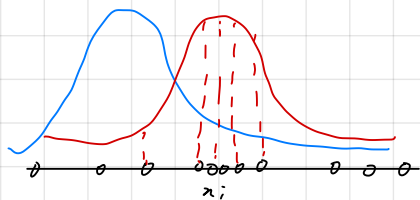
• ex) distribution of wealth.

Idea: Assume data comes from some parametric density

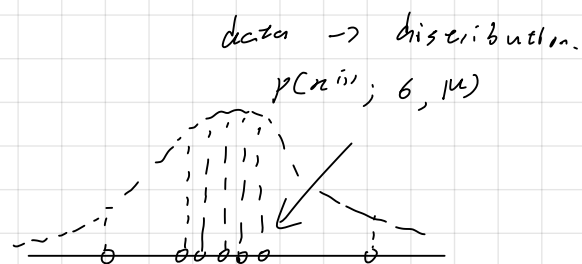
↑
Use data to estimate parameter

Maximum Likelihood Estimation

Use data to estimate parameter → estimating the parameter



- not likely
- more likely



Let p be Gaussian prob. density function.

$p(x^{(i)}; \mu, \sigma)$ quantifies how likely it is to see $x^{(i)}$ if parameter μ and σ used.

$\downarrow$
 $\prod_{i=1}^n p(x^{(i)}; \mu, \sigma)$ quantifies the likelihood of seeing $x^{(1)} \dots x^{(n)}$

• joint density of data

Likelihood of μ and σ , with respect to data $x^{(1)}, \dots, x^{(n)}$.

$\downarrow$

$$\ell(\mu, \sigma; x^{(1)} \dots x^{(n)}) = \prod_{i=1}^n p(x^{(i)}; \mu, \sigma).$$

Find $\arg \max_{\mu, \sigma} \ell(\mu, \sigma; x^{(1)} \dots x^{(n)})$. maximize likelihood.

$\downarrow$
more data points has high values

$\downarrow$

• Log-likelihood $\rightarrow$ easier to work with.

$$\tilde{\ell}(\mu, \sigma) = \ln \ell(\mu, \sigma).$$

$$= \ln \prod_{i=1}^n p(x^{(i)}) \quad \leftarrow \ln(a \cdot b) = \ln(a) + \ln(b)$$

$$= \sum_{i=1}^n \ln(p(x^{(i)}))$$

• Derivate and set to 0.

Solution:

MLE for Gaussian Distribution:

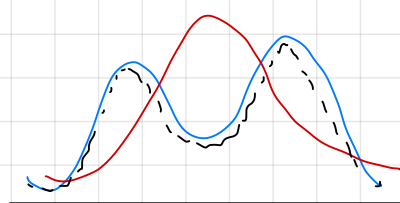
$$\mu_{MLE} = \frac{1}{n} \sum_{i=1}^n x^{(i)} \quad \sigma_{MLE} = \sqrt{\frac{1}{n} \sum_{i=1}^n (x_i - \mu_{MLE})^2}$$

fitting data:

compute μ_{MLE} , σ_{MLE} .

Parametric vs. non-parametric estimation

- Non-parametric approaches can fit arbitrary density, but require lots of data - especially in high dimensions.
- Parametric approaches require less data, provided that they are correctly specified.



- true
- parametric estimation.
- non-parametric estimation

Bayes with multiple features

with k feature, pick y maximizing

$$P_{\vec{x}}(\vec{x} | Y=y) \prod (Y=y)$$

Multivariate Gaussian density estimation

$$X \sim N(\mu, \sigma^2)$$

$$P(x) = \frac{1}{\sqrt{2\pi}\sigma^2} e^{-\frac{1}{2}(x-\mu)^2/\sigma^2}$$

$\Downarrow$

① d indep. r.v. $x_1, \dots, x_d$ from Gaussian, different mean but same variance.

joint density

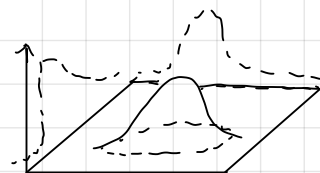
$$p(x_1, \dots, x_d) = p(x_1) p(x_2) \dots p(x_d)$$

$$= \left(\frac{1}{\sqrt{2\pi}\sigma^2} e^{-\frac{1}{2}(x_1-\mu_1)^2/\sigma^2} \right) \dots \dots (\dots)$$

$$= \frac{1}{(2\pi\sigma^2)^{d/2}} \exp \left(\frac{(x_1-\mu_1)^2 + (x_2-\mu_2)^2 + \dots + (x_d-\mu_d)^2}{2\sigma^2} \right)$$

$$= \frac{1}{(2\pi\sigma^2)^{d/2}} \exp \left(\frac{-\|\vec{x} - \vec{\mu}\|^2}{2\sigma^2} \right)$$

$\uparrow$
controls the width



spherical Gaussians

2-dimensional example

② $x_1, \dots, x_d$ indep Gaussian
different mean and variances

$$p(x_1, \dots, x_d) = \dots$$

$$= \frac{1}{(2\pi)^{d/2} \sigma_1 \dots \sigma_d} \exp\left(-\frac{1}{2} \left[\frac{(x_1 - \mu_1)^2}{\sigma_1^2} + \dots + \frac{(x_d - \mu_d)^2}{\sigma_d^2} \right]\right)$$

Define:

$$C = \begin{pmatrix} \sigma_1^2 & 0 & \dots & 0 \\ 0 & \sigma_2^2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \sigma_d^2 \end{pmatrix}$$

$$p(\vec{x}) = \frac{1}{(2\pi)^{d/2} |C|^{\frac{1}{2}}} \exp\left(-\frac{1}{2} (\vec{x} - \vec{\mu})^T C^{-1} (\vec{x} - \vec{\mu})\right)$$

where $|C|$ is the determinant of C .

↓
product of diagonal of C .

③ $x_1, \dots, x_d$ not indep.

Covariance:

$$\text{Cov}(x_i, x_j) = E[(x_i - \mu_i)(x_j - \mu_j)]$$

Note: $\text{Var}(x_i) = \text{Cov}(x_i, x_i)$.

Covariance measure how much two quantities vary together.

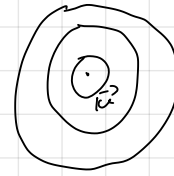


$$\text{Cov}(x_i, x_j) = E[(x_i - \mu_i)(x_j - \mu_j)]$$

Covariance Matrix:

$$C = \begin{pmatrix} \text{Var}(x_1) & \text{Cov}(x_1, x_2) & \dots & \dots \\ \text{Cov}(x_1, x_2) & \text{Var}(x_2) & & \\ \vdots & & \ddots & \\ \vdots & & & \text{Var}(x_d) \end{pmatrix}$$

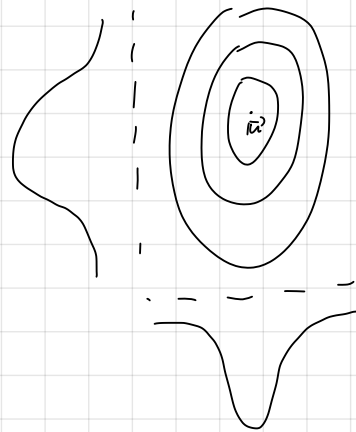
↓



- Contours are (hyper)spheres
- every slice through middle gives same gaussian

Axis-Aligned Gaussians

- Contours are axis aligned (hyper)ellipses
- C is the covariance matrix
- Diagonal
- Entries are variances.



C symmetry, $\text{Cov}(x_i, x_j) = \text{Cov}(x_j, x_i)$

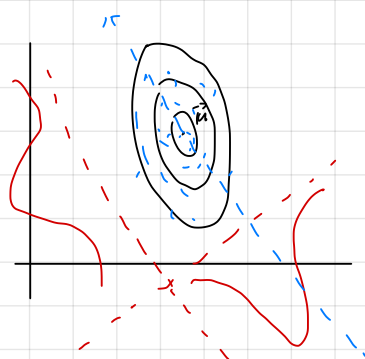
$$p(\vec{x}) = \frac{1}{(2\pi)^{d/2} |C|^{1/2}} \exp\left(-\frac{1}{2}(\vec{x} - \vec{\mu})^T C^{-1}(\vec{x} - \vec{\mu})\right)$$

rotated axis-aligned Gaussian

$C \rightarrow$ covariance matrix

$$\begin{pmatrix} \cdot & \text{Cov} \\ \text{Cov} & \cdot \end{pmatrix}$$

negative



slope negative

negative covariance

which matrix are valid covariance matrices:

• if C is the rotation of some diagonal covariance matrix C_0 . That is $C = R C_0$.

• C is symmetric, positive semi-definite, $x C x^T \geq 0$.

↓
eigenvalues ≥ 0 .

3 Cases

① C diagonal, with all same entries

② C diagonal, different entries

③ C not diagonal

MLE for multivariate Gaussian

$\vec{\mu}$: controls Gaussian's location

C : controls Gaussian's shape

$$\vec{\mu}_{MLE} = \frac{1}{n} \sum_{i=1}^n \vec{x}^{(i)}$$

For C , make assumption.

more parameter estimate

- Spherical
- Axis-aligned
- General

more-likely
to overfitting

① Spherical :

$$\sigma_{MLE}^2 = \frac{1}{n} \sum_{i=1}^n \| \vec{x}^{(i)} - \vec{\mu}_{MLE} \|^2$$

variance of the data



② Axis-aligned

σ_1^2 = sample variance of heights

σ_2^2 = sample variance of weights

③ General Gaussians.

MLE for covariance i and j :

$$C_{ij} = \left(\frac{1}{n} \sum_{k=1}^n \vec{x}_i^{(k)} \vec{x}_j^{(k)} \right) - \mu_i \mu_j$$

How: make matrix

$$\begin{pmatrix} x_1^{(1)} & x_1^{(2)} \\ \vdots & \vdots \\ x_n^{(1)} & x_n^{(2)} \end{pmatrix} \Rightarrow X = \begin{pmatrix} x_1^{(1)} - \bar{x}^{(1)} & x_1^{(2)} - \bar{x}^{(2)} \\ \vdots & \vdots \\ x_n^{(1)} - \bar{x}^{(1)} & x_n^{(2)} - \bar{x}^{(2)} \end{pmatrix} \Rightarrow C = \frac{1}{n} X^T X$$

Discriminant Analysis

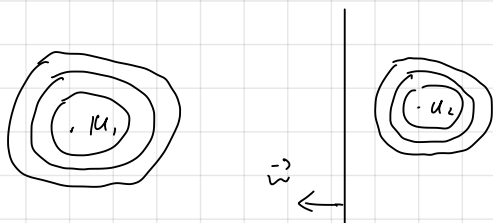
Bayes classifier

1. fit Gaussian for $p(\vec{x} | Y=y)$

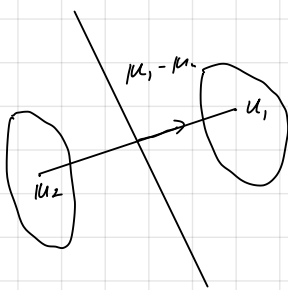
2. for new point, predict y maximizing $p(\vec{x} = \vec{x} | Y=y) \mathbb{1}_{(Y=y)}$

Decision Boundary:

surface between different classification



$$\text{choose class 1 if } \vec{w} \cdot \frac{(\vec{\mu}_1 - \vec{\mu}_2)}{\sigma^2} \geq \theta$$



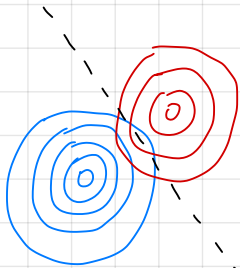
covariance matrix identical, diagonal

↓
axis-aligned Gaussians

predict class 1 if $\vec{x} \cdot \vec{w} \geq 0$.

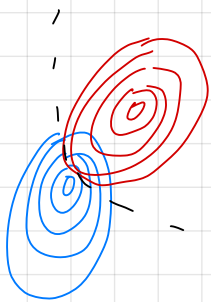
$$C = \frac{n_1 C_1 + n_2 C_2}{n_1 + n_2}$$

← compute covariance for each class,
perform weighted average



Decision boundary

- Covariance matrix equal, decision boundary linear
- Linear Discriminant Analysis (LDA).



Decision boundary

- Covariance matrix C_1, C_2 different, non-diagonal.
- Decision boundary is quadratic

↓
Quadratic Discriminant Analysis.
(QDA)

Conditional Independence:

A and B independent if

$$\mathbb{P}(A \cdot B) = \mathbb{P}(A) \mathbb{P}(B)$$

⇔

$$\mathbb{P}(A|B) = \mathbb{P}(A).$$

check if $\mathbb{P}(A|B) = \mathbb{P}(A)$

$$\mathbb{P}(B|A) = \mathbb{P}(B)$$

↙

A and B are independent if learning B does not influence your belief that A happens.

Real world data are always not independent =

↓

but some are very close!

A, B conditionally independent given C if

$$\mathbb{P}(A, B | C) = \mathbb{P}(A | C) \cdot \mathbb{P}(B | C)$$

$\Uparrow$

$$\mathbb{P}(A | B, C) = \mathbb{P}(A | C)$$

$\uparrow$

learning B will not influence A given C.

learning B happens in addition to C does not influence your belief that A happens given C.

ex) Survival and class not indep.

$$\mathbb{P}(\text{Survived} = 1) = 0.408$$

$$\mathbb{P}(\text{Survived} | \text{Pclass} = 1) = 0.657$$

they're close to conditionally independent given ticket price:

$$\mathbb{P}(\text{Survived} = 1 | \text{Pclass} = 1, \text{Fare} > 50) = 0.708$$

$$\mathbb{P}(\text{Survived} = 1 | \text{Fare} > 50) = 0.696$$

learning $\text{Pclass} = 1$ will not influence since fare already telling Pclass

$X_1, \dots, X_d$ mutually conditionally independent given Y if

$$\mathbb{P}(X_1, \dots, X_d | Y) = \mathbb{P}(X_1 | Y) \cdot \mathbb{P}(X_2 | Y) \cdots \mathbb{P}(X_d | Y)$$

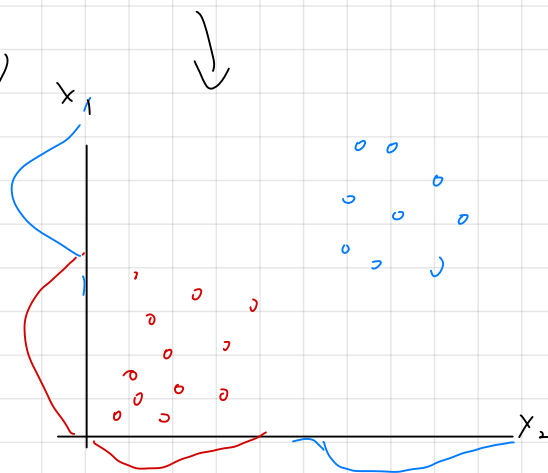
Curse of dimensionality

$\downarrow$

$$p(\vec{x} | Y=y) = p(x_1, \dots, x_d | Y=y)$$

However, if features are mutually conditionally independent given Y

$$\downarrow$$
$$= p_1(x_1 | Y=y) p_2(x_2 | Y=y) \cdots p_d(x_d | Y=y)$$



For Bayes classifier,

estimate

$$p(x_1, \dots, x_d | Y = y_i)$$

$$= p_1(x_1 | Y = y_i) \cdots p_d(x_d | Y = y_i) \quad \text{if } x_1, \dots, x_d \text{ mutually conditionally independent given } Y$$

↓

reduced the number of bins needed to cover the input space for histogram estimate

Naive Bayes Algorithm.

① Assume $x_1, \dots, x_d$ mutually independent given the class label.

② Estimate one-dimensional densities $p_1(x_1 | Y = y_i) \cdots p_d(x_d | Y = y_i)$

③ pick y_i maximizes $p_1(x_1 | Y = y_i) \cdots p_d(x_d | Y = y_i) \mathbb{P}(Y = y_i)$.

Gaussian Naive Bayes

$p(x_1 | Y) \cdots p(x_d | Y)$ a product of 1-d Gaussian with different mean and variance

$$C = \begin{pmatrix} \sigma_1^2 & & 0 \\ & \ddots & \\ 0 & & \sigma_d^2 \end{pmatrix} \quad \begin{array}{l} d\text{-dimensional gaussian} \\ \text{with diagonal covariance matrix} \end{array}$$

each class has own diagonal covariance matrix

↓

⇒ QDA with diagonal covariance.

Naive Bayes work well in high-dimensions

practical issues

$\mathbb{P}(x_1 | Y) \cdots \mathbb{P}(x_d | Y)$ multiplying lots of small probability

potential for underflow

Trick: $\log [\mathbb{P}(X_1 = x_1 | Y = y_i) \cdots \mathbb{P}(X_d = x_d | Y = y_i) \mathbb{P}(Y = y_i)]$

$$= \left(\sum_{j=1}^d \log \mathbb{P}(X_j = x_j | Y = y_i) \right) + \log \mathbb{P}(Y = y_i)$$

↑
summation
not multiplication.

Regression in the sense of Probability

ERM $\rightarrow$ least square
(minimize error)

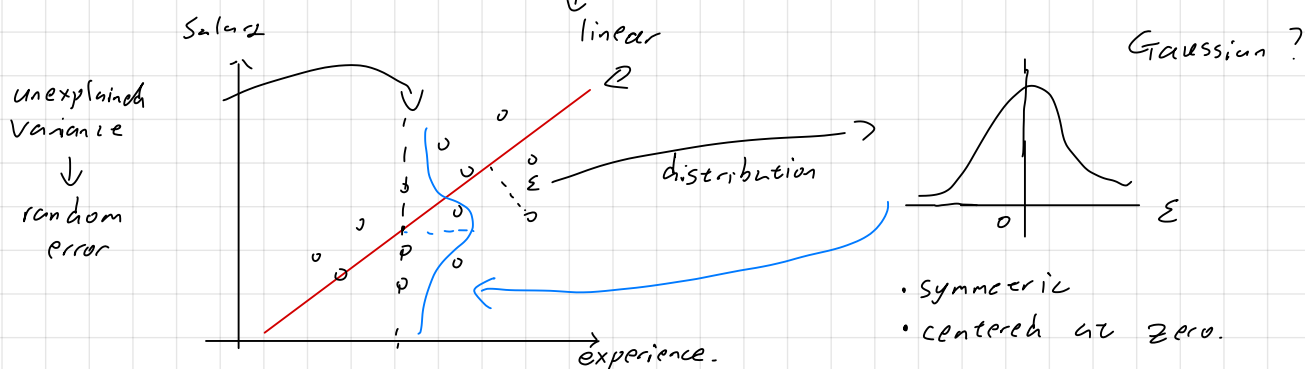
↓
"curve fitting" approach.

From probabilistic view:

There is uncertainty in the data generation process.

Modeling uncertainty

$$\text{Salary} = \underbrace{w_0 + w_1 \times (\text{Experience})}_{\text{linear}} + \underbrace{\epsilon}_{\text{random error}}$$



$$\text{Salary} = w_0 + w_1 \times (\text{Experience}) + \underbrace{N(0, \sigma^2)}_{\epsilon}$$

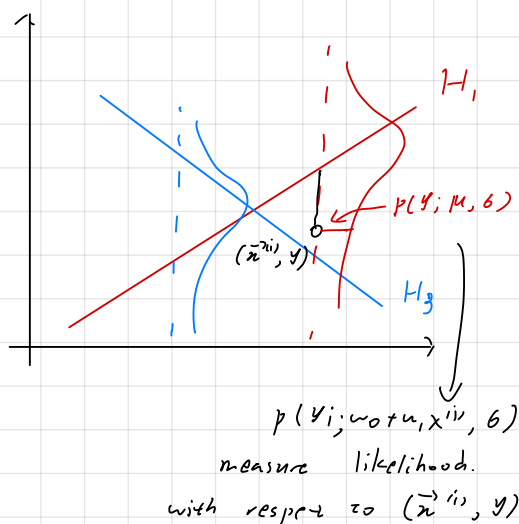
$$\Leftrightarrow \text{Salary} \sim N(w_0 + w_1 \times (\text{Experience}), \sigma^2)$$

In general, $Y \sim N(\text{Aug}(\vec{x}) \cdot \vec{w}, \sigma^2)$.

for any feature vector $\vec{x}$, target Y is drawn from Gaussian centered at $\text{Aug}(\vec{x}) \cdot \vec{w}$.

Given some data, $\Rightarrow$ parameters

estimate $\xrightarrow{\quad}$ by maximizing likelihood.



likelihood:

$p(y; \mu, \sigma)$ gaussian pdf:

$$p(y; \mu, \sigma) = \frac{1}{\sqrt{2\pi} \sigma} e^{-\frac{(y-\mu)^2}{2\sigma^2}}$$

observe data set $\{(\vec{x}^{(ii)}, y_i)\}$.

likelihood

So

$$L(\vec{w}, \sigma) = \prod_{i=1}^n p(y_i; \text{Aug}(\vec{x}^{(ii)}) \cdot \vec{w}, \sigma)$$

$$L(\vec{w}, \sigma) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi} \sigma^2} e^{-\frac{1}{2} \frac{(y_i - \vec{w} \cdot \vec{x}^{(ii)})^2}{\sigma^2}}$$

$$\begin{aligned} \tilde{L}(\vec{w}, \sigma) &= \ln \prod_{i=1}^n \frac{1}{\sqrt{2\pi} \sigma^2} e^{-\frac{1}{2} \frac{(y_i - \vec{w} \cdot \vec{x}^{(ii)})^2}{\sigma^2}} \\ &= \sum_{i=1}^n \ln \left(\frac{1}{\sqrt{2\pi} \sigma^2} e^{-\frac{1}{2} \frac{(y_i - \vec{w} \cdot \vec{x}^{(ii)})^2}{\sigma^2}} \right) \\ &= \sum_{i=1}^n \ln \left(\frac{1}{\sqrt{2\pi} \sigma^2} \right) + \ln \left(e^{-\frac{1}{2} \frac{(y_i - \vec{w} \cdot \vec{x}^{(ii)})^2}{\sigma^2}} \right) \end{aligned}$$

$$= \sum_{i=1}^n \ln \left(\frac{1}{\sqrt{2\pi} \sigma^2} \right) - \frac{1}{2} \frac{(y_i - \vec{w} \cdot \vec{x}^{(ii)})^2}{\sigma^2}$$

= ...

$$= -\frac{1}{2\sigma^2} \sum_{i=1}^n (\text{Aug}(\vec{x}^{(ii)}) \cdot \vec{w} - y_i)^2 + \frac{n}{2} \ln \frac{1}{\sigma^2} - \frac{n}{2} \ln(2\pi)$$

$$\arg \max_{\vec{w}} \left[-\frac{1}{2\sigma^2} \sum_{i=1}^n (\text{Aug}(\vec{x}^{(ii)}) \cdot \vec{w} - y_i)^2 + \frac{n}{2} \ln \frac{1}{\sigma^2} - \frac{n}{2} \ln(2\pi) \right]$$

constant

$$= \arg \max_{\vec{w}} \left[-\frac{1}{2\sigma^2} \sum_{i=1}^n (\text{Aug}(\vec{x}^{(ii)}) \cdot \vec{w} - y_i)^2 \right]$$

constant

$$= \arg \max_{\vec{w}} \left[-\frac{1}{n} \sum_{i=1}^n (\text{Aug}(\vec{x}^{(ii)}) \cdot \vec{w} - y_i)^2 \right]$$

$$= \arg \min_{\vec{w}} \left[\frac{1}{n} \sum_{i=1}^n (\text{Aug}(\vec{x}^{(ii)}) \cdot \vec{w} - y_i)^2 \right]$$

minimizing mean squared error.

Probabilistic view of regularization

$$\vec{w}^* = \arg \min_{\vec{w}^2} \frac{1}{n} \sum_{i=1}^n (H(\vec{x}^{ii}; \vec{w}^2) - y_i)^2 + \lambda \|\vec{w}^2\|^2.$$

helps control overfitting.

↑
assign belief and updating belief

we believe it is more likely to be small (close to zero) than large.

$$w_i \sim N(0, \sigma^2)$$

By buyer's rule:

$$\underset{\bar{w}^?}{\operatorname{argmax}} [p_{\bar{w}^?}(\bar{w}^? | \bar{x}^?, \mathcal{V})]$$

$$= \underset{\vec{w}}{\operatorname{argmax}} \ln [p_Y(y|\vec{w}, \vec{x}) p_{\vec{w}}(\vec{w})]$$

$$= \arg \min_{\vec{w}} [-\ln p_Y(y|\vec{w}, \vec{x}) - \ln p_{\vec{w}}(\vec{w})]$$

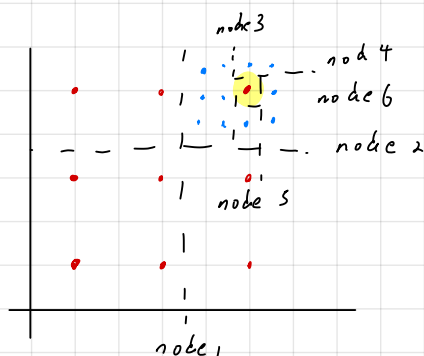
$$= \underset{\vec{w}}{\operatorname{argmin}} [\operatorname{MSE}(\vec{w}) - \ln p_{\vec{w}}(\vec{w})]$$

$$\downarrow$$

$$= C + \frac{1}{2\sigma^2} \|\vec{w}\|^2$$

Decision Tree

- A rooted tree
- Internal node ask yes/no question.
- Leaf node are decisions (class labels).
- Path from root is a sequence of "and"s:
- To make prediction $\rightarrow$ traverse tree.



start with single node containing all data points

- repeat greedy procedure

- look at all possible questions (splits)
- pick the one that most reduce uncertainty

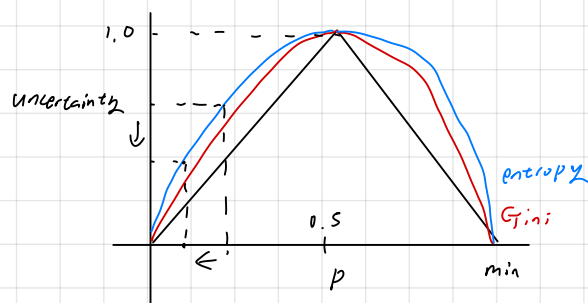
intuitive:

a good question splits the data by class.

Measuring Uncertainty:

Suppose our node contains proportions

p from class +
 $(1-p)$ from class -

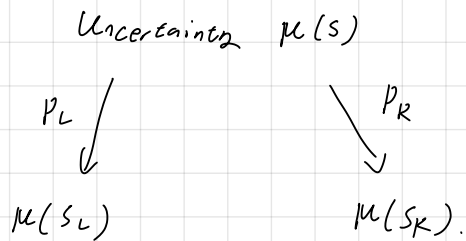


- Misclassification rate: $\min \{p, 1-p\}$
- Gini index: $2p(1-p)$
- Entropy: $p \log \frac{1}{p} + (1-p) \log \frac{1}{1-p}$

minimizes uncertainty.

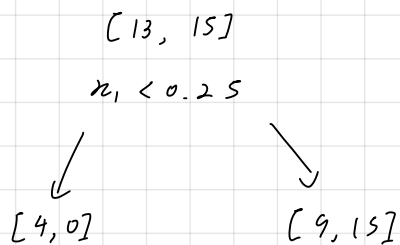
$\Downarrow$
 minimize proportion in one class
 p or $(1-p)$.

Let $\mu(S)$ be the uncertainty score for a set of labeled points, S .



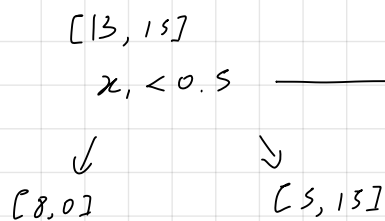
$$P_L \mu(S_L) + P_R \mu(S_R)$$

ex)



initial G_{int} : $2 \times \frac{13}{28} \times \frac{15}{28}$

$$P_L \mu(S_L) + P_R \mu(S_R) = \frac{4}{28} \cdot 0 + \frac{24}{28} \cdot 2 \cdot \frac{9}{24} \cdot \frac{15}{24} = \frac{45}{112}$$



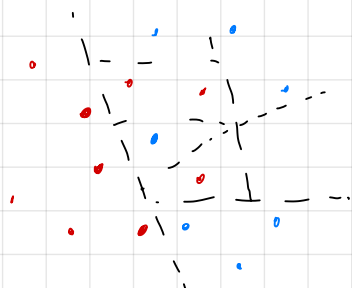
good
question.

minimize
uncertainty.

$$P_L \mu(S_L) + P_R \mu(S_R) = \frac{8}{28} \cdot 0 + \frac{20}{28} \cdot 2 \cdot \frac{5}{20} \cdot \frac{15}{20} = \frac{30}{112}$$

Recursively ask questions that minimize uncertainty.

Overfitting



separate perfectly for all training points

Regularization

[10, 20]



• Pruning: simplify already-constructed tree

• Early-stopping: stop early.



stop recursion when:

- node is "pure enough" (uncertainty is low).
- tree is too deep.

• replace node by most common class

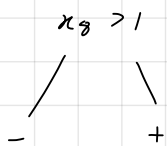
• if change reduces validation error. keep it, otherwise, reserve it.

• Decision tree properties:

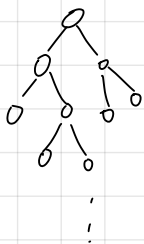
- can accommodate any type of data
- can accommodate any number of classes.
- can perfectly fit any data set.
- very expressive.

• Weak learner - one which performs only a little better than chance

• Decision stump - is a weak learner



• Full decision tree is a strong learner

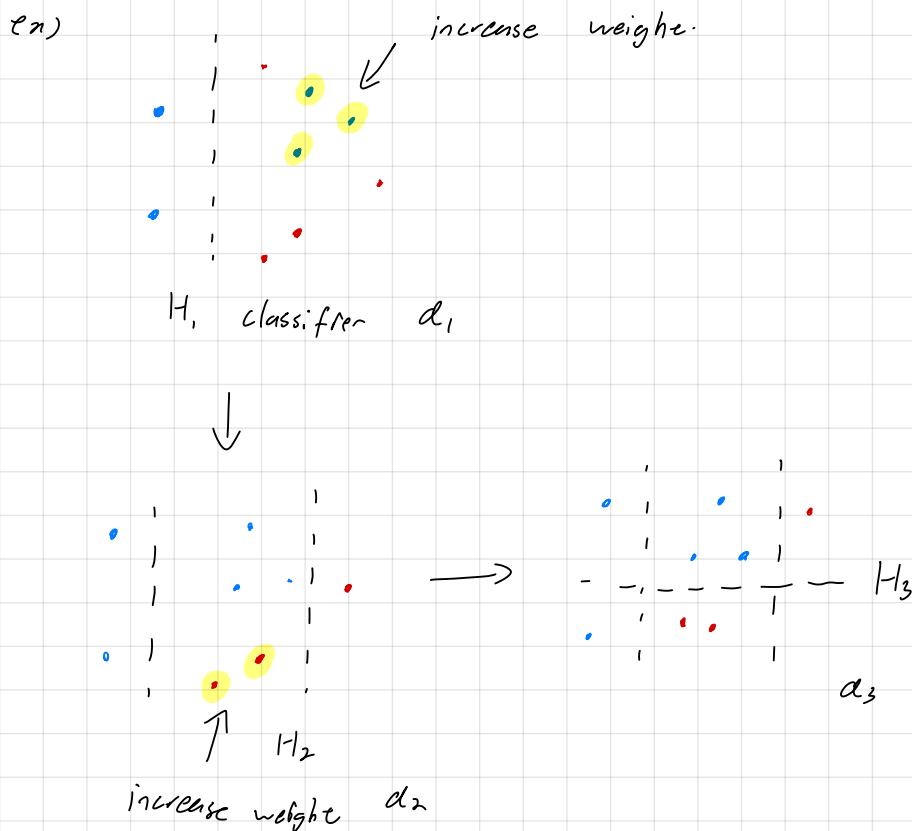


Boosting

• Train a weak classifier, $H_1: x \rightarrow [-1, 1]$.

• Increase weight of misclassified points, train another classifier H_2 .

- Repeat, creating more classifier, updating weights.
- Final classifier: a linear combination of $H_1, \dots, H_k$.



Final classifier: $\text{sign}(d_1 H_1(x) + d_2 H_2(x) + d_3 H_3(x))$.

AdaBoost

- Measuring performance:

normalized.

- suppose weights at step t are in $\vec{w}^{(t)}$,

- use weights to learn a classifier:

$$H_t: \mathcal{X} \rightarrow [-1, 1].$$

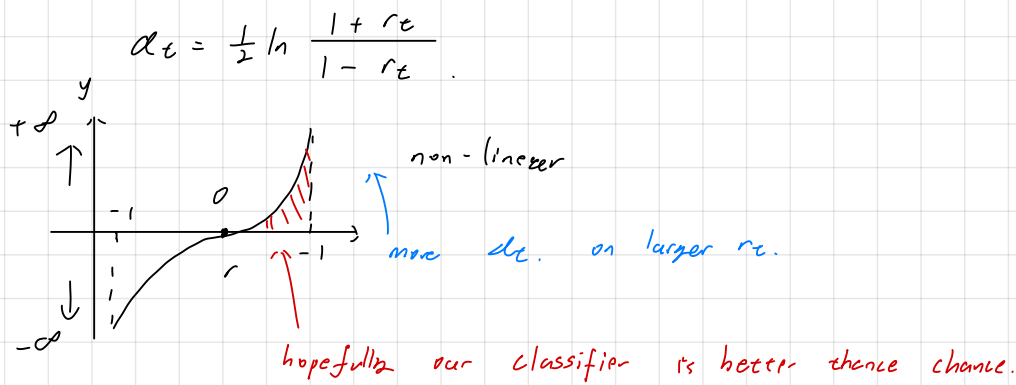
- The "margin":

$$r_t = \sum_{i=1}^n w_i^{(t)} y_i H_t(\vec{x}^{(i)}) \in [-1, 1].$$

more weight

on good prediction $\rightarrow$ good margin. $\rightarrow$ the larger r_t , the better H_t is doing on "important" points.
 on bad prediction less weight

- performance of H_t :



- updating weights:

weight misclassified point more heavily.

$$y_i H_t(\vec{x}^{(i)}) < 0$$

$$w_i^{(t+1)} \propto w_i^{(t)} \cdot \exp(-\alpha_t y_i H_t(\vec{x}^{(i)}))$$

b/c normalize w

misclassified $\rightarrow y_i H_t(\vec{x}^{(i)}) < 0$

$$-\alpha_t y_i H_t(\vec{x}^{(i)}) > 0$$

more weight

- Final classifier:

$$H_t(\vec{x}) = \sum_{t=1}^T \alpha_t H_t(\vec{x})$$

- Algorithm:

Given data $(\vec{x}^{(1)}, y), \dots, (\vec{x}^{(n)}, y_n)$, labels in $\{-1, 1\}$.

- initialize weight vector $\vec{w}^{(1)} = (\frac{1}{n}, \frac{1}{n}, \dots, \frac{1}{n})^T$

- repeat:

- given data and weights $\vec{w}^{(t)}$ to weak learner.
- receive a classifier, $H_t: \mathcal{X} \rightarrow \{-1, 1\}$ back.

$$\alpha_t = \frac{1}{2} \ln \frac{1+r_t}{1-r_t}$$

$$\vec{w}^{(t+1)} \propto \vec{w}^{(t)} \cdot \exp(-\alpha_t y_i H_t(\vec{x}^{(i)}))$$

$$H_t(\vec{x}) = \sum_{t=1}^T \alpha_t H_t(\vec{x})$$

To learn decision stump.

- try all features, threshold.

- choose that which maximizes the margin:

$$r_t = \sum_{i=1}^n w_i^{(t)} y_i; H_t(\vec{x}^{(i)}) \in [-1, 1].$$



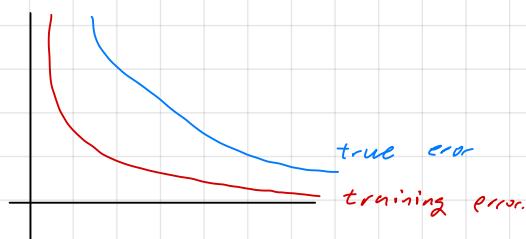
$\Leftrightarrow$ maximizes performance

$$\alpha_t = \frac{1}{2} \ln \frac{1+r_t}{1-r_t}$$

Theorem: Suppose on each round t , the weak learner returns a rule H_t , whose error on step t weighted data is $\leq \frac{1}{2} - \gamma$.

After T rounds, training error is at most $e^{-\gamma^2 T/2}$.

Boosted decision tree resistant to overfitting:



Why weak learner $\rightarrow$ fast to learn.

Random Forest

Boosted decision tree: $\rightarrow$ slow, train decision with all data / feature.

Prevent decision tree from overfitting by "hiding data" randomly.

Train a bunch of trees, quickly

Average them to make final prediction

For $t = 1$ to T

choose n' training points randomly, with replacement

Fit decision tree H_t .

At each node, restrict to one of k features, chosen randomly.

Final classifier: majority vote of $H_1, \dots, H_T$

common setting: $n' = n$ (bootstrap)

$$k = \sqrt{d}.$$