

CLASSIFIERS

$\{x: f(x)=0\}$: isosurface for 0
 Decision boundary: (d-1) dim surface
 linear C: decision b. line/plane

LINALG

unit vec: $\frac{x}{\|x\|}$
 $f(x) = w \cdot x + \alpha$, $H = \{w \cdot x = -\alpha\}$
 $\tilde{w} \perp H$, $w^T(y-x) = 0$, y, x on H
 $\tilde{w}$ unit vector $\Rightarrow f(x)$ signed dist from x to H
 α : Distance from origin

CENTROID METHOD: mean of μ_c, μ_{not-c}
 $f(x) = (\mu_c - \mu_{not-c}) \cdot x - \frac{\mu_c + \mu_{not-c}}{2}$

PERCEPTRON METHOD: for linearly separable
 uses gradient descent $O(n^2/y^2)$ iter
 $y_i = \begin{cases} 1 & \text{in } c \\ -1 & \text{not in } c \end{cases}$
 $x_i \cdot w \geq 0, y_i = 1$
 $x_i \cdot w \leq 0, y_i = -1$
 ADD 1 bias term
 $y_i x_i \cdot w \geq 0$

LOSS FUNCTION: $L(z, y_i) = \begin{cases} 0 & y_i \geq z \\ -y_i z & \text{otw} \end{cases}$

RISK: $\frac{1}{n} \sum L(x_i \cdot w, y_i) = \frac{1}{n} \sum -y_i x_i \cdot w$

$\nabla R(w) = -\sum y_i x_i$

while some $y_i x_i \cdot w < 0$, $w \leftarrow w + \sum y_i x_i$

Fictitious dim: Add α to w , 1 to x

SYM

margin: dist from decision b to nearest training point

dist from x_i : $\frac{w \cdot x_i + \alpha}{\|w\|}$

margin: $\min \frac{1}{\|w\|} |w \cdot x_i + \alpha| \geq \frac{1}{\|w\|}$

HARD MARGIN:

min $\|w\|^2$ (smooth func)
 s.t. $y_i(x_i \cdot w + \alpha) \geq 1$

Optimal: margin = $\frac{1}{\|w\|}$

Slab: $\frac{2}{\|w\|}$

linearly separable, outliers $\Rightarrow$

SOFT MARGIN:

min $\|w\|^2 + C \sum \xi_i$
 s.t. $y_i(x_i \cdot w + \alpha) \geq 1 - \xi_i$
 $\xi_i \geq 0$

small C: max margin $\frac{1}{\|w\|}$, min $\|w\|$, underfit
 less sensitive, more flat

big C: $\xi_i \rightarrow 0$, small margin, overfit
 very sensitive, approach hard margin

Parabolic Lifting Map: $\Phi(x) = \begin{bmatrix} x \\ \|x\|^2 \end{bmatrix}$

Ellipsoid: axis aligned, no cross terms
 $Ax^2 + Bx^2 + Cx^2 + Dx_1x_2 + \dots + Tx_3 + \alpha = 0$

$f(x) = [A \ B \ C \ \dots \ T] \cdot \Phi(x) + \alpha$
 quadric $\Rightarrow$ conic in 2D
 overfit if too high

Raising: linearly separable, wider boundary, robust

Edge Detection: Find regions with high gradients

- convex func: no min, 1 min, ∞ min
- strictly convex: 1 global min
- high ellipticity of contours ($\lambda_{max} \gg \lambda_{min}$), no good learning rate
- convex polytope satisfies linear program (if non empty, linearly separable)

DECISION THEORY

Bayes Theorem:
 $P(Y=1|X) = \frac{P(X|Y=1)P(Y=1)}{P(X)}$
 posterior

$R(r) = E[L(r(x), Y)]$
 $= \sum_x [L(r(x), 1)P(Y=1|X) + L(r(x), -1)P(Y=-1|X)]P(X)$

$r^*(x) = \begin{cases} 1 & \text{if } L(-1, 1)P(Y=1|X) > L(1, -1)P(Y=-1|X) \\ -1 & \text{otw} \end{cases}$

When L symmetric, pick class with biggest posterior

$E[g(x)] = \int g(x)f(x)dx$ $\text{Var}(x) = E[(x - E(x))^2]$

$\mu = E[x] = \int x f(x) dx$ $\sigma^2 = E[(x - \mu)^2] = E[x^2] - \mu^2$

Bayes Decision Boundary:
 $P(X|Y=1)P(Y=1) = P(X|Y=0)P(Y=0)$

Bayes Risk: Area under min of functions

GAUSSIAN DIS. ANALYSIS

$X \sim N(\mu, \sigma^2)$: $f(x) = \frac{1}{(\sqrt{2\pi}\sigma)^d} \exp\left(-\frac{\|x - \mu\|^2}{2\sigma^2}\right)$
 concave

Bayes maximises $f(X|Y=C)\pi_C$

Equivalently, $Q_C(x) = -\frac{\|x - \mu\|^2}{2\sigma^2} - d \ln \sigma + \ln \pi_C$

Models $P(Y=y|x)$ as a logistic $\frac{1}{1+e^{-s}}$

QDA SAMPLE variances

pick class to maximise Q_C

2 class: $Q_C(x) - Q_D(x) = 0$ (DB)

$P(Y=C|X) = \frac{f_{X|Y=C} \pi_C}{f_{X|Y=C} \pi_C + f_{X|Y=D} \pi_D} = s(Q_C - Q_D)$

2 class: DB: $w \cdot x + \alpha = 0$, $P(Y=C|x) = s(DB)$

MLE

$\hat{\mu} = \frac{1}{n} \sum x_i$ $\hat{\sigma}_2 = \frac{1}{n} \sum \|x_i - \hat{\mu}\|^2$

LDA: Same means and priors, one variance

$\hat{\sigma}_2 = \frac{1}{n} \sum \|x_i - \hat{\mu}_c\|^2$

QDA: Same mean $\hat{\mu}$, calc $\hat{\sigma}_c^2$, $\tilde{\pi}_c = \frac{n_c}{n}$

$\mathcal{L}(\mu, \sigma^2; X_1, \dots, X_n) = f(X_1)f(X_2)\dots f(X_n)$

$\mathcal{L}(\dots) = \ln f(X_1) + \dots + \ln f(X_n)$

$= \sum \left(-\frac{\|x_i - \mu\|^2}{2\sigma^2} - d \ln \sqrt{2\pi} - d \ln \sigma \right)$

EIG

$\vec{v}_i, \lambda_i$ for $A \rightarrow \vec{v}_i, \lambda_i^k$ for A^k
 $\rightarrow \vec{v}_i, \frac{1}{\lambda_i}$ for A^{-1}

Spectral Theorem: Real, symmetric M has real λ_i , n eigenvectors which are mutually orthogonal
 $M = V \Lambda V^T$ (UNIT eigenvectors)

PSD:
 $x^T A x \geq 0$, $\lambda_i \geq 0$, $A = U U^T$
 A^2 exists

VISUAL

$x^T A^{-1} x = \|A^{-1/2} x\|_2^2$ is an ellipsoid

axes $\vec{v}_i, \lambda_i$, radii $\lambda_i, \{A\}$

A diagonal: axis aligned

MULTIV. GAUSSIAN

$X \sim N(\mu, \Sigma)$

$f(x) = \frac{1}{\sqrt{(2\pi)^d |\Sigma|}} \exp\left(-\frac{1}{2}(x - \mu)^T \Sigma^{-1}(x - \mu)\right)$

Σ : covariance, Σ^{-1} : precision

Isocontours of $(x - \mu)^T \Sigma^{-1}(x - \mu)$ determined by eigenpairs of $\Sigma/2$

Covariance

$\text{Cov}(R, S) = E[(R - E[R])(S - E[S])^T]$
 $= E[RS^T] - \mu_R \mu_S^T$

$\text{Var}(R)$: PSD matrix

R_i, R_j ind $\Rightarrow \text{Cov}(R_i, R_j) = 0$

$\text{Cov}(R_i, R_j) = 0$ and multivariate norm $\Rightarrow R_i, R_j$ independent

All features ind $\Rightarrow \text{Var}(R)$ diagonal

$\text{Var}(R)$ diagonal + joint normal $\Rightarrow f(x) = f(x_1) \dots f(x_n)$

$\Rightarrow$ ellipsoids axis aligned, squared radii on Σ diagonal

ANISOTROPIC

LDA: $\frac{1}{n} \sum_{c=1}^C \sum_{x \in c} (x - \hat{\mu}_c)(x - \hat{\mu}_c)^T$

QDA: $\frac{1}{n} \sum_{c=1}^C \sum_{x \in c} (x - \hat{\mu}_c)(x - \hat{\mu}_c)^T$

$Q_C - Q_D$ is quadratic

DB quadratic

$-\frac{1}{2}(x - \mu_c)^T \Sigma_c^{-1}(x - \mu_c) - \frac{1}{2} \ln |\Sigma_c| + \ln \pi_c$

Multi-LDA: $\max \mu_c^T \Sigma^{-1} x + \frac{\mu_c^T \Sigma^{-1} \mu_c}{2} + \ln \pi_c$

LDA underfits, QDA overfits

$\tilde{X} = X - \mu$

$\text{Var}(x) = \frac{1}{n} \tilde{X}^T \tilde{X}$

Decorrelate X : $Z = X V$, $\text{Var}(R) = V \Lambda V^T$
 $X \sim N(\mu, \Sigma)$ $Z \sim N(0, \Lambda)$

Whitening: center + sphering
 sphering X : $W = X \text{Var } R^{-1/2}$

REGRESSION

1. linear: $w \cdot x + \alpha$
2. poly
3. logistic: $s(w \cdot x + \alpha)$

Loss functions

- A. $(z - y)^2$ squared
- B. $|z - y|$ absolute
- C. $-y \ln z - (1 - y) \ln(1 - z)$ logistic (0-1 labels)

Cost functions

- a. $\frac{1}{n} \sum L(h(x_i), y_i)$ mean
- b. $\max L(h(x_i), y_i)$ max
- c. $\sum w_i L(h(x_i), y_i)$ weighted
- d. $a/b/c + \lambda \|w\|^2$ ℓ_2
- e. $a/b/c + \lambda \|w\|_1$ ℓ_1

LEAST SQUARES

$$\min \sum_i (x_i \cdot w + \alpha - y_i)^2$$

$$\min \|Xw - y\|^2$$

$$\nabla \text{RSS}: X^T X w = X^T y \quad X^T X \text{ always PSD}$$

$$w = (X^T X)^{-1} X^T y$$

$\hat{g} = Xw = X X^T y = H y$. Ideal $H = I$
 Unique when $X^T X$ singular ✓
 Sensitive to outliers ✗

LOGISTIC REGRESSION cost func CONVEX
 $J = \sum_i L(s(x_i \cdot w), y_i)$
 $= - \sum_i (y_i \ln s_i + (1 - y_i) \ln(1 - s_i))$
 LINEAR ab for linear input

$s'(x) = s(x)(1 - s(x))$
 $\nabla_w J = -X^T (y - s)$ $\nwarrow$ PSD
 $\nabla_w^2 J = X^T \Omega X$ $\Omega = \text{diag}(s_i(1 - s_i))$
 BGD: $w \leftarrow w + \epsilon X^T (y - s(Xw))$
 SGD: $w \leftarrow w + \epsilon (y_i - s_i(x_i \cdot w)) X_i$
 separates linearly separable points

WEIGHTED LEAST SQUARES

$$\min (Xw - y)^T \Omega (Xw - y)$$

$$\Rightarrow X^T \Omega X w = X^T \Omega y$$

NEWTON'S METHOD

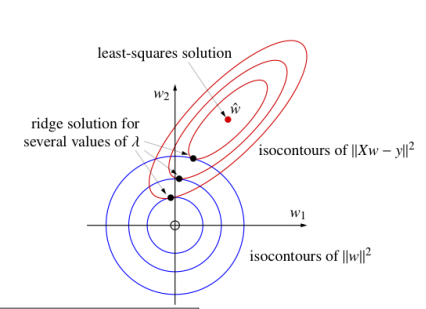
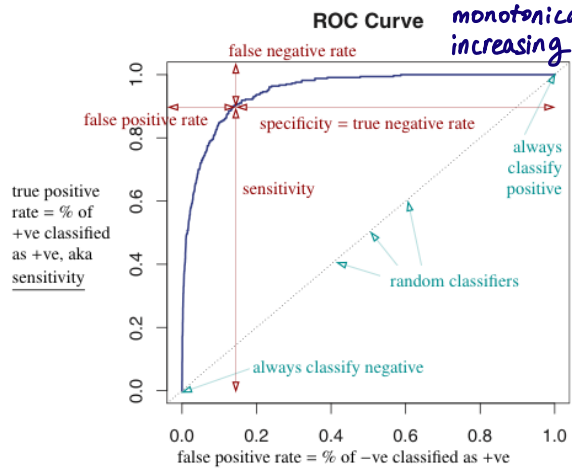
Iterative optimisation method for smooth $J(w)$
 Works when $J(w)$ is quadratic
 $w \leftarrow w - (\nabla^2 J)^{-1} (\nabla J)$
 $w \leftarrow w + e, (X^T \Omega X) e = X^T (y - s)$
 small weight in $W \Rightarrow$ more emphasis
 faster than SGD, esp when d small
 NOT on perceptron

$$\nabla_{\alpha} \hat{y} = (\nabla_{\alpha} \hat{y})^T + \alpha \nabla_{\alpha} \hat{y}$$

$$\nabla_{\alpha} f(y) = (\nabla_{\alpha} \hat{y}) (f'(\hat{y}))$$

$$\nabla_{\alpha} (\hat{y} \cdot \hat{z}) = (\nabla_{\alpha} \hat{y}) \cdot \hat{z} + (\nabla_{\alpha} \hat{z}) \cdot \hat{y}$$

$$\nabla_{\alpha} g(\hat{y}) = \nabla_{\alpha} (g(y)) \nabla_y (g(y))$$



Knob: posterior prob threshold for GDA/Log Reg
 Based on real test data

BIAS-VARIANCE

Bias: Inability of h fitting g
 Variance: Random noise in data
 Overfitting: Bias $\downarrow$ Variance $\uparrow$
 Underfitting: Bias $\uparrow$ Variance $\downarrow$
 + Feature: Variance $\uparrow$ Bias $=, \downarrow$
 $E[L(h(z)), y] = E[(h(z) - y)^2]$

$$= E[h(z)^2] + E[y^2] - 2E[yh(z)]$$

$$= \text{Var}(h(z)) + E[h(z)]^2 + \text{Var}(y) + E[y]^2 - 2E[yh(z)]$$

$$= \underbrace{(E[h(z)] - g(z))^2}_{\text{bias}^2} + \underbrace{\text{Var}(h(z))}_{\text{var}} + \underbrace{\text{Var}(y)}_{\text{noise}}$$



Variance more harmful to test error
 hyperparameters (λ, c, d) affect b-v

RIDGE

$\min \|Xw - y\|^2 + \lambda \|w'\|^2$
 ALWAYS unique optimum, ONLY one
 $\lambda \rightarrow \infty, \text{var} \rightarrow 0, \text{bias} \uparrow, w^* \rightarrow 0$
 $(X^T X + \lambda I) w = X^T y$
 $\nwarrow$ PD
 Bayesian Justification: $w' \sim N(0, c^2)$
 $f(w') \propto e^{-\|w'\|^2 / 2c^2}$
 $\max \ln f(y) + \ln f(w') - \text{const}$
 $= -\text{const} \|Xw - y\|^2 - \text{const} \|w'\|^2$
 $= \min \|Xw - y\|^2 + \lambda \|w'\|^2$
 Gaussian $w \sim \mathcal{N}(0, I)$ prior

LASSO

$\min \|Xw - y\|^2 + \lambda \|w'\|_1$, quadratic sparsity (intersects w axes) prog
FEATURE
 feature $\uparrow$, variance $\uparrow$
FORWARD: Start with 0 features, keep few $\hookrightarrow$ adding until val error $\uparrow$
 won't find best 2-f model
BACKWARD: Start with d , remove if more $\hookrightarrow$ removing $\downarrow$ val error
 Train data $\downarrow$, Train Accuracy $\uparrow$

Bayes classifier has nothing to do w sample points, needs to know abt distribution
 Ellipse eqn requires different quad coeffs
 $P(x \neq t) \leq \frac{E[\chi^2]}{t}$
 $\Sigma = E[(z - \mu)(z - \mu)^T]$
 $P(A|B) = P(A|B) P(B)$
 $R(r(x)) = i | \pi$
 $= \sum L(r(x) = i, y = j) P(y = j | \pi)$
 bias = $E[\hat{\theta}] - \theta$

$w^* = X^T y$ (least squares soln. of LEAST NORM)

LDA convex polytope (S subset of union of hyperplanes, S connected)

$$\text{rank}(X) + \dim(\text{N}(A)) = n$$

$$\dim(\text{row}(X)) = \dim(\text{col}(X^T))$$

$$= \text{rank}(X^T)$$

$$= \text{rank}(X)$$

transforming data changes scale
 shifting does nothing
 $\langle s, t \rangle \leq \|s\| \|t\|$
 $A \Rightarrow B, B \Rightarrow A, \neg A \Rightarrow \neg B$
 $\Sigma \text{ diag} \Rightarrow \text{diag}(\sigma_1^2, \sigma_2^2, \dots, \sigma_n^2)$

DECISION TREES

↳ classification + regression
↳ Tree with 2 nodes (internal node 1 feature + leaf)

Entropy: $H(s) = -\sum p_c \log_2 p_c$
Weighted Avg: $H_a = \frac{|S_L|}{|S_L| + |S_R|} H(s_L) + \frac{|S_R|}{|S_L| + |S_R|} H(s_R)$

Info gain: $H(s) - H_a$
Test: $\leq O(d)$, usually $\leq O(\log N)$
Train: $\leq O(nd)$

VARIATIONS
Regression: Average sample points at each node
Stopping Early: limits tree depth, tree size
Multivariate Split: Non-axis aligned splits w other algs. Classification ↑ Building Speed ↓
Ensemble Learning: Reduces variance in DTs
Bagging: Same learning alg. on random subsamples (WITH replacement), T learners
Random Forests: Bagging + feature subset m

Kernel $\Phi: \mathbb{R}^d \rightarrow \mathbb{R}^D$ d input features
Degree p-polynomials blow upto $O(d^p)$ features
In many algs, $w = X^T a = \sum a_i x_i$ (linear combo of pts)
and can use inner product of $\Phi(x)$ only.
Optimize n dual weights $\tilde{a}$ instead of Δ primal $\tilde{w}$
Solve nxn linear system instead of dxd
Kernel Ridge Regression:
($X^T X + \lambda I$) $\tilde{w} = X^T y$
($X X^T + \lambda I$) $a = y$
Then $\tilde{X} y = (X^T X + \lambda I) X^T a$ [works if you center X]
↳ $w = X^T a$ is a solution.
Obj: $\min_a \|X^T X a - y\|^2 + \lambda \|X^T a\|^2$
Train: Solve $(X X^T + \lambda I) a = y$
Test: $h(z) = w^T z = a^T X z = \sum a_i (X_i^T z)$

Kernel fn: $K(x, z) = x^T z$. $K = X X^T (n \times n)$
 $K_{ij} = K(X_i, X_j)$
Alg: $K_{ij} \leftarrow K(X_i, X_j)$ $O(n^2 d)$
Solve $(K + \lambda I) a = y$ $O(n^3)$
Test: $h(z) = \sum a_i K(X_i, z)$ $O(nd)$

Polynomial Kernel: Degree p in
 $K(x, z) = (x^T z + 1)^p = \phi(x)^T \phi(z)$
Compute $\phi(x)^T \phi(z)$ in $O(d)$ instead of $O(d^p)$
SGD dual: $a_i \leftarrow a_i + \epsilon (y_i - s(K a))$

Gaussian Kernel: generate feature vecs in ∞ space
 $\phi: \mathbb{R}^d \rightarrow \mathbb{R}^\infty$. $K(x, z) = e^{-\frac{\|x - z\|^2}{2\sigma^2}}$ $O(d)$
inner products of similar features large, small for dissimilar features

Key: $h(z) = \sum a_j K(X_j, z)$ is linear combo of Gaussians centered at training pts
Dual weights coefficients of linear combo

Why? Very smooth h, like K-NN but smoother, oscillates less than polynomials
larger $\sigma \rightarrow$ wider gaussian $\rightarrow$ bias ↑↑ variance ↓↓

NEURAL NETWORKS

↳ classification + regression
↳ perceptrons failed @ XOR

1 hidden layer:
Input: $x_1, \dots, x_d$ Layer 1 weights: $m \times (d+1)$
Hidden: $h_1, \dots, h_m$ Layer 2 weights: $K \times (m+1)$
Output: $z_1, \dots, z_K$

Training: SGD/Batch. Cost function NOT convex
Initialize random weights to break symmetry.
Backprop: $O(\text{edges})$, dp alg
Vanishing Gradient: s close to 0 or 1, s' $\rightarrow 0$
Gradients change too slowly. Layers ↑ Vanishing ↑

Sigmoid $\rightarrow$ ReLU: $\max\{0, y\}$, might explode
Output Unit

① Regression: linear output units + squared loss
② Classification: classes ≥ 2 , softmax units
 $z_i(t) = \frac{e^{t_i}}{\sum_j e^{t_j}}$, cross entropy $-\sum y_i \ln z_i$

linear/ReLU $\rightarrow$ squared error
sigmoid/softmax $\rightarrow$ cross entropy

OPTIMIZATIONS

Fix vanishing gradient $\rightarrow$ SGD on large sets
Normalize training points: center, then scale
Use small rand. subsample to choose ϵ
Dif ϵ for dif layers of weights
Emphasize schemes (repeat rare examples)
Acceleration scheme for SGD
Add enough weights, test error ↓↓ "double descent"

Heuristics:
① Weight decay: Add $\lambda \|w\|^2$ to loss
step $\Delta w_i = -\epsilon \frac{\partial J}{\partial w_i}$ has extra term $-2\epsilon \lambda w_i$
② Ensemble of NNs or dropout in 1 NN (p=0.5)

CNN

Local Connectivity $\rightarrow$ hidden unit connects to small patch from previous layer
Shared Weights $\rightarrow$ group of hidden units share set of weights "kernel/filters"
Convolution: same linear transf. applied to different patches of input by shifting

out = $1 + \frac{\text{input} - \text{kernel height} + 2p}{\text{stride}}$

① Local Connectivity $\rightarrow$ hidden unit connects to small patch from previous layer
② Shared Weights $\rightarrow$ group of hidden units share set of weights "kernel/filters"

Convolution: same linear transf. applied to different patches of input by shifting

Given q, find K training points near q. dist metric up to you. Regression (any label) or classification (majority or histogram w probabilities)
nn $\uparrow \rightarrow$ underfitting | $n \rightarrow \infty$, 1-NN err $< 2B - B^2$
nn $\downarrow \rightarrow$ overfitting | 2-clones, $\leq 2B - B^2$
draw pts independently from same prob. distribution.
 $n \rightarrow \infty, K \rightarrow \infty, K/n \rightarrow 0$, err converges to B

Exhaustive K-NN Alg: Given q, scan through all n points, computing squared distances to q. Maintain a max-heap with K shortest distances. Every time $O(nd + n \log K)$
2-S d $\rightarrow$ Voronoi
Medium d $\rightarrow$ k-d trees
Large d $\rightarrow$ exhaustive k-NN can use PCA/rand. proj

Voronoi: X point set. Vor $w = \{ \|p - w\| \leq \|p - v\|, \forall v \in X \}$
Voronoi cell always convex polyhedron/polytope
Size $\in O(n^{d+1})$, often $O(n)$
For 2D: $O(n \log n)$ to compute v.d. + trapezoidal map for pc. location. $O(\log n)$ query time.
d-D: Use binary space partition tree $\epsilon = 0 \rightarrow$ exact $\epsilon > 0 \rightarrow$ approx
K-d Trees:

1. choose splitting feature w greatest width. ft. ϵ in $\max(X_j - X_k)$, make boxes cubical
2. choose splitting value. median pt. for feature or $(X_j + X_k)/2$. Median guarantees $O(\log n)$ depth. $O(nd \log n)$ building time. Each internal node stores a training pt.
Goal: Given q, find w so $\|q - w\| \in (1+\epsilon)\|q - s\|$
Maintain: K-nearest so far, binary min-heap of dist^2

NEAREST NEIGHBOURS

Given q, find K training points near q. dist metric up to you. Regression (any label) or classification (majority or histogram w probabilities)
nn $\uparrow \rightarrow$ underfitting | $n \rightarrow \infty$, 1-NN err $< 2B - B^2$
nn $\downarrow \rightarrow$ overfitting | 2-clones, $\leq 2B - B^2$
draw pts independently from same prob. distribution.
 $n \rightarrow \infty, K \rightarrow \infty, K/n \rightarrow 0$, err converges to B

Exhaustive K-NN Alg: Given q, scan through all n points, computing squared distances to q. Maintain a max-heap with K shortest distances. Every time $O(nd + n \log K)$
2-S d $\rightarrow$ Voronoi
Medium d $\rightarrow$ k-d trees
Large d $\rightarrow$ exhaustive k-NN can use PCA/rand. proj

Voronoi: X point set. Vor $w = \{ \|p - w\| \leq \|p - v\|, \forall v \in X \}$
Voronoi cell always convex polyhedron/polytope
Size $\in O(n^{d+1})$, often $O(n)$
For 2D: $O(n \log n)$ to compute v.d. + trapezoidal map for pc. location. $O(\log n)$ query time.
d-D: Use binary space partition tree $\epsilon = 0 \rightarrow$ exact $\epsilon > 0 \rightarrow$ approx
K-d Trees:

1. choose splitting feature w greatest width. ft. ϵ in $\max(X_j - X_k)$, make boxes cubical
2. choose splitting value. median pt. for feature or $(X_j + X_k)/2$. Median guarantees $O(\log n)$ depth. $O(nd \log n)$ building time. Each internal node stores a training pt.
Goal: Given q, find w so $\|q - w\| \in (1+\epsilon)\|q - s\|$
Maintain: K-nearest so far, binary min-heap of dist^2

Voronoi: X point set. Vor $w = \{ \|p - w\| \leq \|p - v\|, \forall v \in X \}$
Voronoi cell always convex polyhedron/polytope
Size $\in O(n^{d+1})$, often $O(n)$
For 2D: $O(n \log n)$ to compute v.d. + trapezoidal map for pc. location. $O(\log n)$ query time.
d-D: Use binary space partition tree $\epsilon = 0 \rightarrow$ exact $\epsilon > 0 \rightarrow$ approx
K-d Trees:

1. choose splitting feature w greatest width. ft. ϵ in $\max(X_j - X_k)$, make boxes cubical
2. choose splitting value. median pt. for feature or $(X_j + X_k)/2$. Median guarantees $O(\log n)$ depth. $O(nd \log n)$ building time. Each internal node stores a training pt.
Goal: Given q, find w so $\|q - w\| \in (1+\epsilon)\|q - s\|$
Maintain: K-nearest so far, binary min-heap of dist^2

Voronoi: X point set. Vor $w = \{ \|p - w\| \leq \|p - v\|, \forall v \in X \}$
Voronoi cell always convex polyhedron/polytope
Size $\in O(n^{d+1})$, often $O(n)$
For 2D: $O(n \log n)$ to compute v.d. + trapezoidal map for pc. location. $O(\log n)$ query time.
d-D: Use binary space partition tree $\epsilon = 0 \rightarrow$ exact $\epsilon > 0 \rightarrow$ approx
K-d Trees:

1. choose splitting feature w greatest width. ft. ϵ in $\max(X_j - X_k)$, make boxes cubical
2. choose splitting value. median pt. for feature or $(X_j + X_k)/2$. Median guarantees $O(\log n)$ depth. $O(nd \log n)$ building time. Each internal node stores a training pt.
Goal: Given q, find w so $\|q - w\| \in (1+\epsilon)\|q - s\|$
Maintain: K-nearest so far, binary min-heap of dist^2

UNSUPERVISED LEARNING

Sample points, no labels $\rightarrow$ discover structure in data

PCA

Dimensionality reduction: find K directions that capture most of the variation
Why: reducing # dimensions, remove irrelevant dimensions, find small basis for repr $\frac{x^T w}{\|w\|^2}$
 $X^T X$: square, symmetric, PSD, (λ_i, v_i) : corresponding unit principal components

① Fit Gaussian to data with MLE
Choose K Gaussian axes of greatest variance
MLE estimates $\hat{\Sigma} = \frac{1}{n} X^T X$

1. Center X, normalize X
2. Compute unit eigenvectors of $X^T X$
3. Choose K, choose $v_{d-K+1}, \dots, v_d$
4. Compute $\tilde{x} \cdot \tilde{v}_i$ for each training point

② Find direction $\tilde{w}$ that maximizes sample variance of projected data
 $\max_w \frac{1}{n} \sum (X_i \cdot \frac{w}{\|w\|})^2 = \frac{1}{n} \frac{\|X w\|^2}{\|w\|^2} = \frac{1}{n} \frac{w^T X^T X w}{w^T w}$
 $\rightarrow v_d$ achieves max variance λ_d / n

③ Find direction $\tilde{w}$ that minimizes mean squared projection distance
(proj orthogonal to proj hyperplane) $\rightarrow$ PCA

$\min_w \sum_i \|X_i - \tilde{X}_i\|^2 = \sum_i \left\| X_i - \frac{X_i \cdot w}{\|w\|^2} w \right\|^2$
 $= \sum_i \left(\|X_i\|^2 - \left(\frac{X_i \cdot w}{\|w\|} \right)^2 \right)$
: constant - n (Var from ②)

SVD
Problem: Computing $X^T X$ takes $\Theta(nd^2)$ time, eigenvectors inaccurate
 $X = U D V^T$. When $n \geq d$
 $X = U D V^T = \sum_i \delta_i u_i v_i^T$
 $\delta_i = \text{rank } i$
outer product matrix
 $U^T U = I = V^T V$ (orthonormal)
D (diagonal)

D: $\delta_1, \dots, \delta_d \rightarrow$ Nonnegative, singular values of X
Non-zero $\delta_i = \text{rank}(X)$
 $X^T X = V D V^T U D V^T = V D^2 V^T$ (v_i has eigenvalue δ_i^2)
Find K greatest singular values in $O(n \log K)$ time
Since $X = U D V^T \Rightarrow X V = U D \Rightarrow \text{Row}_i(U D) = X_i \cdot v_j$
principal coordinates

CLUSTERING
Partition data into clusters so points in clusters are more similar than across clusters
Why? Discovery, hierarchy, graph partitioning, quantization

K-means clustering: partition n points into K disjoint clusters
Cluster i mean: $\mu_i = \frac{1}{n_i} \sum X_j$

Find $\min \sum_i \sum_{j \in Y_i} \|X_j - \mu_i\|^2$ squared distances from points to means
NP-hard, solve in $O(nK^2)$ time

CLUSTERING
Partition data into clusters so points in clusters are more similar than across clusters
Why? Discovery, hierarchy, graph partitioning, quantization

K-means clustering: partition n points into K disjoint clusters
Cluster i mean: $\mu_i = \frac{1}{n_i} \sum X_j$

Find $\min \sum_i \sum_{j \in Y_i} \|X_j - \mu_i\|^2$ squared distances from points to means
NP-hard, solve in $O(nK^2)$ time

CLUSTERING
Partition data into clusters so points in clusters are more similar than across clusters
Why? Discovery, hierarchy, graph partitioning, quantization

K-means clustering: partition n points into K disjoint clusters
Cluster i mean: $\mu_i = \frac{1}{n_i} \sum X_j$

Find $\min \sum_i \sum_{j \in Y_i} \|X_j - \mu_i\|^2$ squared distances from points to means
NP-hard, solve in $O(nK^2)$ time

CLUSTERING
Partition data into clusters so points in clusters are more similar than across clusters
Why? Discovery, hierarchy, graph partitioning, quantization

K-means clustering: partition n points into K disjoint clusters
Cluster i mean: $\mu_i = \frac{1}{n_i} \sum X_j$

Find $\min \sum_i \sum_{j \in Y_i} \|X_j - \mu_i\|^2$ squared distances from points to means
NP-hard, solve in $O(nK^2)$ time

CLUSTERING
Partition data into clusters so points in clusters are more similar than across clusters
Why? Discovery, hierarchy, graph partitioning, quantization

K-means clustering: partition n points into K disjoint clusters
Cluster i mean: $\mu_i = \frac{1}{n_i} \sum X_j$

Find $\min \sum_i \sum_{j \in Y_i} \|X_j - \mu_i\|^2$ squared distances from points to means
NP-hard, solve in $O(nK^2)$ time

CLUSTERING
Partition data into clusters so points in clusters are more similar than across clusters
Why? Discovery, hierarchy, graph partitioning, quantization

K-means heuristic:

1. Fixed y_j , Update μ_i
 2. Fixed μ_i , Update y_j
- Halt when step ② changes no labels.
Both steps minimize cost, but not all vars.
Alg never returns to prev assignment. MUST terminate. Finds local minimum

Starting:

- ↳ Forgy Method: Choose K rand points to be initial $\mu_i \rightarrow$ Go to step ② better
- ↳ Random Partition: Randomly assign pts to a cluster $\rightarrow$ Go to step ①

Equivalent: Min within-cluster variation

$$\min_y \sum_{i=1}^K \frac{1}{n_i} \sum_{y_j: y_m=i} \|X_j - X_m\|^2$$

K-Medoids Clustering:

- ↳ Specify distance fn $d(x, y)$
- ↳ Replace mean with medoid - pt that min total distance to other pts in same cluster
- ↳ Medoid ALWAYS input point, more robust to outliers

HIERARCHICAL CLUSTERING

- K no longer a hyperparameter $O(n^3)$ time
- ① Point Set $\rightarrow$ Bottom-up 'Agglomerative' start w a point, keep joining upwards
- ② Graphs $\rightarrow$ Top-down 'Divisive' start w everything in a cluster, keep splitting

Complete Linkage: $\max \{d(a, b)\}$

Single linkage: $\min \{d(a, b)\}$

Average Linkage: $\frac{1}{|A||B|} \sum_{a \in A} \sum_{b \in B} d(a, b)$

Centroid linkage: $d(\mu_A, \mu_B)$ *median distance*

HIGH DIMENSIONS

- ↳ Intuition wrong in higher dim
- ↳ $\|p\|^2 = p_1^2 + p_2^2 + \dots + p_d^2$ in d dim
 $p_i \sim N(0, 1)$ $p_i^2 \sim \chi^2(1)$ $E[p_i^2] = 1$ $\text{Var}(p_i^2) = 2$
 $\|p\|^2 \sim \chi^2(d)$
- ↳ For large d $\|p\|$ is concentrated in thin shell around radius $\sqrt{d}$ with thickness $\propto \sqrt{d}$
- ↳ In high dim, 1- nn and 1000- nn don't differ much $\Rightarrow$ K-means & nn not as effective
- $\cos \theta = \frac{p \cdot q}{\|p\| \|q\|}$ $E[\cos \theta] = 0$, $SD(\cos \theta) = \frac{1}{\sqrt{d}}$

↳ d large $\rightarrow \cos \theta$ tends to 0 / θ tends to 90°
 ↳ d large $\Rightarrow$ 2 rand vectors almost always $\perp$

Random Projection: All to PCA, preserves dist. b/w points. Project very high- d space to rand medium- d subspace

$K = \left\lceil \frac{2 \ln(1/\delta)}{\epsilon^2/2 - \epsilon^2/3} \right\rceil$. For some point q ,

$\tilde{q}$ is orthogonal proj of q onto S , multiplied by $\sqrt{\frac{d}{K}}$
 Choosing projection direction: choose each component from a univariate Gaussian dist. then normalise vectors to unit length. K rand directions, then Gram-Schmidt orthogon.

PSEUDOINVERSE

D : diagonal $n \times d$ matrix
 Transpose $D \rightarrow$ replace nonzero entry w reciprocal- D^T
 $DD^T D = D$ $D^T DD^T = D^T$
 SVD of $X = UDV^T$. $\text{rank } D = \text{rank } X$
Moore-Penrose Pseudoinv. of X : $X^+ = VD^+U^T$

1. $XX^+ = UDV^T VD^+U^T = U(DD^+)U^T$, PSD
 all eigenvalues 1 or 0
 2. $X^+X = VD^+U^T UDV^T = V(D^+D)V^T$, PSD
 3. $D, D^+, DD^+, D^+D, X, X^+, XX^+, X^+X$, ranks equal
 4. X has rank n , $XX^+ = I_n$, X^+ right inv.
 5. X has rank d , $X^+X = I_d$, X^+ left inv.
 6. $XX^+X = X$ [$UDD^+U^T UDV^T = UDV^T = X$]
 7. $X^+XX^+ = X^+$ [symmetric proof ↑]
- Theorem: $X^T X w = X^T y$; $w = X^+ y$
 Proof: $X^T X w = X^T X X^+ y = VD^+U^T UDV^T VDD^+U^T y = VD^+D^+U^T y = VD^+U^T y = X^+ y$

If normal eqns. have multiple solns, then $w = X^+ y$ is the least norm soln (minimises $\|w\|$ among all solutions)

↳ Useful when $X^T X$ is singular

LEARNING THEORY

Range Space / Set System: Pair (P, H) where
 P is set of all possible test/training pts
 H is set of hypotheses/classifiers
 ↳ $h \in P$ that specifies which points h predicts are in class C

1. Power-set classifier: H power set of P containing 2^K subsets of P
2. Linear classifier: $P = \mathbb{R}^d$, H is set of all halfspaces $\{x: w \cdot x \geq -\alpha\}$ / every possible linear classifier in d dimensions

Power-set classifier sucks at learning
Risk (Generalization Error): $R(h)$ of h is probability that h misclassifies a random point x drawn from D

∞ test data $\Rightarrow$ risk = test error
Empirical Risk (Training Error): $\hat{R}(h)$ is % of X misclassified by h
 $n \rightarrow \infty$, $\hat{R}(h)$ approximates $R(h)$

Hoeffding's Inequality:
 $\Pr(|\hat{R}(h) - R(h)| > \epsilon) \leq 2e^{-2\epsilon^2 n}$
 Choose $\hat{h} \in H$ that minimizes $\hat{R}(\hat{h})$

Dichotomies: Dich of X is $X \wedge h$, $h \in H$
 For n points, 2^n dichotomies.
 Dichotomy assigns points to class C or not C
 If H allows all 2^n , $\hat{R}(\hat{h}) = 0$
 H induces Π dichotomies.

$\Pr(\text{at least one dich has } |\hat{R} - R| > \epsilon) \leq \delta$,
 $\delta = 2\Pi e^{-2\epsilon^2 n}$
 $|\hat{R}(h) - R(h)| \leq \epsilon = \sqrt{\frac{1}{2n} \ln \frac{2\Pi}{\delta}}$

Smaller Π , larger $n \Rightarrow$ training error $\approx$ true risk (variance)
 Smaller $\Pi \Rightarrow$ less likely to overfit (bias)
 Sample complexity: #training pts needed for prob(ϵ)
 $n \geq \frac{1}{2\epsilon^2} \ln \frac{2\Pi}{\delta}$

SHATTER FUNCTION

of dichotomies: $\Pi_H(X) = |\{X \wedge h: h \in H\}|$, $n \in |X|$
 Shatter function: $\Pi_H(n) = \max_{|X|=n, X \subseteq P} \Pi_H(X)$ $\Pi_H(2) = 3$
 $\Pi_H(n) = 14$

For all range spaces, $\Pi_H(n)$ is polynomial in n , or $\Pi_H(n) = 2^n$

Cover's Theorem: linear classifiers in $\mathbb{R}^d$ allow $\Pi_H(n) = 2 \sum_{i=0}^d \binom{n-1}{i}$ dichotomies of n points
 $n \leq d+1$, $\Pi_H(n) = 2^n$ $n \geq d+1$, $\Pi_H(n) \approx 2 \left(\frac{en}{d}\right)^d$

To achieve $R(\hat{h}) \leq \hat{R}(h) + \epsilon \leq R(h^*) + 2\epsilon$ w.p. $\geq 1-\delta$.
 $n \geq \frac{1}{2\epsilon^2} \left(d \ln \frac{1}{\delta} + d + \ln \frac{1}{\delta} \right)$. Linear classifiers only need $n \in \Theta(d)$ training points for train err $\approx$ test err

VC DIM. $Vc(H) = \max \{n: \Pi_H(n) = 2^n\}$

H shatters set X of n pts if $\Pi_H(X) = 2^n$. $Vc(H)$ is largest X H can shatter. (X allows all 2^n dichotomies)
 $\Pi_H(n) \leq \sum_{i=0}^n \binom{n}{i}$. $n \geq Vc(H)$, $\Pi_H(n) \leq \left(\frac{en}{Vc(H)}\right)^{Vc(H)}$
 Vc -dim upper bound on exponent of polynomial $O(Vc(H))$ pts. suffice for accuracy
 Vc finite $\Rightarrow$ sample complexity grows w. # features
 $Vc \infty \Rightarrow$ no amt. of training data will make it generalise well
 $Vc(H) = 3$. $\Pi_H(n) \leq \frac{e^3}{27} n^3$. $O(1)$ sample complexity

ADA-BOOST

ensemble, classif + regression
 reduces bias *correct* $w \downarrow$
 Train T classifiers $G_1, \dots, G_T$ *wrong* $w \uparrow$
 Weight for X_i in G_t grows according to how many classifiers misclassified it
 Train G_t to try harder to correctly classify X_i with $w \uparrow$
 Metalearner is linear combo of learners
 Test point z , $M(z) = \sum B_t G_t(z)$. G_t is ± 1 , M continuous. Return sign of $M(z)$.

Iteration T : Find G_T, B_T so
 $\min \text{Risk} = \frac{1}{n} \sum_{i=1}^n L(M(X_i); y_i)$, $L(X_i) = \sum_{t=1}^T B_t G_t(X_i)$
 Metalearner uses exponential loss:
 $L(p, d) = e^{-p \cdot d} = \begin{cases} e^{-p} & d = +1 \\ e^p & d = -1 \end{cases}$ *emphasizes misclassifs*

cost fn. for G_t = # misclassified points
 $n \cdot \text{Risk} = \sum_{i=1}^n e^{-y_i M(X_i)} = \sum_{i=1}^n \prod_{t=1}^T e^{-B_t y_i G_t(X_i)}$
 $= e^{-B_T \sum w_i^T} + e^{B_T \sum w_i^T}$
 $w_i^{t+1} = \begin{cases} w_i^t e^{-B_T y_i G_t(X_i)} & \text{correct} \\ w_i^t e^{B_T} & \text{misclassified} \end{cases}$
 B_T . Set $\frac{\partial}{\partial B_T} (\text{Risk}) = 0$
 $B_T = \frac{1}{2} \ln \left(\frac{1 - \text{err}_T}{\text{err}_T} \right)$ $\text{err}_T = 0$, $B_T = 0$ (perfect learner)
 $\text{err}_T = 1/2$, $B_T = 0$

Algorithm:
 1. Initialize $w_i \leftarrow 1/n$
 2. For $t \leftarrow 1$ to T . 1) Train G_t with weights w_i 2) Calc. err
 3) Reweight points $w_i \leftarrow w_i \cdot e^{B_t}$ or $w_i \cdot e^{-B_t}$
 3. Return $M(z) = \text{sign} \left(\sum B_t G_t(z) \right)$

Why?: Reduces bias reliably, no hyperparameter search needed, fast, short trees + ADA = feature subset selection
 Linear decision boundaries don't boost well

More: $P(Y=1|z) \approx 1/(1 + e^{-2M(z)})$
 Exponential loss vulnerable to outliers
 If all learners beat error $< 50\%$, train err $\rightarrow 0$ wolly