

What do we need to learn from data?

- a pattern $\leftarrow$ not hard
- not solvable analytically $\leftarrow$ not hard.
- data

Lecture 2

Fundamentals of learning feasibility

Is learning feasible?

probabilistic! can we extrapolate?

bin - red & green marbles

$$\begin{array}{l} \text{P}[\text{red}] = \mu \\ \text{P}[\text{green}] = 1 - \mu \end{array}$$

μ is unknown

N picks independently to gen. sample

fraction of red in sample is v

does v help with μ ? yes obv but is
sampling problem

possible vs probable $\leftarrow$ almost certain is
good enough

What does v say about μ ?

in large N , v converges to μ within ϵ

FORMALLY: $P[|v - \mu| > \epsilon] \leq 2e^{-2\epsilon^2 N}$
 ϵ squared term kills N

HOEFFDING'S INEQUALITY

all extrapolation is driven by tolerance and sample size. " $\mu = v$ " is probably approximately correct.

- valid for all N and ϵ
- bound independent of μ
- tradeoffs between N and ϵ and the bound.

smaller N , bigger ϵ

- $v \approx \mu \Rightarrow \mu \approx v$ (when you run exp. μ is unknown. μ affects v , not vice versa, but symmetric so our exps work)

HOW DOES THIS CONNECT TO LEARNING?

unknown μ which is a function

$$f: X \rightarrow y$$

each marble is point $x \in \mathcal{X}$

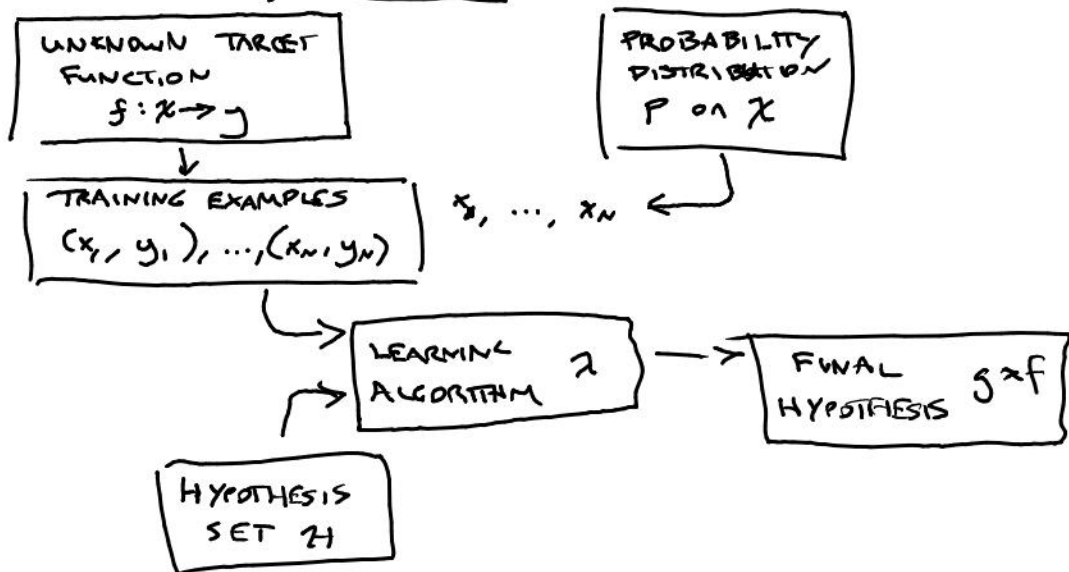
bin is our input space, so is $\mathcal{X}$

$\circ: h(x) = f(x)$ (HYPOTHESIS CORRECT)

$\bullet: h(x) \neq f(x)$

so we have a mapping for how accurate we are. introduces probability to learning
we are not restricting probability dists.
we don't have to know what it is.
hoeffding lets us do this because it's all sample & tolerance.

DIAGRAM



but we are not done!

h is fixed!

for this h , v generalises to μ

THIS IS VERIFICATION, NOT LEARNING

DOES MY HYPOTHESIS WORK, NOT IS THIS OPTIMAL

WE NEED TO FIND HYPOTHESIS IN WHOLE SPACE

NO guarantee that v is small and we're
picking from H

multiple bins? bin $h_1, \dots, h_N$ generating
different μ
try to drive $n \rightarrow \infty$

is learning just picking the right bin?

seems silly for searches through infinite
hypothesis space

v (in sample) denoted by $E_{in}(h)$

μ (out of sample) denoted by $E_{out}(h)$

$$P[|E_{in}(h) - E_{out}(h)| > \epsilon] \leq 2e^{-2\epsilon^2 N}$$

1 SFE S M E T H I N K B A D C O M I N G !

Hoeffding doesn't apply to multiple bins
bin number drives probability of sample w/in
tolerance, as bins $\rightarrow$ we guarantee bad
hypothesis is found!

if you toss 1 coin 10 times, prob of
10 heads is $\sim 0.1\%$

if you toss 1000 coins 10x each, $P(10H) = 63\%$
i'm just quickly calc'ing...

$$1 - (1 - 0.5^{10})^{1000} \rightarrow 62.35\% \text{ OK}$$

Solution! S is final hypothesis

$$P[|E_n(h) - E_{\text{out}}(h)| > \epsilon] \leq P[|E_n(h) - E_{\text{out}}(h_1)| > \epsilon]$$

UNION BOUND

$$\dots$$
$$|E_n(h_m) - E_{\text{out}}(h_m)| > \epsilon]$$

$$\leq \sum_{m=1}^M P[|E_n(h_m) - E_{\text{out}}(h_m)| > \epsilon]$$

THIS SUCKS THO?

$$P[|E_n(h) - E_{\text{out}}(h)| > \epsilon] \leq \sum_{m=1}^M 2e^{-2\epsilon^2 N}$$

this is a generalisation problem! bigger N decouples
in-sample from out

PROBLEM SET:

1) d

2) a

3) d (if black bull 2/3 chance it's in the 2-black bag)

$$4) (1 - 0.55)^{10} = 3.405 \times 10^{-4} \quad b$$

$$5) 1 - (1 - (1 - 0.55)^{10})^{1000} \Rightarrow 0.288 \quad c$$

$$6) X = \{0, 1\}^3$$

$$y_n = f(x_n) \text{ for } n=1..5$$

	x_n			y_n
D	0	0	0	0
	0	0	1	1
	0	1	0	1
	0	1	1	0
	1	0	0	1
ID	1	0	1	?
	1	1	0	?
	1	1	1	?

WRONG. WHY? $\rightarrow$ [C] $\leftarrow$ but I don't feel great about it

- 7) b
- 8) c
- 9) b
- 10) b

LECTURE 3: THE LINEAR MODEL 1

g is final hypothesis

$$P[|E_{in}(g) - E_{out}(g)| > \epsilon] \leq M 2e^{-2\epsilon^2 N}$$

WORST CASE, UNION BOUND

FOR INDEPENDENT EVENTS $\leftarrow$ OBY LINE OF ATTACK

THIS IS THE PRINCIPLE THAT THROUGH LEARNING
WE CAN GENERALISE, NOT THE FINAL RESULT

LET'S DO SOMETHING PRACTICAL BEFORE FINISHING
THEORY

- INPUT REPRESENTATION
- LINEAR CLASSIFICATION
- LINEAR REGRESSION
- NON-LINEAR TRANSFORMATION

A REAL DATA SET: [DATA 16x16 IMAGE OF NUMBERS]

HUMAN ERROR ~ 2.5%



HOW DO WE REPRESENT INPUT?

$16 \times 16 \Rightarrow 256$ pixels

raw input $x = (x_0 \dots x_{256})$

linear model: $(w_0 \dots w_{256})$ 256 dimensional space is too big!

INPUT REPRESENTATION IS SIMPLIFICATION

FEATURES: EXTRACT USEFUL INFORMATION

MAYBE INTENSITY AND SYMMETRY

$$x = (x_0, x_1, x_2)$$

WE LOSE INFORMATION, BUT PROBABLY

USELESS INFORMATION, AND WE'VE

DROPPED 254 DIMENSIONS

SO $x = (x_0, x_1, x_2)$ x_0 INTENSITY x_2 SYMMETRY

ACCEPT SOME ERROR!

PLA: EVALUATION OF E_{in} & E_{out} BY ITERATION

OUR DATA IS NOT LINEARLY SEPARABLE $\leftarrow$ NEVER CONVERGES

CUT OFF AFTER 1K RUNS. FINAL HYPOTHESIS IS NOT GOOD. LET'S MODIFY

THE 'POCKET' ALGORITHM $\leftarrow$ PICKS 'BEST' INTERMEDIATE HYPOTHESIS (E_{in} MIN).
MUCH BETTER

LINEAR REGRESSION

CLASSIFICATION: CREDIT APPROVAL (Y/N)

REGRESSION: CREDIT LINE (dollar amount)

INPUT x = age
 salary
 ...
 debt

LINEAR REGRESSION

OUTPUT: $h(x) = \sum_{i=0}^d w_i x_i = w^T x$

MAKES ALGORITHM
LINEAR! $\nearrow$

DATASET: HISTORICAL DECISIONS

$(x_1, y_1) \dots (x_n, y_n)$

$y_n \in \mathbb{R}$ is credit line for customer n

HOW TO MEASURE ERROR

how well does $h(x) = w^T x$ approximate $f(x)$?

use squared error in linear regression $(h(x) - f(x))^2$

(just standard rather than considered... for now)

$$n\text{-sample error: } E_n(h) = \frac{1}{N} \sum_{n=1}^N (h(x_n) - y_n)^2$$

LINEAR REGRESSION TRIES TO PRODUCE A LINE
THAT MINIMISES ERROR. HYPERSPLANE OF D-1

$$E_n(w) = \frac{1}{N} \sum_{n=1}^N (w^T x_n - y_n)^2$$

$$= \frac{1}{N} \|Xw - y\|^2$$

where

$$X = \begin{bmatrix} -x_1^T - \\ -x_2^T - \\ \dots \\ -x_n^T - \end{bmatrix} \quad y = \begin{bmatrix} y_1 \\ y_2 \\ \dots \\ y_n \end{bmatrix}$$

FEW PARAMETERS, MANY
EXAMPLES

MINIMISING

$$E_n(w) = \frac{1}{N} \|Xw - y\|^2$$

$$\nabla E_n(w) = \frac{2}{N} X^T (Xw - y) = 0 \leftarrow \text{VECTORS!}$$

WHERE DOES THE TRANSPOSE COME FROM?

INVERTIBLE, $\rightarrow X^T X w = X^T y$
PROBABLY

$$w = X^+ y \quad \text{where} \quad X^+ = (X^T X)^{-1} X^T$$

$\uparrow$
THE PSEUDO-INVERSE OF X

$$\left(\underbrace{\begin{bmatrix} \end{bmatrix}}_{d+1 \times N} \underbrace{\begin{bmatrix} \\ \\ \end{bmatrix}}_{N \times d+1} \right)^{-1} \underbrace{\begin{bmatrix} \end{bmatrix}}_{d+1 \times N}$$

$\Downarrow$

$$\left(\underbrace{\begin{bmatrix} \\ \end{bmatrix}}_{d+1 \times d+1} \right)$$

THE LINEAR REGRESSION ALGORITHM

1. CONSTRUCT THE MATRIX X AND VECTOR y
FROM THE DATA SET $(x_1, y_1) \dots (x_n, y_n)$

$$X = \begin{bmatrix} -x_1^T & - \\ \vdots & \\ -x_n^T & - \end{bmatrix}, \quad y = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix}$$

input data matrix target vector

2. COMPUTE THE PSEUDO-INVERSE $X^+ = (X^T X)^{-1} X^T$
3. RETURN $w = X^+ y$

WHY IS X INVERTIBLE? WOULD HAVE TO BE
SUPER UNLUCKY TO HAVE ALL YOUR DATA POINTS
DEPENDENT ON EACH OTHER

LINEAR REGRESSION FOR CLASSIFICATION

LINEAR REGRESSION LEARNS A REAL-VALUED FUNCTION

$$y = f(x) \in \mathbb{R}$$

BINARY FUNCTIONS ARE REAL-VALUED! $\pm 1 \in \mathbb{R}$

USE LINEAR REGRESSION TO GET w WHERE $w^T x_n \approx y_n \in \pm 1$
IN THIS CASE, $\text{sign}(w^T x_n)$ IS LIKELY TO AGREE W/ $y_n \in \pm 1$
GOOD MIT WEIGHTS FOR CLASSIFICATION

NON LINEAR TRANSFORMATION

LINEAR IS LIMITED

DATA: MESSY, NON-SEPARABLE. CAN WE FIND A TRICK?



CREDIT LINE

years in residence

nonlinear $[x_i < 1]$ and $[x_i > 5]$

can we do this with linear models

LINEAR IN WHAT?

linear regression implements

$$\sum_{i=0}^d w_i x_i$$

linear classification implements

$$\text{sign}\left(\sum_{i=0}^d w_i x_i\right)$$

linear in x , yes but more importantly

linear in w !

data is transformable because weights given to non-linear features have linear dependency
transform from constant to other constant to make new constant separable

THERE IS A CATCH!!

WHERE X^T COMES FROM WHEN
WE DIFFERENTIATE $\|Xw - y\|^2$

LET'S IGNORE CHAIN RULE FIRST &
EXPAND

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

$$\underbrace{w^T X^T X w}_{(1)} - \underbrace{w^T X^T y}_{(3)} - \underbrace{y^T X w}_{(3)} - \underbrace{y^T y}_{(4)}$$

d/dw of each term

① quadratic matrix differentiation

$d/dx (x^T A x) = 2Ax$ when A is symmetric.

$X^T X$ is always symmetric.

$$\therefore d/dw (w^T X^T X w) = \underline{2X^T X w}$$

② SCALAR TRANSPOSE TRICK! THE SCALAR

PRODUCED BY $w^T a$ IS THE SAME AS

$a^T w$. SO $w^T X^T y$ IS THE FORM $a^T w$,

DERIVATIVE IS $X^T y$. ③ SAME TRICK

④ TRIVIAL

$$\text{①..④ : } 2X^T(Xw - y)$$

TWO KEY THINGS:

- 1) MATRIX MULT IS NOT COMMUTATIVE
BUT WHEN RESULT IS SCALAR WE CAN
TRANSPOSE AND FLIP
- 2) DIFFERENTIATING A LINEAR SCALAR FUNCTION
 $a^T w$ WITH RESPECT TO w RETURNS THE
COEFFICIENT VECTOR a — THE THING MULTIPLYING
 w , EXTRACTED & RETURNED IN VECTOR FORM.
THE TRANSPOSE APPEARS IN DIFFERENTIATION WORK
BECAUSE EXTRACTING LEFT-SIDED x HERE
REQUIRED TRANSPOSITION

$$d/dw (Xw) = X^T$$

$$d/dw (w^T X) = X$$

CHECK MATRIX SHAPE BEFORE PERFORMING
OPERATIONS!!

ROW MULTIPLICATION

$$\begin{bmatrix} - & - & - \end{bmatrix} \begin{bmatrix} - & - & - \\ - & - & - \\ - & - & - \end{bmatrix} \checkmark$$
$$\begin{bmatrix} - & - & - \\ - & - & - \end{bmatrix} \begin{bmatrix} - & - & - \end{bmatrix} \times$$

MATRIX MULTIPLICATION RULE

COLS OF FIRST MUST EQUAL

ROWS OF SECOND

RESULT WILL BE ROWS OF FIRST $\times$
COLS OF SECOND

IF MULTIPLICATION IS INVALID THE ARE
TRICKS -

- TRANSPOSITION
- RESHAPING
- BROADCASTING
- PSEUDO-INVERSE

LECTURE 4

$$\text{signal: } \sum_{i=0}^d w_i x_i = w^T x$$

$$w = \underbrace{(x^T x)^{-1} x^T}_{x^\dagger} y \quad (\text{one-step learning})$$

transforms $w^T x$ linear in w , so x can
be transformed $\rightarrow x$ is a stack of constants
 $(x_1, x_2) \xrightarrow{\phi} (x_1^2, x_2^2)$ might help separate
data

ERROR AND NOISE

NON LINEAR TRANSFORMATION

$$\text{original: } x_n \in \mathcal{X} \quad \boxed{\begin{matrix} \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot \end{matrix}} \xrightarrow{\phi} \boxed{\begin{matrix} \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot \end{matrix}}$$

$$\text{transformed: } z_n = \phi(x_n) \in \mathcal{Z}$$

new data is linearly separable. then we
apply ' ϕ^{-1} ' to $\hat{g}(z) = \text{sign}(\tilde{w}^T z)$
in x -space: $\hat{g}(x) = \text{sign}(\tilde{w}^T \phi(x))$

WHAT TRANSFORMS TO WHAT?

$$x = (x_0 \dots x_d) \xrightarrow{\phi} \underbrace{z(z_0 \dots z_d)}$$

POTENTIALLY HARD TO GENERALISE

$$x_1 \dots x_n \xrightarrow{\phi} z_1 \dots z_n$$

$$y_1 \dots y_n \xrightarrow{\phi} y_1 \dots y_n \leftarrow \text{untouched}$$

$$\text{No weights in } \mathcal{X} \quad \tilde{w} = (w_0 \dots w_d)$$

SO WHAT IS HYPOTHESIS?

$$g(x) = \text{sign}(\tilde{w}^T \phi(x))$$

ERROR MEASURES AND NOISY TARGETS



ERROR MEASURE

what does $h \neq f$ mean?

ERROR MEASURE: $E(h, f)$ ← we try to minimise

ALMOST ALWAYS POINTWISE DEFINITION: $c(h(x), f(x))$

EG: SQUARED ERROR: $c(h(x), f(x)) = (h(x) - f(x))^2$

BINARY " : $c(h(x), f(x)) = \begin{bmatrix} h(x) \neq f(x) \end{bmatrix}$
T/F

1 if true, 0 if false

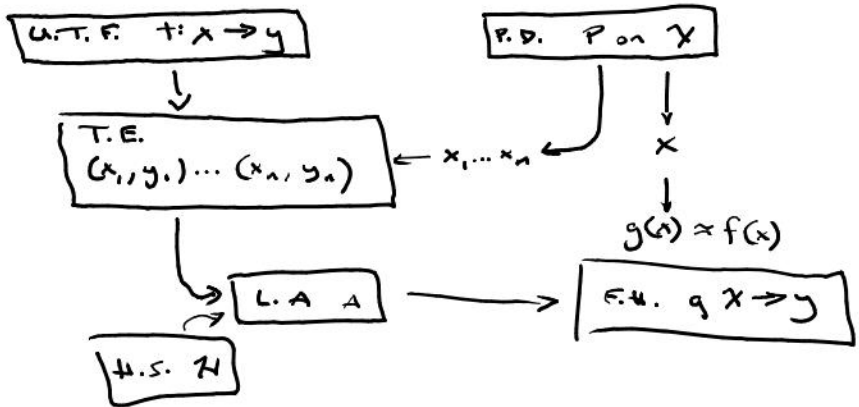
POINTWISE TO OVERALL → OVERALL IS AVERAGE OF POINTWISE ERRORS

In-sample: $E_{in}(h) = \frac{1}{N} \sum_{n=1}^N c(h(x_n), f(x_n))$
↑ LOOKS LIKE SOMETHING WE'VE SEEN

OUT OF SAMPLE:

$$E_{out}(h) = \mathbb{E}_x [e(h(x), f(x))]$$

ADD POINTWISE ERROR TO LEARNING DIAGRAM.



ERROR COMES FROM PROBABILITY DISTRIBUTION THAT DRIVES TRAINING DATA. OR HOPEFULLY DOESN'T WORK

CHOOSE ERROR MEASURE

FINGERPRINT VERIFICATION

$$p_{\text{out}} \rightarrow f \rightarrow \begin{cases} +1 & \text{You} \\ -1 & \end{cases}$$

error types: false positive, false negative

how to penalise?

		f	
		+1	-1
h	+1	no error	false pos
	-1	false neg	no error

← no inherent merit in picking one error function over another

IN A SUPERMARKET

COST OF FALSE POSITIVE LOW

COST OF FALSE NEGATIVE HIGH

		f	
		+1	-1
h	+1	0	1
	-1	1	0

IN THE CIA

COST OF FALSE POSITIVE HIGH

COST OF FALSE NEGATIVE LOWER

		f	
		+1	-1
h	+1	0	1000
	-1	1	0

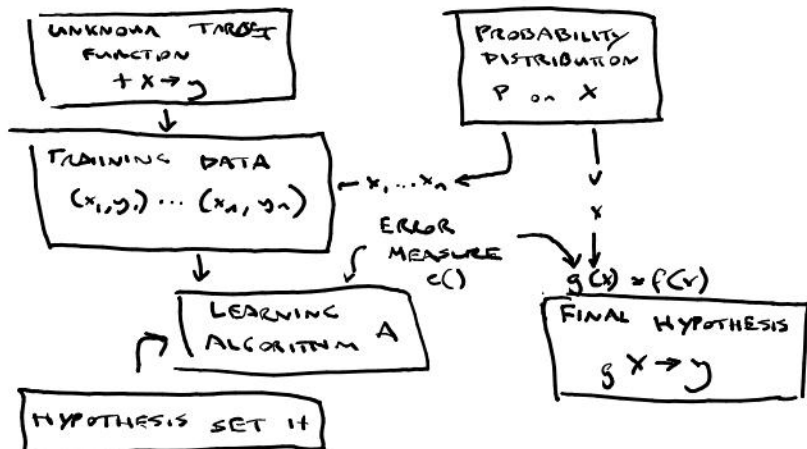
PRACTICAL LEARNING PROBLEM $\rightarrow$ ERROR MEASURE SPECIFIED BY USER

NOT ALWAYS POSSIBLE. WHAT ARE ALTERNATIVES?

PLAUSIBLE MEASURES: squared error $\hat{=}$ (gaussian noise)

FRIENDLY MEASURES: closed-form solution, convex optimisation

MODIFY LEARNING DIAGRAM AGAIN



NOISY TARGETS

THE TARGET FUNCTION IS NOT ALWAYS A FUNCTION
(OBSOLETE!)

CREDIT APPROVAL EXAMPLE $\rightarrow$ TWO IDENTICAL CUSTOMERS
IN TRAINING SET HAVE
DIVERGENT OUTCOMES

INSTEAD OF $y = f(x)$ WE USE TARGET DISTRIBUTION
 $P(y|x)$

(x, y) is now generated by the joint distribution
 $P(x)P(y|x)$

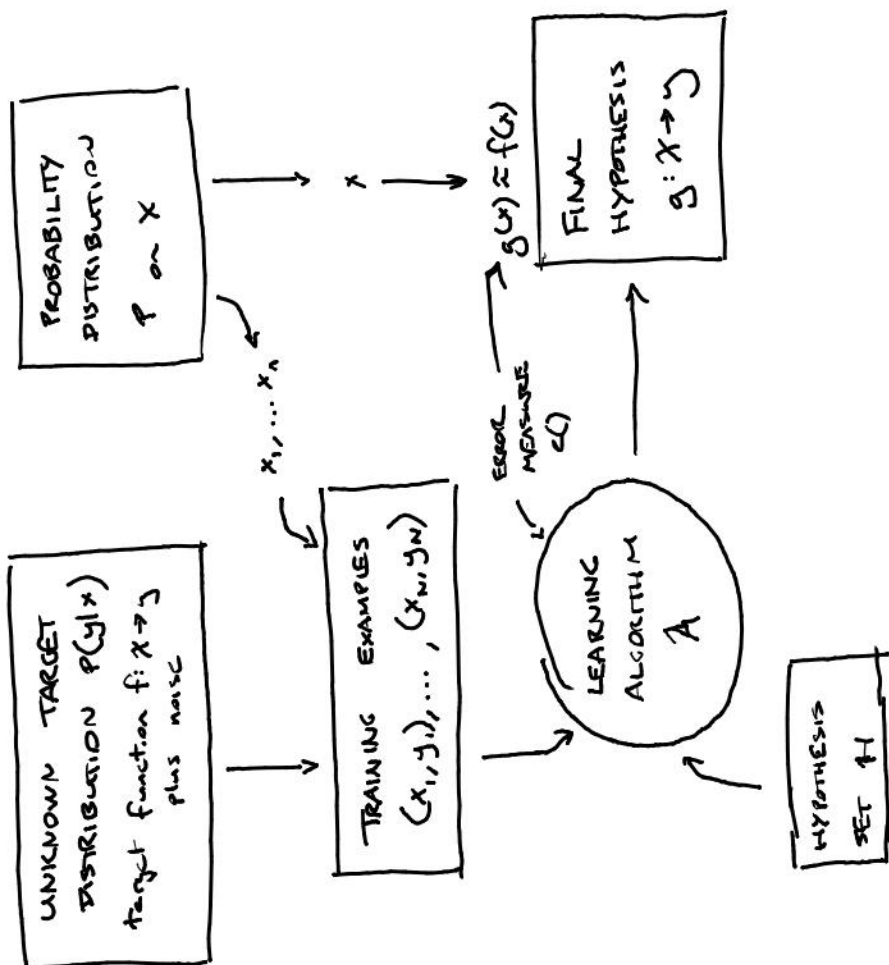
noisy target = deterministic target $f(x) = E(y|x)$ plus noise
 $y - f(x)$

deterministic function is a special case of

noisy target $P(y|x) = 0$ unless $y = f(x)$

FINAL LEARNING DIAGRAM!

(for supervised learning)



DISTINCTION BETWEEN $P(y|x)$ AND $P(x)$

BOTH TARGET DISTRIBUTION AND $P(x)$ ARE UNKNOWN
CONVEYING PROBABILISTIC ASPECTS OF x AND y

TARGET DISTRIBUTION $P(y|x)$ IS WHAT WE'RE
TRYING TO LEARN.

THE INPUT DISTRIBUTION $P(x)$ QUANTIFIES RELATIVE
IMPORTANCE OF POINT x . WE'RE NOT TRYING TO
LEARN IT

MERCING $P(x)P(y|x)$ AS $P(x,y)$ MERGES CONCEPTS.
DANGEROUS!

PREAMBLE TO THEORY OF LEARNING

• Learning is feasible. it is likely that

$$E_{out}(g) \approx E_{in}(g)$$

is this learning?

we need $g \neq f$, which means $E_{in}(g) \neq 0$

so what is $E_{out}(g) \approx E_{in}(g)$? generalisation

$E_{in}(g)$ is a proxy

THE 2 QUESTIONS

$E_{out}(g) \approx$ ACHIEVED THROUGH

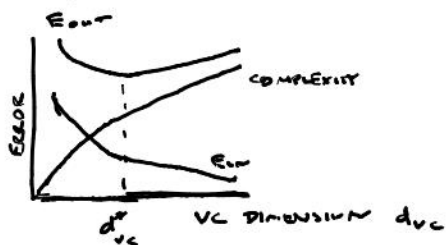
$$E_{out}(g) \approx E_{in}(g) \quad \text{and} \quad E_{in}(g) \approx 0$$

LEARNING IS THEREFORE TWO QUESTIONS.

1. CAN WE MAKE SURE $E_{out}(g)$ IS CLOSE TO $E_{in}(g)$

2. CAN WE MAKE $E_{in}(g)$ SMALL ENOUGH?
BIASED TERM!

WHAT THE THEORY WILL ACHIEVE



characterise the feasibility of learning for
infinite M

characterise the tradeoff $\leftarrow$ REGULARIZATION

model complexity	$\uparrow$	E_{in}	$\downarrow$
"	$\uparrow$	$E_{out} - E_{in}$	$\uparrow$

HOMEWORK #2

1. b

2. d

3. OK WHEN $1-\mu$ ON λ

" μ ON $1-\lambda$

ERROR WHEN μ ON λ

$1-\mu$ ON $1-\lambda$

e

4. b $\leftarrow$ if your function is binary and you have a 0.5 error you're just guessing!

5. c

6. c

7. a $\leftarrow$ obvious but I should build

8. d

9. -1, 0, 0, 0.05, 1.6, 1.5 $\leftarrow$ a

10. b

LECTURE 5: TRAINING AND TESTING

ERROR MEASURES SHOULD BE USED SPECIFIED

BUT WE CAN ALSO LOOK ANALYTICALLY/PRACTICALLY.

THEN WE PLUG INTO IN-SAMPLE AND OOS, E_y IS
EXPECTED VALUE ACROSS ^{known} PROBABILITY DIST.

REVISION

NOISE FACTOR $\rightarrow P(y|x)$ makes target distribution

$$E_{\text{out}}(h) = E_{x,y} [e(h(x), y)]; \quad (x_1, y_1) \dots (x_n, y_n) \text{ general}$$

by $P(x, y) = P(y|x) \cdot P(x)$

THIS LECTURE

- FROM TRAINING TO TESTING
- KEY NOTION: BREAK POINT

TRAINING TO TESTING

- REVISION! SAMPLE PROBLEMS TO PREPARE YOU FOR A TEST. WHY DON'T WE JUST PRACTICE ON THE FINAL? B/C FINAL IS JUST WAY OF GAUGING LEARNING, NOT THE GOAL
- WHAT DOES THIS LOOK LIKE FORMALLY?

TESTING: PLAIN Hoeffding

$$\mathbb{P}[|E_n - E_{out}| > \epsilon] \leq 2 e^{-2\epsilon^2 N}$$

TRAINING: M-modified

$$\mathbb{P}[|E_n - E_{out}| > \epsilon] \leq 2 M e^{-2\epsilon^2 N}$$

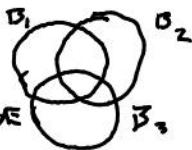
WE WOULD LIKE TO GET RID OF M BECAUSE M IS ALWAYS INFINITE!

WHERE DOES M COME FROM?

BAD EVENTS B_n ARE WHEN E_n DOES NOT TRACK E_{out} : " $|E_n(h_n) - E_{out}(h_n)| > \epsilon$ "

THEN UNION BOUND APPLIED, BUT OBV BAD EVENTS ARE NOT INDEPENDANT!

U.S. $\mathbb{P}[B_1 \text{ or } B_2 \text{ or } \dots B_n]$



UNION BOUND IS CONSERVATIVE

$\sum_{n=1}^N \mathbb{P}[B_n] \leftarrow$ ASSUMES NO OVERLAP

NON-CONSERVATIVE BOUND NEEDS TO CHARACTERISE THE OVERLAP

CAN WE IMPROVE ON M ?

YES! BAD EVENTS VERY OVERLAPPING.

PERCEPTRON HAS INFINITE M BUT THEY'RE
TIGHTLY CORRELATED

FOR SMALL LINE CHANGE...



ΔE_{out} : change in area of error regions
 ΔE_{in} : change in labels of data points } BOTH VERY
AND
CORRELATED

$$|E_{in}(h_n) - E_{out}(h_n)| \approx |E_{in}(h_{n+1}) - E_{out}(h_{n+1})|$$

OVERLAPPING!

M REPLACEMENT TIME!

COMPLEXITY! INSTEAD OF WHOLE INPUT
SPACE, RESTRICT ATTENTION TO
THE SAMPLE. CONSIDER FINITE
INPUT, AND COUNT NUMBER OF
TH. WE CALL THESE DICHOTOMIES

DICHOTOMIES: $\min_i -h$

$$h: \mathcal{X} \rightarrow \{-1, +1\} \leftarrow \text{hypothesis}$$

$$h: \{x_1, \dots, x_n\} \rightarrow \{-1, +1\} \leftarrow \text{dichotomies}$$

$|H|$ can be infinite for hypotheses

number of dichotomies $|H(x_1, x_2, \dots, x_N)| \rightarrow 2^N$

$$= 2 \cdot 2^N - 2\epsilon^2 N \leftarrow \text{HOEFFDING}$$

THE GROWTH FUNCTION

THE GROWTH FUNCTION COUNTS MAX DICHOTOMIES
FOR ANY HYPOTHESIS SET GIVEN N POINTS

$$m_H(N) = \max_{x_1, \dots, x_N \in \mathcal{X}} \underbrace{|H(x_1, \dots, x_N)|}_{\substack{\text{SET OF} \\ \text{DICHOTOMIES} \\ \text{(COUNT)}}}$$

SATISFIES

$$m_H(N) \leq 2^N$$

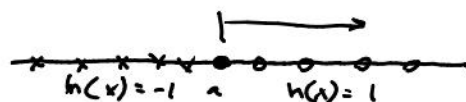
$$m_H(3) = 8$$

$$m_H(4) = 14 \leftarrow \text{WE CAN'T REACH SOME CONFIGURATIONS}$$

BREAK POINT! SET OF ALL POSSIBLE HYPOTHESES

STRONGER THAN H PERCEPTRON

POSITIVE RAYS



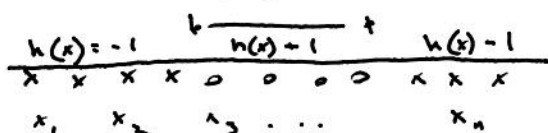
$x_1 \quad x_2 \quad x_3 \quad \dots \quad x_n$

$\mathcal{H}$ is set of $h: \mathbb{R} \rightarrow \{-1, +1\}$

$$h(x) = \text{sign}(x - a)$$

$$m_{\mathcal{H}}(N) = N + 1$$

POSITIVE INTERVALS



$\mathcal{H}$ is set $h: \mathbb{R} \rightarrow \{-1, +1\}$

$$m_{\mathcal{H}}(N) = \binom{N+1}{2} + 1 = \frac{1}{2}N^2 + \frac{1}{2}N + 1$$

CONVEX SETS

$\mathcal{H}$ is set of $h: \mathbb{R}^2 \rightarrow \{-1, +1\}$

$h(x) = +1$ is convex

$$m_{\mathcal{H}}(N) = ?$$



COLLAPSE INTO CIRCLE SO GROWTH FUNCTION IS

$$2^N$$

WEAKNESS OF GROWTH FUNCTION AS MAX IS
THAT BOUND IS NOT AS TIGHT AS POSSIBLE GIVEN
TRAINING DATA.

"SHATTERING" IS WHEN YOU GET ALL POSSIBLE
PHILOSOPHIES

GROWTH FUNCTIONS

POSITIVE RAYS: $N+1$

" INTERVALS: $\frac{1}{2}N^2 + \frac{1}{2}N + 1$

CONVEX SET: 2^N

MORE SOPHISTICATED $\rightarrow$ BIGGER GROWTH FUNCTION

WHAT HAPPENS WHEN YOU REPLACE M WITH $m_H(N)$?

$$P[|E_{in} - E_{out}| > \epsilon] \leq 2M e^{-2\epsilon^2 N}$$

vs $-2\epsilon^2 N$

$$2m_H(N) e^{-2\epsilon^2 N}$$

if $m_H(N)$ polynomial the exponential wins
WE NEED THIS!

BREAK POINT!

THE POINT AT WHICH AN $\mathcal{H}$ SET CAN'T
REACH ALL POSSIBLE HYPOTHESES.

IF NO DATASET OF SIZE k CAN BE SHATTERED
BY $\mathcal{H}$, THEN k IS A BREAK POINT FOR $\mathcal{H}$

$$m_{\mathcal{H}}(k) < 2^k$$

→ PERCEPTRON? $k=4$

BREAK POINT DRIVES LEARNING BEHAVIOR!

POSITIVE RAYS $k=2$

INTERVALS $k=3$

CONVEX SET $k=\infty$

MAIN RESULT

NO BREAK POINT $\Rightarrow m_{\mathcal{H}}(N) = 2^N$

ANY " $\Rightarrow m_{\mathcal{H}}(N)$ is

WASN'T TOLD
US WHY
YET
polynomial in N

SO ANY HYPOTHESIS SET WITH A BREAK POINT
IS GENERALISABLE!

LECTURE 6: THE THEORY OF GENERALIZATION

• DICHOTOMIES: RESTRICT ATTENTION TO HYPOTHESES
THAT IMPACT DATASET (FINITE)

GROWTH FUNCTION: $m_H(N) = \max_{x_1, \dots, x_N \in X} |H(x_1, \dots, x_N)|$
most possible dichotomies

BREAK POINT: DATA POINTS THAT CANNOT BE
SHATTERED BY HYPOTHESES.

CHARACTERISES GENERALISATION

- ① • PROOF THE GROWTH FUNCTION $m_H(N)$ IS POLYNOMIAL
- ② • PROOF THAT $m_H(N)$ CAN REPLACE M

① TO SHOW: $m_H(N)$ IS POLYNOMIAL

WE SHOW: $m_H(N) \leq \dots \leq \dots \leq$ a polynomial

KEY QUANTITY: $B(N, k)$: MAXIMUM NUMBER OF
DICHOTOMIES ON N POINTS WITH
BREAK POINT k

$B(N, k)$ IS PURELY COMBINATORIAL

RECURSIVE BOUND

ISOLATE
LAST POINT

CONSIDER THE TABLE

	# ROWS	x_1	x_2	...	x_{n-1}	x_n
S_1	α	+1	+1	...	+1	+1
		-1	+1	...	+1	-1
S_2^+	β	+1	-1	...	+1	+1
		-1	-1	...	+1	+1
S_2^-	β	+1	-1	...	+1	-1
		-1	-1	...	+1	-1

S_1 : ROWS THAT APPEAR
ONLY ONCE (FOR
 $x_1, \dots, x_{n-1}$)

$$B(N, k) = \alpha + 2\beta$$

ESTIMATE α AND β

IF WE FOCUS ON COLS $x_1 \rightarrow x_{n-1}$, α REMAINS
BUT WE HAVE β TWICE. $\alpha + \beta \leq B(N-1, k)$

NOW FOCUS ON $S_2 = S_2^+ \cup S_2^-$ ROWS

s_2^+	β	+1 -1 ... +1				+1
		-1 -1 ... +1				+1
		$\vdots$				
		+1 -1		+1		+1
s_2^-	β	-1 -1		-1		+1
		+1 -1		+1		-1
		-1 -1		+1		-1
		$\vdots$				
		-1 -1		+1		-1
		-1 -1		-1		-1

LOOK AT $s_N^+ \rightarrow$ MATRIX HAS BP OF $k-1$

IF YOU HAVE ALL POSSIBLE PATTERNS ON $k-1$ COLUMNS,
 ADDING x_n COL TAKES YOU TO k

$$\beta \leq B(N-1, k-1)$$

PUT IT TOGETHER AND DO IT AGAIN

$$B(N, k) = \alpha + \beta \quad \alpha + \beta \leq B(N-1, k)$$

$$\beta \leq B(N-1, k-1)$$

$$\therefore B(N, k) = B(N-1, k) + B(N-1, k-1)$$

SOLUTION IS UPPER BOUND FOR ANY
GROWTH FUNCTION WITH BREAK POINT

$$B(N, k) = B(N-1, k) + B(N-1, k-1)$$

		k					
		1	2	3	4	5	6 ...
N	1	1	2	2	2	2	2
	2	1	3	4	4	4	4
	3	1	4	7			
	4	1	5	11	11		
	5	1	6	16	22		
	6	1	7	24	38		
	...						

PUZZLE FROM LAST LECTURE → (4) in row 3, column 2

Diagram: A triangle with vertices at (3,4), (4,4), and (4,5) containing the text 2^N .

ANALYTIC SOLUTION:

THEOREM: $B(N, k) = \sum_{i=0}^{k-1} \binom{N}{i}$

BOUNDARY CONDITIONS: EASY

INDUCTION STEP: $N-1$ $\rightarrow$ N

$$\sum_{i=0}^{k-1} \binom{N}{i} = \sum_{i=0}^{k-1} \binom{N-1}{i} + \sum_{i=0}^{k-2} \binom{N-1}{i} ?$$

$$= 1 + \sum_{i=1}^{k-1} \binom{N-1}{i} + \sum_{i=1}^{k-1} \binom{N-1}{i-1}$$

$$= 1 + \sum_{i=1}^{k-1} \left[\binom{N-1}{i} + \binom{N-1}{i-1} \right]$$

REDUCE WITH COMBINATORIAL ARGUMENT

CHOOSE 10 PEOPLE FROM N IN A ROOM

N CHOOSE 10

OR CHOOSE 10 PEOPLE FROM N-1

N-1 CHOOSE 10

PLUS CHOOSE 9 PEOPLE FROM N-1

SO

$$1 + \sum_{i=1}^{k-1} \binom{N}{i} = \sum_{i=0}^{k-1} \binom{N}{i} \quad \swarrow \text{POLYNOMIAL!}$$

GIVEN

FOR $\mathbb{N}$, BREAK POINT IS FIXED

$$m_2(N) \leq \sum_{i=0}^{k-1} \binom{N}{i} \quad \text{max power is } N^{k-1}$$

k is constant;

this is polynomial in N

THREE EXAMPLES

$$\sum_{i=0}^{k-1} \binom{N}{i}$$

POSITIVE RAYS: $k=2$

$$m_H(N) = N+1 \leq \binom{N}{0} + \binom{N}{1} = 1+N$$

POSITIVE INTERVAL: $k=3$

$$m_H(N) = \frac{1}{2}N^2 + \frac{1}{2}N + 1 \leq \frac{1}{2}N^2 + \frac{1}{2}N + 1$$

2D PERCEPTION: $k=4$

$$m_H(N) = ? \leq \frac{1}{6}N^3 + \frac{5}{6}N^2 + 1$$

YOU SHOULD DO THESE REDUCTIONS

PROOF $m_H(N)$ CAN REPLACE M

INSTEAD OF

$$\mathbb{P}[|E_{in} - E_{out}| > \epsilon] \leq 2M e^{-2\epsilon^2 N}$$

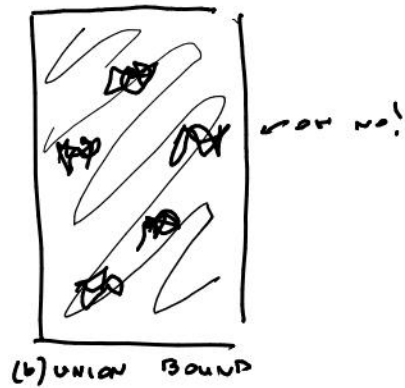
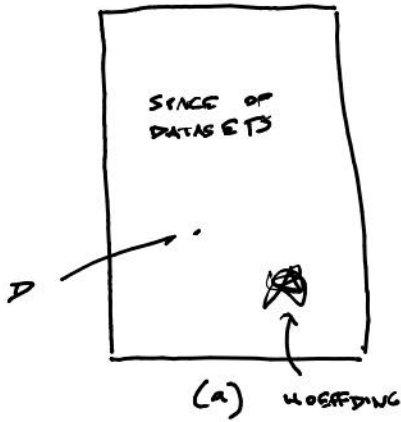
WE WANT

$$\dots \leq 2m_H(N) e^{-2\epsilon^2 N}$$

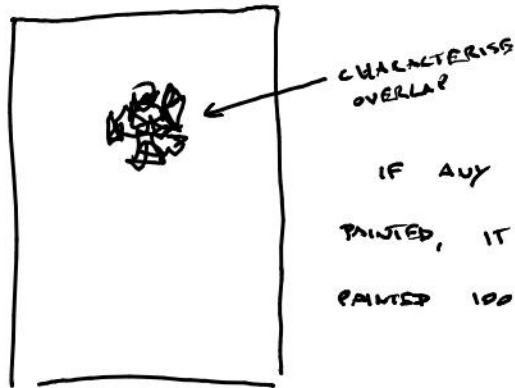
FORMAL PROOF IS A LOT!

- HOW DOES $m_H(N)$ RELATE TO OVERLAPS?
- WHAT TO DO ABOUT E_{out} ?
- PUT IT TOGETHER

PICTURES!



VC BOUND (c)



IF ANY POINT IS
PRINTED, IT WILL BE
PRINTED 100 TIMES. $A \leq \frac{1}{100}$

WE COLLAPSE HYPOTHESES INTO DICHOTOMIES, EXCEPT
Font is a thing.

$$E_{out}(h)$$



$$E_{in}(h) \dots \dots \dots$$

so $E_{in}(h)$ tracks $E_{in}'(h)$

$$E_{in}'(h) \dots \dots \dots$$

BECAUSE BOTH TRACK $E_{out}!$

$$E_{in}''(h) \dots \dots \dots$$

it's the same problem! but in E_{in} only so
it's back to dichotomies!

WE DON'T HAVE

$$P[|E_{in}(h) - E_{out}(h)| > \epsilon] \leq 2m_H(\epsilon) e^{-2\epsilon^2 N}$$

BUT RATHER:

..

$$\leq 4m_H(2N) e^{-1/8 \epsilon^2 N}$$

↑ BECAUSE OF THE $E_{in} \rightarrow E_{out}$

VAPNIK - CHERNOVENKIS
INEQUALITY

why $1/8$? ϵ ends up as $\epsilon/4$ in

proof

IT'S UGLY BUT SHOWS LEARNING IS FEASIBLE

PROBLEM SET #3

$$1. P[|E_{in}(g) - E_{out}(g)| > \epsilon] \geq 2Me^{-2\epsilon^2 N}$$

$$\text{set } \epsilon = 0.05, M = 1$$

$$\text{what is } N \text{ for } 0.03$$

$$-2(0.05^2)N$$

$$0.03 \geq 2e$$

$$-0.005N$$

$$0.015 \geq e$$

$$\ln(0.015) \geq -0.005N$$

$$-4.2 \geq -0.005N$$

$$\therefore N \leq \boxed{840}$$

b

✓

$$2. 0.03 \geq 20e^{-2(0.05)^2 N}$$

$$-0.005N$$

$$0.0015 \geq e$$

$$\therefore N \geq 1301$$

c

✓

$$3. 0.00015 \geq e^{-0.005N}$$

$$\therefore N \geq 1761$$

d

✓

4. Break point for a perceptron in 2D is

$$4. \text{ what about 3D? } \boxed{5}$$

b ✓

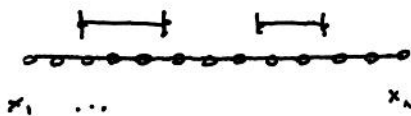
5. Which of the following are possible formulae for a growth function $m_H(N)$:

- i) $1+N$ ii) $2^{\lfloor N/2 \rfloor}$ ← not poly
 iii) $1+N+\binom{N}{2}$ iv) 2^N
 v) $\sum_{i=0}^{\lfloor \sqrt{N} \rfloor} \binom{N}{i}$ ← not poly

where $\lfloor u \rfloor$ is the biggest integer $\leq u$ and $\binom{M}{n} = 0$ when $n > M$

i, ii, v b ✓

6. $h: \mathbb{R} \rightarrow \{-1, +1\}$ and $h(x) = +1$ is point w/in intervals and -1 otherwise. smallest break point?



we know pos interval is

3; this gives us 2 more degrees of freedom. 5

c ✓

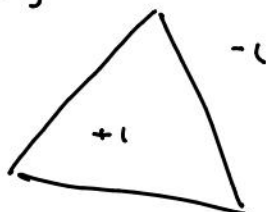
7. What is the growth function for the two-interval hypothesis set? place interval ends in 4 of $N+1$ spots

$$\frac{N+1}{4} + \frac{N+1}{2} \quad \checkmark$$

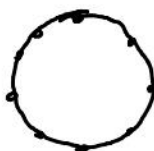
8. d ✓

9. triangle model:

$$\mathbb{R}^2 \rightarrow \{-1, +1\}$$



THE VERTICES SPLIT THE DATAPPOINTS ON THE CIRCLE
SELECTED FOR MAX DICHOIMIES:



so BREAK POINT IS $2 \cdot V + 1$

$\therefore 7$ ✓

10. COMPUTED THE GROWTH FUNCTION $m_H(N)$ for the
learning model made up of two concentric circles
in $\mathbb{R}^2$. SPECIFICALLY, Z_1 contains the functions which
are +1 for $a^2 \leq x_1^2 + x_2^2 \leq b^2$ and -1 otherwise

is this different from positive intervals? I don't
think so. $\therefore m_H(N) = \frac{N+1}{2} + 1$

this wouldn't work w/ non-zero origins or ellipses
but no info is encoded in angular dimension

LECTURE 7: THE VC DIMENSION

RECAP: VC INEQUALITY:

$$P[|E_n(g) - E_{out}(g)| > \epsilon] \leq 4 m_H(2N) e^{-\frac{1}{8} \epsilon^2 N}$$

where $m_H(N)$ is polynomial when k exists;

$$\text{i.e. } m_H(N) \leq \sum_{i=0}^{k-1} \binom{N}{i}$$

YOU WILL REMEMBER THE VC DIMENSION!

- DEFINITION
- VC DIMENSION OF PERCEPTRONS
- INTERPRETATION
- GENERALISATION BOUNDS ← KEY

VC DIMENSION IS A NUMBER DEFINED FOR H , DENOTED BY $d_{VC}(H)$; THE MOST POINTS H CAN SHATTER. THE LARGEST VALUE OF N FOR WHICH $m_H(N) = 2^N$.

is this just $k-1$ or am I missing something?

$N \leq d_{VC}(H) \Rightarrow H$ can shatter N points

$N > d_{VC}(H) \Rightarrow N$ is a break point of H

THE GROWTH FUNCTION

in terms of k

$$m_H(N) \leq \sum_{i=0}^{k-1} \binom{N}{i}$$

in terms of d_{VC}

$$m_H(N) \leq \sum_{i=0}^{d_{VC}} \binom{N}{i}$$

d_{VC} is order of polynomial which bounds growth function

EXAMPLE: POSITIVE RAYS: $d_{VC} = 1$

2D PERCEPTION: $d_{VC} = 3$

2D CONVEX SETS: $d_{VC} = \infty \leftarrow$ PESSIMISTIC UPPER BOUND

THE VC DIMENSION AND LEARNING

FINITE $d_{VC} \Rightarrow g \in \mathcal{H}$ will generalise
(FINAL HYPOTHESIS)

• INDEPENDENT OF THE LEARNING ALGORITHM

• " INPUT DISTRIBUTION

(BECAUSE WE ARE PICKING SAMPLES TO MAXIMISE DISCREPANCIES)

• " TARGET FUNCTION

• $d_{VC}, N \nmid g$ ARE ALL THAT'S LEFT
↓
DRIVES VC INEQUALITY/ PROBABALISTICALLY

d_{VC} FOR PERCEPTRONS

For $d=2$, $d_{VC}=3$

IN GENERAL, $d_{VC} = d+1 \leftarrow$ LET'S SHOW WHY

$$d_{VC} \leq d+1 \quad (1)$$

$$d_{VC} \geq d+1 \quad (2)$$

(1) A SET OF $N = d+1$ POINTS IN $\mathbb{R}^d$ SHATTERED BY THE PERCEPTRON

$$X = \begin{bmatrix} -x_1^T \\ -x_2^T \\ \vdots \\ -x_{d+1}^T \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 & \dots & 0 \\ 1 & 1 & 0 & \dots & 0 \\ 1 & 0 & 1 & \dots & 0 \\ \vdots & & & & \\ 1 & 0 & 0 & 0 & 1 \end{bmatrix}$$

CHOOSE SUCH THAT X IS INVERTIBLE

FOR ANY $y = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_{d+1} \end{bmatrix} = \begin{bmatrix} \pm 1 \\ \pm 1 \\ \vdots \\ \pm 1 \end{bmatrix}$, can we find a vector w satisfying $\text{sign}(xw) = y$

DROP THE SIGN! $xw = y$ and X is invertible

$$\therefore w = X^{-1}y$$

SO WE CAN SCATTER THESE POINTS.

THIS IMPLIES WHAT?

$$d_{vc} \geq d+1 \quad \textcircled{1} \quad \checkmark$$

② SHOW THAT $d_{vc} \leq d+1$ BY SHOWING WE CANNOT SCATTER ANY SET OF $d+2$ POINTS

FOR ANY $d+2$ POINTS,

$$x_1, \dots, x_{d+1}, x_{d+2}$$

MORE POINTS THAN DIMENSIONS $\therefore$

$$x_j = \sum_{i \neq j} a_i x_i \quad \leftarrow \text{AT LEAST ONE POINT IS LINEAR COMBO OF OTHERS}$$

WHERE NOT ALL THE a_i 'S ARE ZERO

(we know because first component of perceptron is fixed 1)

CONSIDER DICHOTOMY:

x_i 's WITH NON-ZERO a_i , get $y_i = \text{sign}(a_i)$

and x_j gets $y_j = -1$

NO PERCEPTRON CAN IMPLEMENT THIS

$$x_j = \sum_{i \neq j} a_i x_i \Rightarrow w^T x_j = \sum_{i \neq j} a_i w^T x_i$$

if $y_i = \text{sign}(w^T x_i) = \text{sign}(c_i)$, then $a_i w^T x_i > 0$
 this forces $\sum_{i \neq j} a_i w^T x_i > 0 \leftarrow \text{uh oh!!}$

$y_j = \text{sign}(w^T x_j) = +1 \leftarrow \therefore$ DICHOTOMY DOESN'T WORK,
 CAN'T SHATTER

$$\textcircled{1} \nleftrightarrow \textcircled{2} \Rightarrow d_{VC} = d+1$$

THIS IS THE NUMBER OF PARAMETERS IN
 PERCEPTRON $w_0, w_1, \dots, w_{d+1}$

INTERPRETATION OF VC DIMENSION

① DEGREES OF FREEDOM!

PARAMETERS CREATE DEGREES OF FREEDOM

of params: analog degrees of freedom

d_{VC} : equivalent 'binary' degrees of freedom

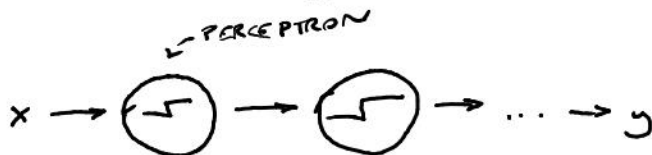
SO ABSTRACTS MECHANICS OF MODEL INTO

DICHOTOMIES

POSITIVE RAYS: ONE DEGREE OF FREEDOM ($d_{vc} 1$)

" INTERVALS: TWO DOF ($d_{vc} 2$)

NOT JUST PARAMETERS! PARAMETERS MAY NOT
CONTRIBUTE TO DEGREES OF FREEDOM



STACKING S PERCEPTRONS GIVES YOU $S(D+1)$ PARAMETERS
BUT DOESN'T ADD D.O.F

d_{vc} MEASURES EFFECTIVE NUMBER OF PARAMETERS

② NUMBER OF DATA POINTS NEEDED

VC INEQUALITY HAS TWO SMALL QUALITIES

$$P \left[|E_{in}(g) - E_{out}(g)| > \epsilon \right] \leq \underbrace{4m_{\mathcal{H}}(2N)^{-\frac{1}{2}\epsilon^2 N}}_{S \sim 2^{d_{vc}}}$$

IF WE WANT PARTICULAR ϵ & δ , HOW DOES
N DEPEND ON d_{vc} ?

← SIMPLE δ WHEN d_{vc}
EXISTS
 $N \sim \frac{d_{vc}}{\epsilon^2}$

Fix $N d^{-N} = \text{small value}$

How does N change with d ?

Bound is $\sim$ proportional to $d!$

PRACTICAL OBSERVATION

THE QUANTITY WE'RE TRYING TO BOUND FOLLOWS
THE BOUND $\leftarrow$ WELL....

RULE OF THUMB: $N \geq 10 d_{VC}$ GETS YOU INTO
INTERESTING REGION

GENERALISATION THEOREM

START FROM VC WEQUALITY

$$\mathbb{P} \left[|E_{in}(g) - E_{out}(g)| > \epsilon \right] \leq \underbrace{4 m_H(2N) e^{-\frac{1}{8} \epsilon^2 N}}_{\delta}$$

GET ϵ IN TERMS OF δ

$$\epsilon = 4 m_H(2N) e^{-\frac{1}{8} \epsilon^2 N} \Rightarrow \underbrace{\sqrt{\frac{8}{N} \ln \frac{4 m_H(2N)}{\delta}}}_{\Omega} = \epsilon$$

WITH PROBABILITY $\geq 1 - \delta$, $|E_{in} - E_{out}| \leq \Omega(N, H, \delta)$

GENERALISATION BOUND

$$E_{out} - E_{in} \leq \underbrace{\Omega(N, H, \epsilon)}_{\text{WE KNOW BEHAVIOUR: WHAT WE CARE ABOUT CONDENSE}} \leftarrow \text{WE PROP ABSOLUTE VALUE BECAUSE } E_{out} \geq E_{in} \text{ FOR}$$

WITH PROBABILITY $\geq 1 - \delta$

$$E_{out} \leq E_{in} + \Omega$$

WE DERIVE REGULARISATION FROM THIS!

(FOR ANOTHER TIME)

EXTRAS FROM LECTURE 7

SHATTERING - MATHEMATICAL DEFINITION

H SHATTERS A SET OF POINTS $\{x_1, \dots, x_n\}$

IF AND ONLY IF FOR EVERY POSSIBLE LABELING

$y \in \{-1, 1\}^n$ THERE EXISTS SOME $h \in H$ SUCH THAT

$h(x_i) = y_i$ FOR ALL i

SO FOR PERCEPTRON WITH γ (LABELS) FIND A

w SUCH THAT $\underbrace{\text{SIGN}(xw)}_{\text{CLASSIFIER}} = y$

$D_K \leq d+1$ FOR PERCEPTRON

$x_s = \sum a_i x_i$ (COMBO OF OTHER POINTS)

$$w^T x_s = \sum a_i (w^T x_i)$$

$$S^+ = \underbrace{i: a_i \geq 0}_{+1}, \quad S^- = \underbrace{i: a_i < 0}_{-1}$$

$$x_s = -1$$

DICHOTOMY: $w^T x > 0$ for $i \in S^+$

$w^T x < 0$ for $i \in S^-$

$$w^T x_s < 0$$

$$\text{SO: } w^T x_s = \sum_{+} a_i (w^T x_i) + \sum_{-} a_i (w^T x_i) + \sum_{0} a_i (w^T x_i)$$

$\therefore w^T x_s$ can't be less than zero, no shattering

LECTURE 8: BIAS-VARIANCE TRADEOFF

RECAP: WE HAVE GENERALISATION BOUND

$$E_{out} \leq E_{in} + \Omega(N, \epsilon, d_{VC})$$

N or d_{VC} is interesting - what drives this?

B-V IS A STANDALONE THEORY w/ DIFFERENT APPROACH THAN VC.

- BIAS $\hat{=}$ VARIANCE ①

- LEARNING CURVES ②

① BIAS AND VARIANCE

WE HAVE BEEN THINKING ABOUT AN APPROXIMATION AND GENERALISATION TRADEOFF -

SMALL E_{out} : good approx. of f out of sample
MORE COMPLEX $H \Rightarrow$ better chance to approx. f
LESS " $\Rightarrow$ better chance of generalising

VC ONE APPROACH: $E_{out} \leq E_{in} + \Omega$

B-V IS ANOTHER: decompose E_{out} into

- ① How well H can approximate f

- ② How well we can hope in H to get h

B-V APPLIES TO REAL-VALUE FUNCTIONS,
UNLIKE OUR WORK SO FAR WITH VC DIMENSION,
AND USES SQUARED ERROR (REQUIRED FOR DECOMP)

START WITH E_{out}

MAKE DEPENDENCE
ON DATASET EXPLICIT

DIFFERENCE BETWEEN HYPOTHESIS & TARGET

$$E_{out}(g^{(D)}) = E_x [(g^{(D)}(x) - f(x))^2]$$

$$E_D [E_{out}(g^{(D)})] = E_D [E_x [(g^{(D)}(x) - f(x))^2]]$$

we want
dependence to
go away!

$$= E_x [E_D [g^{(D)}(x) - f(x)]]$$

CAN SWAP INTEGRATIONS

LET'S FOCUS ON $E_D [(g^{(D)}(x) - f(x))^2]$. TO EVALUATE
WE MUST CONSIDER "AVERAGE HYPOTHESIS" $\bar{g}(x) = E_D [g^{(D)}(x)]$
WHAT g WOULD WE GET ON AVERAGE ON INFINITE
DATASETS?

IMAGINE MANY DATASETS $D_1, D_2, \dots, D_k$ BUT YOU
CAN ONLY LEARN ONE AT A TIME

$$\bar{g}(x) \approx \frac{1}{k} \sum_{i=1}^k g^{(D_i)}(x)$$

USING $\bar{g}(x)$

$$\begin{aligned} E_D [(g^{(D)}(x) - f(x))^2] &= E_D [(g^{(D)}(x) - \bar{g}(x) + \bar{g}(x) - f(x))^2] \\ &= E_D [(g^{(D)}(x) - \bar{g}(x))^2 + (\bar{g}(x) - f(x))^2 \\ &\quad + \underbrace{2(g^{(D)}(x) - \bar{g}(x))(\bar{g}(x) - f(x))}_{\text{CROSS!}}] \end{aligned}$$

TREAT $\bar{g}(x)$ & $f(x)$
AS CONSTANTS; E_D OF
CONSTANT IS ITSELF

BUT! WE CAN REMEMBER THAT
 $E_D g^{(D)}(x) = \bar{g}(x)$ AND WE KNOW
EXPECTED VALUE SUMS;
 $\sum (\bar{g}(x) - \bar{g}(x)) = 0$, BYE BYE
CROSS TERM!

VARIANCE
TERM

$$= E_D [(g^{(D)}(x) - \bar{g}(x))^2] + (\bar{g}(x) - f(x))^2$$

↑
BIAS TERM.

$$E_D [(g^{(D)}(x) - f(x))^2] = \text{var}(x) + \text{bias}(x)$$

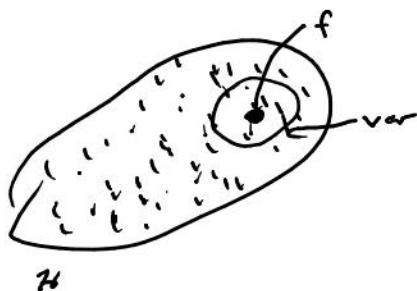
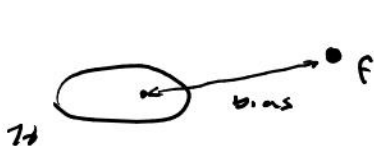
$$\therefore E_D [E_{\text{out}}(g^{(D)})] = E_x [\text{bias}(x) + \text{var}(x)]$$

$$= \text{bias} + \text{var}$$

THE TRADEOFF

$$\text{bias} = \mathbb{E}_x [(\bar{g}(x) - f(x))^2]$$

$$\text{variance} = \mathbb{E}_x [\mathbb{E}_D [g^{(D)}(x) - \bar{g}(x)]^2]$$



SMALL HYPOTHESIS $\Rightarrow$ TO BIGGER MEANS BIAS GOES DOWN $\uparrow$ VARIANCE GOES UP

EXAMPLE: SINE TARGET

$$f(x) = \sin(\pi x) \quad f: [-1, 1] \rightarrow \mathbb{R}$$

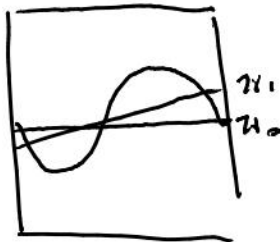
we get two training examples. $N=2$

" hypothesis sets $H_0: h(x) = b$

$$H_1: h(x) = ax + b$$

which is better?

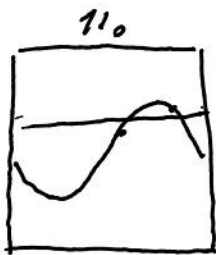
APPROXIMATION:



$$H_0, E_{out} = 0.5$$

$$H_1, E_{out} = 0.2$$

LEARNING:



ERROR DEPENDS ON DATASET!

BUS & VARIANCE - H_0



$\bar{g}(x)$ is the approximation in this case

big bias, small variance

H_1 : LINE CANNAGE!



$$H_0: \text{bias } 0.5$$

$$\begin{array}{r} \text{var } 0.25 \\ \hline 0.75 \end{array}$$

$$H_1: \text{bias } 0.21$$

$$\begin{array}{r} \text{var } 1.69 \\ \hline 1.90 \leftarrow \text{uh oh!} \end{array}$$

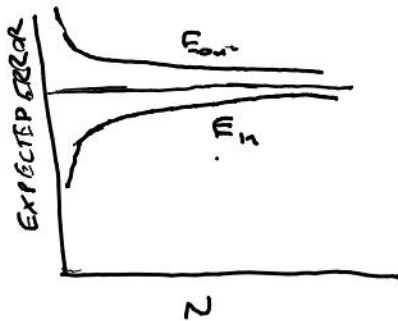
MATCH THE MODEL COMPLEXITY TO THE DATA RESOURCES, NOT THE TARGET COMPLEXITY

② LEARNING CURVES

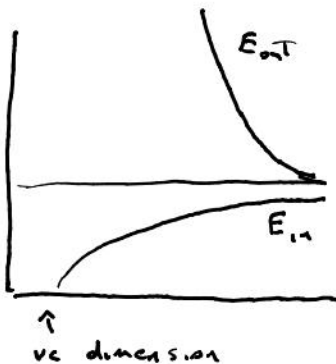
EXPECTED E_{out} AND E_{in}

DATASET D OF SIZE N

$$\begin{aligned} E_{out} &: \mathbb{E}_D [E_{out}(g^{(D)})] \\ E_{in} &: \mathbb{E}_D [E_{in}(g^{(D)})] \end{aligned} \quad \left. \vphantom{\begin{aligned} E_{out} &: \mathbb{E}_D [E_{out}(g^{(D)})] \\ E_{in} &: \mathbb{E}_D [E_{in}(g^{(D)})] \end{aligned}} \right\} \text{how do they vary w/ } N?$$

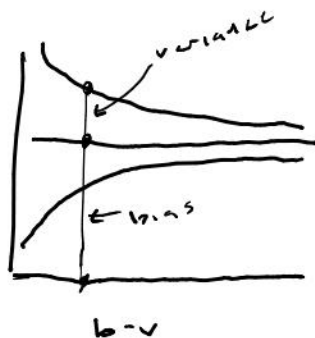
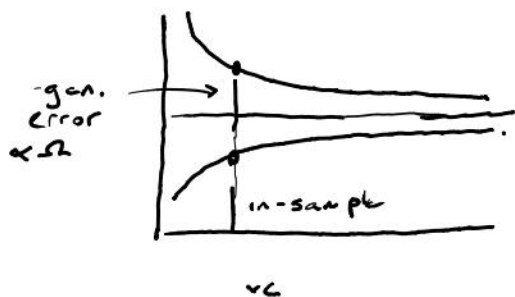


SIMPLE MODEL



complex model

VC vs bias-variance



LINEAR REGRESSION CASE

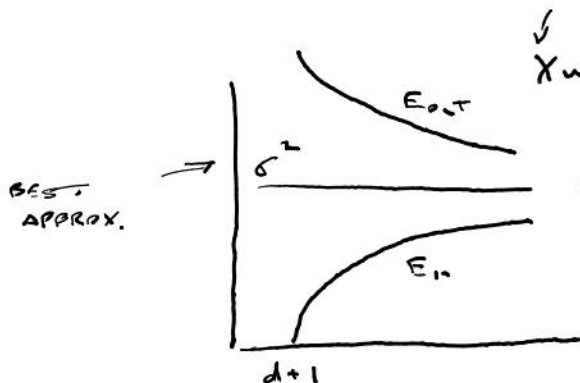
NOISY TARGET $y = w^*T x + \text{noise}$

$D = \{(x_1, y_1), \dots, (x_N, y_N)\}$

$w = X^\dagger y$

in-sample error vector: $Xw - y$

"out of sample" error vector: generate fresh dataset
w/ new noise y'



$Xw - y'$

exp. in-sample = $\sigma^2 \left(1 - \frac{d+1}{N}\right)$

out-sample = $\sigma^2 \left(1 + \frac{d+1}{N}\right)$

gen. error: $2\sigma^2 \left(\frac{d+1}{N}\right)$

VC dimension!

PROBLEM SET #4

$$1: d_{vc} = 10 \quad E < 0.05$$

$$0.05 \leq \sqrt{\frac{8}{N} \ln \frac{4m_H(ZN)}{0.95}}$$

$$0.0025N = 8 \ln \left(\frac{4m_H(ZN)}{0.95} \right) = 8 \ln \left(\frac{8N^{10}}{0.95} \right)$$

$$0.0003125N = \ln(7.6 \cdot N^{10})$$

$$= \ln(7.6) + \ln(N^{10})$$

$$= \ln(7.6) + 10 \cdot \ln(N)$$

$$0.0003125N = 2 + 10 \ln(N)$$

$$N = 6400 + 32,000 \ln(N)$$

$$N - 32,000 \ln(N) = 6400$$

$$420,000 \quad (b)$$

$$2: d$$

$$3: c$$

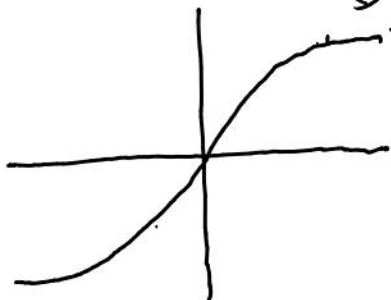
4: $f: [-1, 1] \rightarrow \mathbb{R}$ given by $f(x) = \sin(\pi x)$

getting something like $1.4x \rightarrow (c)$

5: -0.27 bias $\rightarrow (b)$

6: ~ 0.2 variance $\rightarrow (a)$

7:



ignore $ax+b$, b .
& ax^2

ax^2+b will have
a decent bias

but horrible variance

$\therefore (b)$

8: Assume $q \geq 1$ is an integer and let $m_N(1) = 2$

what is the hypothesis set whose growth
function satisfies: $m_H(N+1) = 2m_H(N) - \binom{N}{1}$?

Recall that $\binom{M}{n} = 0$ when $m > M$

growth function starts polynomial regime at

$N = q+1$; $d_{VC} = q$ (c)

9. $\mathcal{H}_{inter} = \mathcal{H}_1 \cap \mathcal{H}_2 \dots \cap \mathcal{H}_N$. if $A \subseteq B \Rightarrow d_{VC}(A) \leq d_{VC}(B)$

$\therefore \min \{d_{VC}(\mathcal{H}_k)\}_{k=1}^N$ is our best upper bound

$\therefore (b)$

for union
10. lower bound is $\max \{d_{H_k}\}_{k=1}^K$

K is a break point for set union;
not obvious contained within sum; $\therefore K-1$ is
an upper bound; $\therefore c$

LECTURE 9: THE LINEAR MODEL II

REMEMBER: $N \propto$ "VC DIMENSION"

WHERE WE ARE:

- LINEAR CLASSIFICATION ✓
- LINEAR REGRESSION ✓
- LOGISTIC REGRESSION
- NONLINEAR TRANSFORMS ~

CLEAR UP NET DIMENSIONALITY

$$x = (x_0, x_1, \dots, x_d) \xrightarrow{\Phi} z = (z_0, z_1, \dots, z_g)$$

Each $z_i = \phi_i(x)$ ← any function; arbitrary length

$$z = \Phi x$$

EXAMPLE: $z = (1, x_1, x_2, x_1 x_2, x_1^2, x_2^2)$

find hypothesis $g(x)$ in χ space:

$$\text{sign}(\tilde{w}^T \Phi(x)) \quad \text{or} \quad \tilde{w}^T \Phi(x)$$

GENERALISATION

$$\begin{array}{ccc} x = (x_0, x_1, \dots, x_d) & \xrightarrow{\Phi} & z = (z_0, z_1, \dots, z_d) \\ \downarrow & & \downarrow \\ w & & \tilde{w} \end{array}$$

$d_{vc} = d+1$ $\tilde{w}$ much bigger than w

$d_{vc} \leq \tilde{d}+1$

TWO NON-SEPARABLE CASES



① use linear model in χ
accept $E_{in} > 0$ (A)

or

insist $E_{in} = 0$; go to high-dimensional χ (B)

⚠ BAD IDEA FOR THIS DATA!

② $z = (1, x_1, x_2, x_1 x_2, x_1^2, x_2^2)$

6-dimensional; 2x examples required?

why not: $z = (1, x_1^2, x_2^2)$? $\leftarrow \hat{d} = 3$

or: $z = (1, x_1^2 + x_2^2)$? $\leftarrow \hat{d} = 2$

or: $z = (x_1^2 - x_2^2 - 0.6)$? $\leftarrow \hat{d} = 1$???

THIS DOESN'T WORK! YOU'RE CHERRY-PICKING
 HYPOTHESIS FROM H ! GENERALISATION STARTS WITH H ,
 NOT PRUNED DOWN; VC INEQUALITY STOPS APPLYING
 IF YOU SNOOP

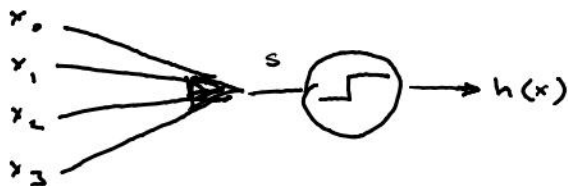
SO HOW DO WE CHOOSE A MODEL?? (FOR LATER)

LOGISTIC REGRESSION

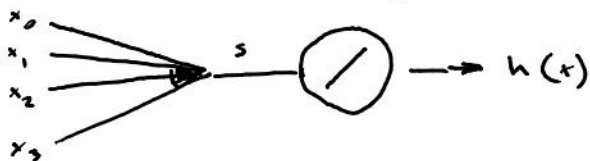
- THE MODEL ①
- ERROR MEASURE ②
- LEARNING ALGORITHM ③

$$\textcircled{1} \quad s = \sum_{i=0}^d w_i x_i$$

linear classification
 $h(x) = \text{sign}(s)$

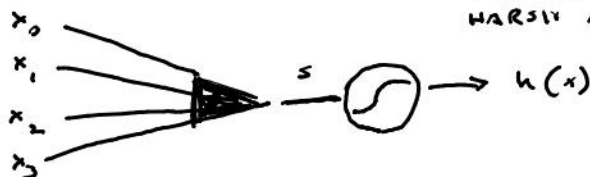


linear regression $h(x) = s$



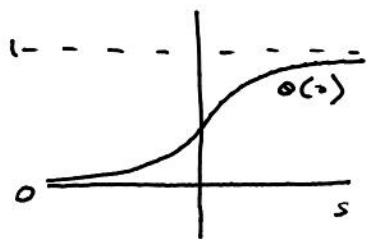
logistic regression $h(x) = \sigma(z)$

NON-LINEAR BUT NOT AS
HARSHLY AS $\text{SIGN}()$



THE MAIN UTILITY IS THAT THE OUTPUT IS
INTERPRETED AS A PROBABILITY

THE LOGISTIC FUNCTION σ



obv can be interpreted
as a probability/uncertainty

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

soft threshold
sigmoid

THIS FORMULA IS CHOSEN FOR
FRIENDLY FUTURE MATHS

PROBABILITY INTERPRETATION

$h(x) = \sigma(s)$ is interpreted as a probability

EXAMPLE: heart attacks

input: age, weight, etc

$\sigma(s)$: probability of h.a. over some time

$s = w^T x$ "risk score"

THIS IS GENUINE PROBABILITY

DATA (x, y) WITH BINARY y , GENERATED BY NOISY TARGET

$$P(y|x) = \begin{cases} f(x) & \text{for } y=1 \\ 1-f(x) & \text{for } y=-1 \end{cases}$$

WE'RE TRYING TO LEARN f , WHICH MUST BE PROB!

$$f: \mathbb{R}^d \rightarrow [0, 1]$$

LEARN $g(x) = \sigma(w^T x) \approx f(x)$

NOW DO WE CHOOSE THE WEIGHTS?

ERROR MEASURE

for each (x, y) , y generated by $f(x)$

plausible measure based on likelihood:

if $h=f$, how likely to get y from x ?

this is never completely clean - live with it!

$$P(y|x) = \begin{cases} h(x) & \text{for } y=1 \\ 1-h(x) & \text{for } y=-1 \end{cases}$$

substitute $h(x) = \sigma(w^T x)$, noting $\sigma(-s) = 1 - \sigma(s)$

use to simplify split: $P(x|y) = \sigma(y w^T x)$

likelihood of $D = (x_1, y_1), \dots, (x_N, y_N)$ is

$$\prod_{n=1}^N P(y_n | x_n) = \prod_{n=1}^N \sigma(y_n w^T x_n)$$

adjusting the weights to compromise between data points captures probability.

$$\text{minimise } -\frac{1}{N} \ln \left(\prod_{n=1}^N \sigma(y_n w^T x_n) \right)$$

$$= \frac{1}{N} \sum_{n=1}^N \ln \left(\frac{1}{\sigma(y_n w^T x_n)} \right) \quad \left[\sigma(s) = \frac{1}{1+e^{-s}} \right]$$

$$E_{in}(w) = \frac{1}{N} \sum_{n=1}^N \underbrace{\ln(1 + e^{-y_n w^T x_n})}_{\substack{e(h(x_n), y_n) \\ \downarrow \\ \text{error measure}}} \quad \text{if sign agrees with risk factor, small error}$$

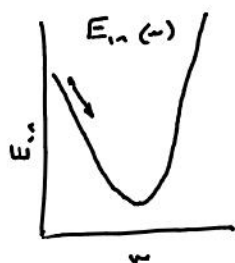
"cross-entropy" error

LOGISTIC REGRESSION ALGORITHM

$$E_{in}(w) = \frac{1}{N} \sum_{n=1}^N \ln(1 + e^{-y_n w^T x_n})$$

for linear regression we had a closed-form solution.

we don't have that here. let's use gradient descent



(generally applicable as long as surface is smooth)

start at $w(0)$; take step along steepest slope

you only use local information

fixed step size: small to avoid higher-order noise

$$w(1) = w(0) + \eta \hat{v}$$

what is the direction $\hat{v}$?

$$\Delta E_{in} = E_{in}(w(0) + \eta \hat{v}) - E_{in}(w(0))$$

$$= \eta \nabla E_{in}(w(0))^T \hat{v} \leftarrow \text{TAYLOR APPROX, IGNORE HIGH-ORDER}$$

$$\Delta = \eta \|\nabla E_{in}(w(0))\| \quad \hat{v} \text{ is a unit vector}$$

$$\hat{v} = - \frac{\nabla E_{in}(w(0))}{\|\nabla E_{in}(w(0))\|}$$

STEP SIZE: TRADEOFF BETWEEN ACCURACY & SPEED

η too small: takes forever

η too large: may not converge, 2nd & 3rd order terms dominate

variable η - just right

η increasing with slope is best

instead of $\Delta w = \eta \hat{J}$

$$\text{use } \Delta w = -\eta \nabla E_n(w(0))$$

THIS IS A FIXED LEARNING RATE, NOT STEP SIZE

LOGISTIC REGRESSION ALGORITHM

1. INITIALISE WEIGHTS AT $t=0$ to $w(0)$
2. for $t=0, 1, 2, \dots$ do
3. compute gradient $\nabla E_n = -\frac{1}{N} \sum_{n=1}^N \frac{\partial w \cdot x_n}{1 + e^{-w^T x_n}}$
4. update weights $w(t+1) = w(t) - \eta \nabla E_n$
5. iterate
6. return w

LECTURE 10: NEURAL NETWORKS

REVIEW: GRADIENT DESCENT

Q: IMPLEMENT GRADIENT?

TERMINATION CRITERIA

NEURAL NETWORKS

- STOCHASTIC GRADIENT DESCENT
- NEURAL NETWORK MODEL
- BACKPROPAGATION

STOCHASTIC GRADIENT DESCENT

GD MINIMISES $E_{in}(w) = \frac{1}{N} \sum_{n=1}^N \underbrace{c(h(x_n), y_n)}_{1/(1+e^{-y_n \cdot w})} \leftarrow i.e.$

by iterative steps along $-\nabla E_{in}$

$$\Delta w = -\eta \nabla E_{in}(-)$$

∇E_{in} based on all examples

"batch" gr.

STOCHASTIC: PICK ONE AT A TIME.

APPLY GD TO $c(h(x_n), y_n) \leftarrow$ LIKE PLA

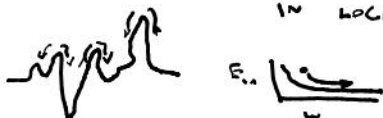
"AVERAGE DIRECTION" $E_n[-\nabla c(h(x_n), y_n)]$

PLUS NOISE $\rightarrow \quad = \frac{1}{N} \sum_{n=1}^N -\nabla c(h(x_n), y_n)$
 $= -\nabla E_{in}$

HOPE NOISE CANCELS

MAIN BENEFIT: CHEAPER COMPUTATION

OTHER BENEFITS: RANDOMISATION — DON'T WANT TO BE TOO DETERMINISTIC, TRAPS IN LOCAL MINIMA / PLATEAUX



SIMPLE! RULE OF THUMB —

$\eta = 0.1$ works

SGD IN ACTION

MOVIE RATINGS — user and movies have own profiles, reverse engineer ratings

user i , movie j , rating: r_{ij}
 $\downarrow$ $\downarrow$
 $u_{i1}, u_{i2} \dots u_{ik}$ $v_{j1}, v_{j2} \dots v_{jk}$

$$\left(r_{ij} - \sum_{k=1}^K u_{ik} v_{jk} \right)^2 = e_{ij} \leftarrow \text{minimising this}$$

going rating by rating is stochastic gradient descent!

step η down gradient

NEURAL NETWORK MODEL

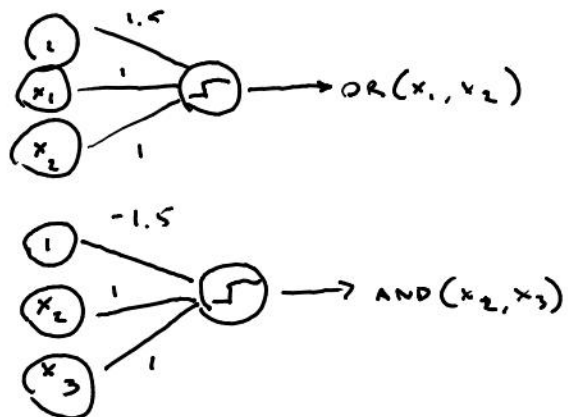
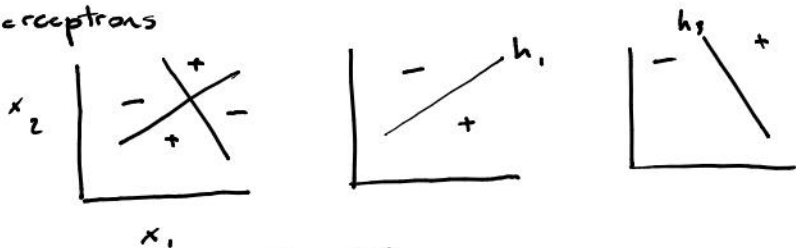
BIOLOGICAL INSPIRATION

function $\rightarrow$ structure NEURONS AS PERCEPTIONS

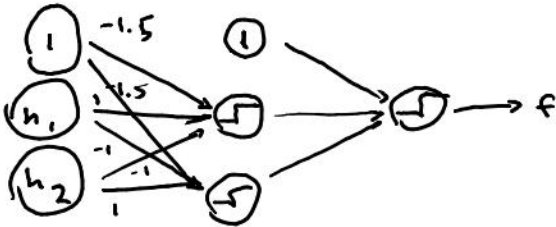
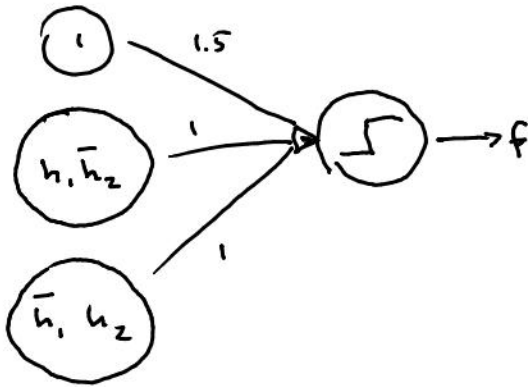
but then pick what is relevant and
take it :D don't get carried away
with biology

what can we do with combos of

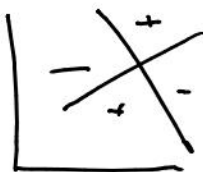
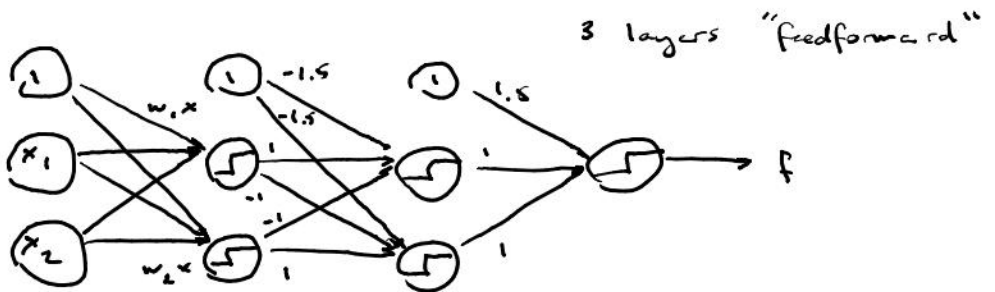
perceptrons



CREATING LAYERS

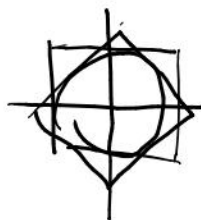
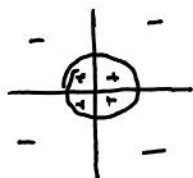


MULTILAYER PERCEPTRON



solved!

POWERFUL!



8 perceptrons

2 red flags — obv generalisation

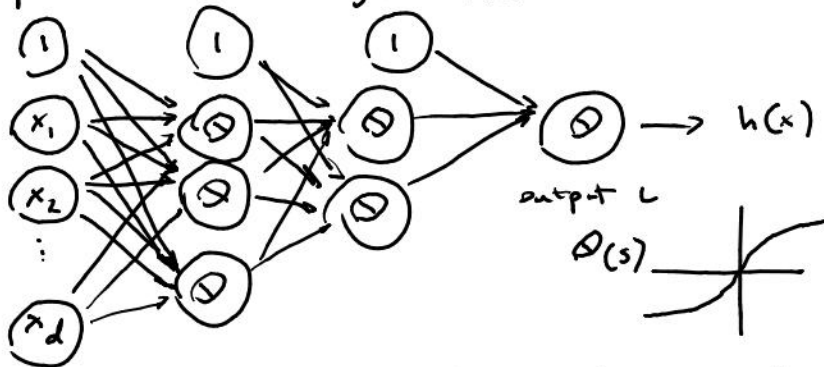
also optimisation

↑ difficult to solve ①

① soften threshold to make error surface smooth

input =

hidden layers $1 \leq l \leq L$



in practice not all nodes use
same σ

How π OPERATES

$$\sigma(z) = \tanh(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}} \quad \begin{array}{l} \text{(looks linear near 0,} \\ \text{hard near } 1/-1) \end{array}$$

$$w_{ij}^{(l)} \begin{cases} 1 \leq l \leq L & \text{layers} \\ 0 \leq i \leq d^{(l-1)} & \text{inputs} \\ 1 \leq j \leq d^{(l)} & \text{outputs} \end{cases}$$

↑ DIMENSION OF LAYER

$$x_j^{(l)} = \sigma(z_j^{(l)}) = \sigma \left(\sum_{i=0}^{d^{(l-1)}} w_{ij}^{(l)} x_i^{(l-1)} \right)$$

APPLY x to $x_1^{(0)} \dots x_{d^{(0)}}^{(0)} \rightarrow \rightarrow \rightarrow x_1^{(L)} = h(x)$

HOW DO WE FIND THE WEIGHTS?

BACKPROPAGATION! APPLY SGD

All weights $w = \{w_{ij}^{(l)}\}$ determine $h(x)$

Error on example (x_n, y_n) is $c(h(x_n), y_n) = c(w)$

need $\nabla c(w)$: $\frac{\partial c(w)}{\partial w_{ij}^{(l)}}$ for all i, j, l

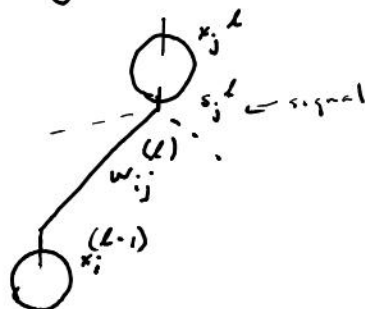
we need to do this efficiently - FFT mentioned

Computing $\partial c(w) / \partial w_{ij}^{(l)}$

we can evaluate one by one: analytically or numerically
but would prefer all at once!

trick:
$$\frac{\partial c(w)}{\partial w_{ij}^{(l)}} = \frac{\partial c(w)}{\partial s_j^{(l)}} \times \frac{\partial s_j^{(l)}}{\partial w_{ij}^{(l)}}$$

chain rule for partials work
because weights are isolated to
layer signal



we have
$$\frac{\partial s_j^{(l)}}{\partial w_{ij}^{(l)}} = x_i^{(l-1)}$$

← TRICKIER

we only need
$$\frac{\partial c(w)}{\partial s_j^{(l)}} = \delta_j^{(l)}$$

δ for last layer $\leftarrow$ start here and back-propagate

$$\delta_j^{(l)} = \frac{\partial c(w)}{\partial s_j^{(l)}} \Rightarrow \delta_1^{(4)} = \frac{\partial c(w)}{\partial s_1^{(4)}} ; c(w) = c(h(x_n), y_n)$$

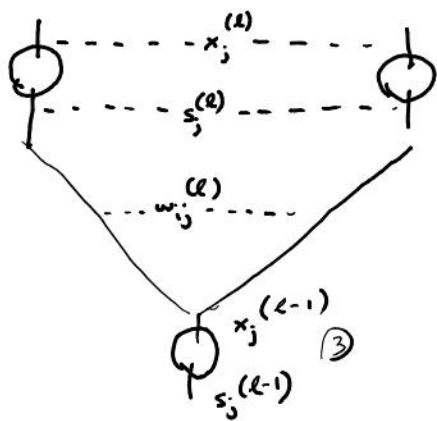
$$h(x_n) = x_1^{(4)}$$

$$e.g. c(w) = (x_1^{(4)} - y_n)^2 \quad \leftarrow \text{for squared error}$$

$$x_1^{(4)} = \sigma(s_1^{(4)})$$

for tanh $\rightarrow \sigma'(s) = 1 - \sigma^2(s)$

BACK PROPAGATE δ



$$\begin{aligned}\delta_i^{(L-1)} &= \frac{\partial c(\sim)}{\partial s_i^{(L-1)}} \quad (2) \\ &= \sum_{j=1}^{d^{(L)}} \frac{\partial c(\sim)}{\partial s_j^{(L)}} \times \frac{\partial s_j^{(L)}}{\partial x_i^{(L-1)}} \times \frac{\partial x_i^{(L-1)}}{\partial s_i^{(L-1)}} \quad (3) \\ &= \sum_{j=1}^{d^{(L)}} (1) \times (2) \times (3)\end{aligned}$$

$$(3) = \sigma'(s_i^{(L-1)})$$

$$(2) = w_{ij}^{(L)}$$

$$(1) = \delta_j^{(L)}$$

$$\delta_i^{(L-1)} = (1 - (x_i^{(L-1)})^2) \sum_{j=1}^{d^{(L)}} w_{ij}^{(L)} \delta_j^{(L)}$$

ALGORITHM

- 1: INITIALISE ALL WEIGHTS $w_{ij}^{(L)}$
- 2: for $t=0, 1, 2, \dots$ do
- 3: Pick $n \in \{1, 2, \dots, N\}$
- 4: FORWARD: COMPUTE ALL $x_j^{(L)}$
- 5: BACKWARD: "
- 6: UPDATE WEIGHTS: $w_{ij}^{(L)} \leftarrow w_{ij}^{(L)} - \eta x_i^{(L-1)} \delta_j^{(L)}$
- 7: ITERATE UNTIL TERMINATION
- 8: RETURN

① DON'T INITIALISE TO ZERO; DO IT RANDOMLY
 ↑
 MATH BREAKS AND CAN'T ADJUST, RANDOM BREAKS SYMMETRY

FINAL REMARK: HIDDEN LAYERS

nonlinear transform (learned, not imposed)

$d_{vc} \sim \# \text{ weights}$

interpretation? hard $\Rightarrow$ explanation is not what we're optimising

HOMEWORK #5

$$1. \quad \overbrace{\mathbb{E}_D [E_{in}(w_{lin})]}^{0.08} = \sigma^2 \left(1 - \frac{d+1}{N}\right)$$

$$\sigma = 0.1$$

$$d = 8$$

$$\hookrightarrow 0.088 = 0.01 \left(1 - \frac{9}{N}\right)$$

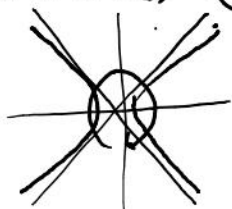
$$0.8 = 1 - \frac{9}{N}$$

$$\frac{9}{N} = 0.2$$

$$N = 45$$

$\therefore c \quad \checkmark$

$$2. \quad \Phi(1, x_1, x_2) = (1, x_1^2, x_2^2)$$



weights have to be
different signs, operating in
2-space

experiment: $x_1 = 0, x_2 = 1 \Rightarrow +1$

$$0 + 1 = 1$$

$$x_1 = -1.5, x_2 = 0 \rightarrow -1$$

$$(-1)225 \rightarrow -1 \quad \checkmark$$

$\therefore d \quad \checkmark$

3. dv_c of $\Phi_4: x \rightarrow (x_1, x_2, \dots, x_n) = 15$

$\therefore c \checkmark$

$$4. E(u, v) = (\underbrace{ue^v - 2ve^{-u}}_{\alpha})^2$$

$$= 2\alpha$$

$$\alpha = (e^v + 2ve^{-u})(ue^v - 2ve^{-u})$$

$c \checkmark$

5. WRITE A PROGRAM IMPLEMENTING GRADIENT DESCENT ON THIS ERROR SURFACE :)

LEARNING RATE = 0.1; starting point is (1, 1)

$$\partial u / \partial x = 2(e^v + 2ve^{-u})(ue^v - 2ve^{-u})$$

$$\partial v / \partial x = 2(ue^v - 2e^{-u})(ue^v - 2ve^{-u})$$

COUNT FOR (e^{-u}) THRESHOLD = 10 $d \checkmark$

$$6. [u, v] = [0.045, 0.023] = c \checkmark$$

7. error on coordinate descent = 0.13 $\rightarrow a \checkmark$

8. WRITE A PROGRAM IMPLEMENTING LOGISTIC REGRESSION - NB THIS IS A WEIRD PERCEPTRON!

REPEAT 100x $w / w^{(0)} = [1, 0, 0]$

TERMINATE WHEN $\|w^{(t+1)} - w^{(t)}\| < 0.01$

CROSS-ENTROPY ERROR =

8. $E_{\text{out}} = 0.104 \rightarrow d \checkmark$

9. Epochs = 330 $\rightarrow a \checkmark$

10. PLA as SGD

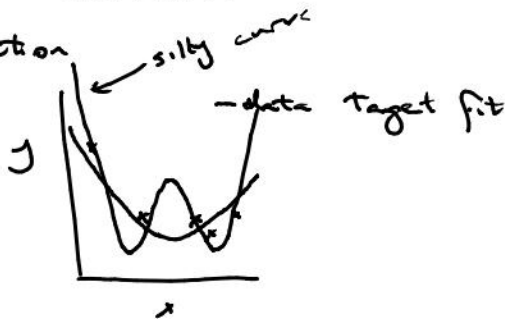
PLA error is $\text{sign}(y w^T x)$
not a nice gradient $\rightarrow c \checkmark$

LECTURE 11: OVERFITTING

- ① • what is overfitting? applies to every machine learning problem; tools to deal with it are universal
- ② • the role of noise
- ③ • deterministic noise
- ④ • dealing with overfitting

① AN ILLUSTRATION OF OVERFITTING

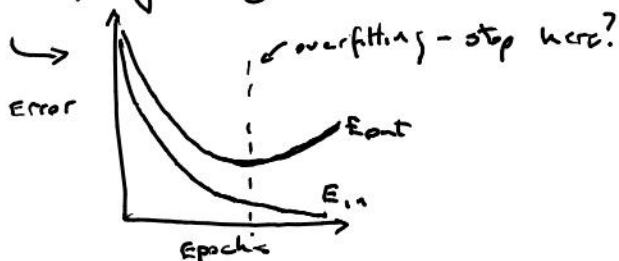
simple target function
five data points
small noise
let's fit w/ 4th order
oh no!!



$E_{\text{in}} = 0$, E_{out} is huge

OVERFITTING WITHIN SAME MODEL

neural net fitting noisy data



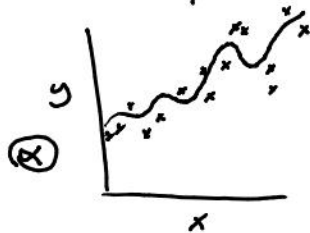
as you explore more of the weight space
generalisation error gets worse

overfitting occurs when $E_{in} \downarrow$ corresponds to $E_{out} \uparrow$

OVERFITTING: FITTING THE DATA MORE THAN IS
WARRANTED

CULPRIT: FITTING THE NOISE - HARMFUL; EXTRAPOLATES
NOISE OUT OF SAMPLE
"HALUCINATION"

CASE STUDY



TENTH ORDER TARGET
+ NOISE

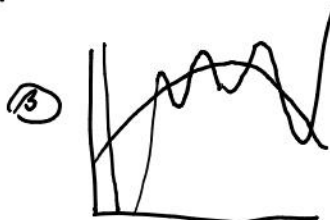
↓ vs
50TH ORDER TARGET
w/o NOISE



TWO FITS FOR EACH TARGET



	2 nd	10 th
E_{in}	0.05	0.034
E_{out}	0.127	9.00



	2 nd	10 th
E_{in}	0.21	~0
E_{out}	0.12	7680 ← lol

we got overfitting w/o noise!

TWO LEARNERS LEARNING A 10th - ORDER TARGET

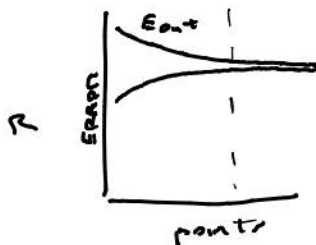
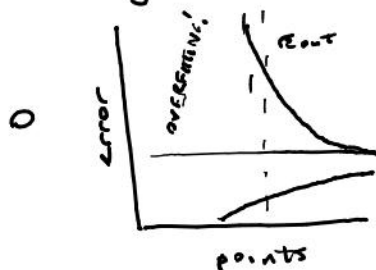
O - 10th ORDER MODEL

R - 2nd ORDER MODEL

"HOW MANY POINTS DO WE HAVE?"

THE DATA MATTERS MORE THAN THE TARGET FUNCTION. MATCH THE RESOURCES!

learning curves



0 loses w/ inadequate data resources

what about our noise-free overfitting?

$H_{10} \hat{=} H_2$ learning 50^{th} order target, no noise!

but is there really no noise?

A DETAILED EXPERIMENT

impact of noise level $\hat{=}$ target complexity

$$y = \underbrace{f(x)}_{\text{target}} + \underbrace{\epsilon(x)}_{\text{noise (independent)}} - \epsilon^2$$
$$y = \sum_{q=0}^{Q_f} \alpha_q x^q + \epsilon(x) \quad \leftarrow \text{letting us pick } f(x) \text{ of arbitrary power}$$

normalised

α not random - use Legendre polynomials so
we can access more of the target space

noise level: ϵ^2

for 10^{th} order

target complexity: Q_f

$Q_f = 10$

data set size: N

$N = 15$

$\epsilon = ?$

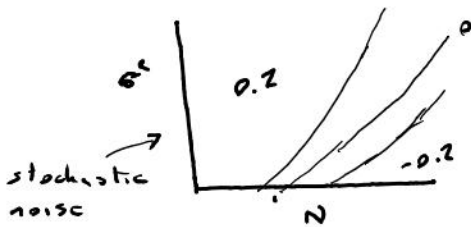
THE OVERFIT MEASURE

we fit the dataset $(x_1, y_1), \dots, (x_n, y_n)$ using
our two models: H_1 & H_2 .

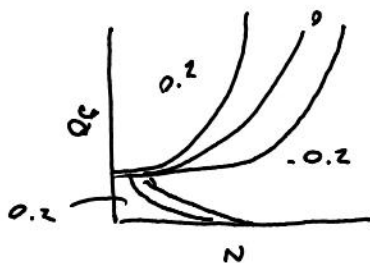
compare E_{out} of $g_2 \in H_2$ & $g_{1,0} \in H_1$

$$\text{overfit measure} = E_{out}(g_{1,0}) - E_{out}(g_2)$$

Let's look at $10c^2$ iterations



positive $\rightarrow$ high overfitting
as we increase noise, overfitting
gets worse



Q_f seems entangled with N
deterministic noise ③

what can we derive from this?

overfitting not determined by conventional noise alone

number of points $\uparrow$ overfitting $\downarrow$

stochastic noise $\uparrow$ " $\uparrow$

deterministic " $\uparrow$ " $\uparrow$

DEFINITION OF DETERMINISTIC NOISE

THE PART OF f THAT H CANNOT CAPTURE
THE NOISE IN THE LEARNING VS THE NOISE
IN THE DATA?

DETERMINISTIC NOISE 1. DEPENDS ON H

2. fixed for given x

neither is true for stochastic noise

IMPACT ON OVERFITTING

deterministic noise $\nexists$ $Q_f \leftarrow$ kicks in when $Q_c < \dim H$
for a finite sample you will model noise

noise $\nexists$ bias-variance

$$\text{decomposition } \mathbb{E}_D [(g^{(D)}(x) - f(x))^2] = \text{var}(x) + \text{bias}(x)$$

if f is noisy?

$$y = f(x) + \epsilon(x) ; \quad \mathbb{E}[\epsilon(x)] = 0$$

$$\mathbb{E}_{D, \epsilon} [(g^{(D)}(x) - y)^2] = \mathbb{E}_{D, \epsilon} [(g^{(D)}(x) - f(x) - \epsilon(x))^2]$$

$\Downarrow$

$$\text{var}(x) + \text{bias}(x) = \mathbb{E}_{\epsilon, x} [(\epsilon(x))^2]$$

DETERMINISTIC

THERE IS NOISE IN THE BIAS :)

AND THE BIAS IS INDEPENDENT OF N !

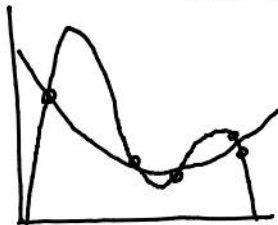
bias & σ^2 push variance around

④ DEALING WITH OVERFITTING

TWO CURVES: REGULARISATION - hit the brakes

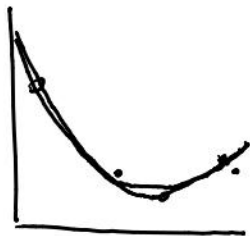
VALIDATION - check the bottom line

①



free fit

!! regularisation restraint/friction; 1%



← much better

LECTURE 12: REGULARISATION

REMEMBER: d_{VC} ANALYSIS ALLOWS FOR OVERFITTING;
IT DOESN'T PREDICT IT. THE VC BOUND IS FOR ALL
POSSIBLE TARGETS

DETERMINISTIC NOISE IS A FUNCTION OF THE
LIMITATIONS OF YOUR MODEL - CANNOT LEARN OUTSIDE
THE CAPACITY OF YOUR HYPOTHESIS SET. BIAS

REGULARISATION - USE IN EVERY ML APPLICATION

- INFORMAL ①
- FORMAL ②
- WEIGHT DECAY ③
- CHOOSING REGULARISER ④

TWO APPROACHES -

MATHEMATICAL

ILL-POSED PROBLEMS IN FUNCTION
APPROXIMATION

IF BAYES USED, THIS DOES IN
THE PRIOR

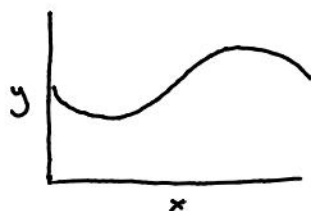
BUT IN MOST CASES THE
ASSUMPTIONS WERE DON'T HOLD!

BEST UTILITY IS TO DEVELOP
MATHEMATICAL INTUITION FOR
REAL-WORLD CASES

HEURISTIC

HANDCAPPING WEIGHTS,
FRICTION, PURELY INTUITIVE

EXAMPLE



instead of crazy lines like before, we restrict weights.
lose perfect E_u but gain E_{est}

w/o REG

bias: 0.21

variance: 1.69

w/REG

bias: 0.23

variance: 0.33

REGULARISATION REDUCES VARIANCE BUT MAY BUMP
BIAS - WE'RE KILLING SOME SIGNAL AS WELL AS NOISE.
BUT THE GAIN IN VARIANCE IS WORTH THE TRADE

THE POLYNOMIAL MODEL

WE USE LEGENDRE POLYNOMIALS $L_1, \dots, L_Q$
MUTUALLY ORTHOGONAL ^(*)

$\mathcal{H}_Q$: polynomials of order Q

$\alpha \quad L_1 = x$

$$z = \begin{bmatrix} 1 \\ L_1(x) \\ \vdots \\ L_Q(x) \end{bmatrix} \quad \mathcal{H}_Q = \left\{ \sum_{q=0}^Q w_q L_q(x) \right\}$$

$$L_2 = \frac{1}{2}(3x^2 - 1)$$

$$L_3 = \frac{1}{2}(5x^3 - 3x)$$

$$L_4 = \frac{1}{8}(35x^4 - 30x^2 + 3)$$

$$L_5 = \frac{1}{8}(63x^5 - \dots)$$

apply linear regression in z -space

REMEMBER - TARGET FUNCTION UNKNOWN!

UNCONSTRAINED SOLUTION

$$\text{given } (x_1, y_1), \dots, (x_n, y_n) \rightarrow (z_1, y_1), \dots, (z_n, y_n)$$

$$\begin{aligned} \text{MINIMISE } E_{lin}(w) &= \frac{1}{N} \sum_{n=1}^N (w^T z_n - y_n)^2 \\ &= \frac{1}{N} (Zw - Y)^T (Zw - Y) \end{aligned} \quad Z = \begin{bmatrix} -z_1, - \\ \vdots \\ -z_n, - \end{bmatrix}; \quad Y = \begin{bmatrix} -y_1, - \\ \vdots \\ -y_n, - \end{bmatrix}$$

$$w_{lin} = (Z^T Z)^{-1} Z^T Y$$

CONSTRAIN THE WEIGHTS

hard constraint: w_2 is constrained version of w_{lin}
($w_1 = 0$ for $q > 2$)

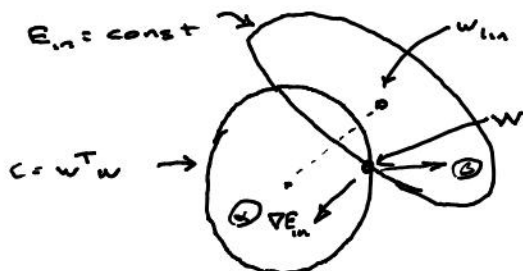
softer version: $\sum_{q=0}^Q w_q^2 \leq C$ "soft-order" constraint; Q is
"budget" subset of $\mathbb{Z}_q$

this constraint pushes weights to be small; it doesn't force any to zero

$$\begin{aligned} \text{MINIMISE: } & \frac{1}{N} (Zw - Y)^T (Zw - Y) = w_{reg} \quad (\text{not } w_{lin}) \\ \text{subject to } & w^T w \leq C \end{aligned}$$

SOLVING FOR w_{reg}

MINIMISE $E_{in}(w) = \frac{1}{N} (Zw - Y)^T (Zw - Y)$; $w^T w \leq C$



if C large enough to encompass w_{lin} ; answer is w_{lin}

① ∇ orthogonal to E_{in} surface

② normal to $w^T w = C$; w_{reg}

i reckon this is $\nabla E_{in} - w = 0$ ✓

$$\nabla E_{in}(w_{reg}) \propto -w_{reg}$$

$$= -2 \frac{\lambda}{N} w_{reg}$$

$$\nabla E_{in}(w_{reg}) + 2 \frac{\lambda}{N} w_{reg}$$

← reformatted for convenience
(checky ins)

MINIMISE $E_{in}(w) + \frac{\lambda}{N} w^T w$
 $\uparrow$
 regularisation
 term

λ is dependent of C , dataset,
 other stuff - but we know one
 exists

$$C \uparrow \lambda \downarrow$$

AUGMENTED ERROR

$$E_{\text{aug}}(w) = \frac{1}{N} (Zw - Y)^T (Zw - Y) + \frac{\lambda}{N} w^T w \quad \text{unconditionally}$$

— solves —

$$E_{\text{in}}(w) = \frac{1}{N} (Zw - Y)^T (Zw - Y) \quad \text{subject to } C \geq w^T w$$

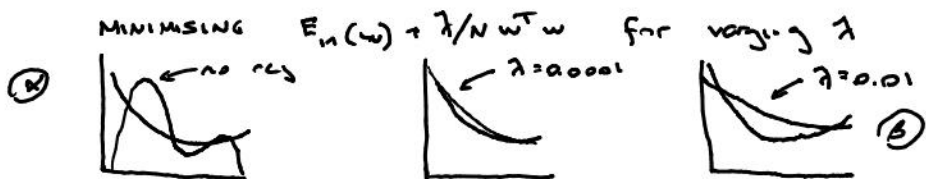
SOLUTION!

$$\nabla E_{\text{aug}}(w) = 0 \Rightarrow Z^T (Zw - Y) + \lambda w = 0$$

$$w_{\text{reg}} = (Z^T Z + \lambda I)^{-1} Z^T Y$$

$$\text{vs. } w_{\text{lin}} = (Z^T Z)^{-1} Z^T Y$$

RESULT



too much λ cares more about driving weights down than fitting data, choice of λ is critical!

(A) OVERFITTING

(B) UNDERFITTING

WEIGHT 'DECAY'

minimising $E_n(w) + \frac{\lambda}{2} w^T w$ is called weight decay. why?

GRADIENT DESCENT: $w(t+1) = w(t) - \eta \nabla E_n(w(t)) - 2\eta \frac{\lambda}{2} w(t)$
 $= w(t) \left(1 - 2\eta \frac{\lambda}{2} \right) - \eta \nabla E_n(w(t))$
 deficits step towards origin

VARIATIONS OF WEIGHT DECAY

EMPHASIS ON CERTAIN WEIGHTS: $\sum_{q=0}^Q \gamma_q w_q^2$
 $\gamma_q = 2^q \Rightarrow$ low order f.i.r
 $\gamma_q = 2^{-q} \Rightarrow$ high order f.i.r

neural nets get different γ in different layers

TIKHONOV REGULARISER: $w^T \Gamma^T \Gamma w$ (proper choice of Γ gives you weight decay etc.)

EVEN WEIGHT GROWTH!



(b) WEIGHT DECAY

(a) WEIGHT GROWTH

MAYBE DON'T CHOOSE WEIGHT GROWTH

PRACTICAL RULE: stochastic noise is high-frequency
 deterministic noise is non-smooth

THEREFORE $\Rightarrow$ constrain learning towards smoother hypothesis

GENERAL FORM OF AUGMENTED ERROR

regulariser $\Omega = \Omega(h)$, minimise $E_{\text{aug}}(h) = E_{\text{in}}(h) + \frac{\lambda}{N} \Omega(h)$

• similar to $E_{\text{out}}(h) \approx E_{\text{in}}(h) + \Omega(h)$ (*)

we're using E_{aug} to estimate E_{out} better than E_{in} can

⊗ GENERALISATION

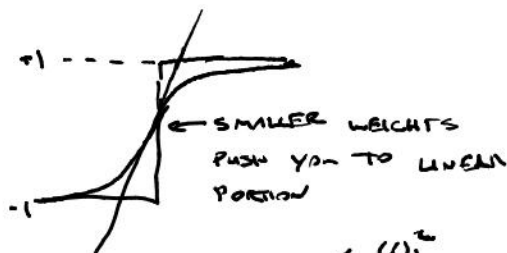
CHOOSING A REGULARISER

- 'PERFECT' Ω :
- CONSTRAINS IN DIRECTION OF $f(x)$
 - HARM NOISE MORE THAN SIGNAL
 - GENERALLY WANTS TO BE SMOOTHER

WHAT IF WE PICK BAD Ω ? IT'S OK, WE HAVE λ :)

WE HAVE TO USE VALIDATION TO CHECK

NN Ω s: WEIGHT DECAY
from linear to logical



WEIGHT ELIMINATION

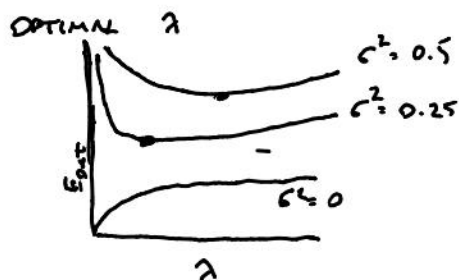
smaller λ ! soft weight elimination

$$\Omega(w) = \sum_{i,j,l} \frac{(w_{ij}^{(l)})^2}{\beta^2 + (w_{ij}^{(l)})^2}$$

EARLY STOPPING - regularisation through the optimiser!

(when to stop?!))

BEST TO SEPARATE CONCERNS :)



Stochastic $\leftarrow$ but also applies to Q_f !!

DETERMINISTIC NOISE IN OVERFITTING BEHAVES
ALMOST EXACTLY LIKE STOCHASTIC NOISE

HOMEWORK #6

1. DETERMINISTIC NOISE $\propto H$

$H' < H$, f is fixed. this means H' is less
powerful than H , so $\hookrightarrow \checkmark$

2. REGULARISATION w/ WEIGHT DECAY

(REMEMBER FORM: $(Z^T Z - \lambda I)^{-1} Z^T Y = w$)

$\phi(x_1, x_2) = (1, x_1, x_2, x_1^2, x_2^2, x_1 x_2, |x_1 - x_2|, |x_1 + x_2|)$

w/o DECAY: $E_{in} \sim 0.3$, $E_{out} \sim 0.8 \Rightarrow a \checkmark$

3. w/ $k=3$: $\Rightarrow d \checkmark$

4. w/ $k=3$: $E_{in} = 0.37$, $E_{out} = 0.43 \Rightarrow c \checkmark$

5. best k : $-1 \Rightarrow d \checkmark$

6. $E_{out} = 0.036 \Rightarrow b \checkmark$

7. $\mathcal{H}(Q, C, Q_0)$
 constraint $\nwarrow$ stopping $\nearrow$ max dimension = 1

$$\mathcal{H}_Q: \{h \mid h(x) = w^T z = \sum_{q=0}^Q w_q L_q(x)\}$$

a) $\underbrace{\mathcal{H}(10, 0, 3)}_{\mathcal{H}_2} \cup \underbrace{\mathcal{H}(10, 0, 4)}_{\mathcal{H}_3} \Rightarrow \mathcal{H}_5$

b) $\underbrace{\mathcal{H}(10, 0, 3) \cup \mathcal{H}(10, 1, 4)}_{\text{BOTH WEIRD, } \therefore \mathcal{H}_3}$

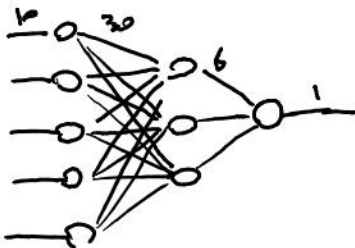
c) $\underbrace{\mathcal{H}(10, 0, 3)}_{\mathcal{H}_2} \cap \underbrace{\mathcal{H}(10, 0, 4)}_{\mathcal{H}_3} \Rightarrow \mathcal{H}_2 \checkmark$

⋮

8. NEURAL NETWORK - fully connected

$$L=2; d^{(0)}=5, d^{(1)}=3, d^{(2)}=1$$

what does this look like?



47! c ✓

NOTE: EACH HIDDEN LAYER ALSO RECEIVES
A BIAS WEIGHT. SO

$$w = 10a_1 + (a_1 + 1)a_2 + \dots + (a_{n-1} + 1)a_n$$

min 2 nodes per hidden layer.

$\therefore$ min weights is $n = 17$; 1 computing node / layer

$$w = 10 \times 2 + \dots + 2 \cdot 1$$

$$20 + 17 \times 2 + 2 = 46 \quad a \quad \checkmark$$

Maximum would be $n = 2$, 2 bias nodes $\hat{=}$ 34 compute

$$10a_1 + (a_1 + 1)a_2 + (a_2 + 1)$$

$$10a_1 + a_1a_2 + 2a_1 + 1$$

$$10a_1 + a_1(34 - a_1) + 2(34 - a_1) + 1$$

$$10a_1 + 34a_1 - a_1^2 + 68 - 2a_1 + 1$$

$$42a_1 - a_1^2 + 69 = 0$$

$$-2a_1 + 42 = 0 \Rightarrow a_1 = 21, a_2 = 13$$

$$210 + 286 + 14 = 510 \Rightarrow c \quad \checkmark$$

LECTURE 13: VALIDATION

- TO SOLVE λ in regularisation we need to take the Lagrangian
- remember E_{avg} is a better proxy than E_{in}
- TODAY WE'LL FOCUS ON CHOOSING λ in
$$E_{\text{avg}}(h) = E_{\text{in}}(h) + \lambda/N \Omega(h) - \text{validation}$$

VALIDATION

- used in pretty much every machine learning application
- the validation set
- model selection
- cross-validation

VALIDATION VS. REGULARISATION

$$E_{\text{out}}(h) = E_{\text{in}}(h) + \text{penalty}$$

regularisation: estimates overfit penalty

validation: estimates $E_{\text{out}}(h)$

Analysing the estimate

on o-o-s point (x, y) , the error is $e(h(x), y)$

squared: $(h(x) - y)^2$

binary: $\mathbb{I}[h(x) \neq y]$

$$\mathbb{E}[e(h(x), y)] = E_{\text{out}}(h)$$

← unbiased error (but not necessarily accurate)

$$\text{var}[e(h(x), y)] = \sigma^2 \leftarrow \text{estimate on one point}$$

FROM POINT TO SET

on a validation set $(x_1, y_1), \dots, (x_k, y_k)$,
the error $E_{\text{val}}(h) = \frac{1}{k} \sum_{k=1}^k e(h(x_k), y_k)$ ← $k = \text{size of validation set}$

$$\mathbb{E}[E_{\text{val}}(h)] = \frac{1}{k} \sum_{k=1}^k \mathbb{E}[e(h(x_k), y_k)] = E_{\text{out}}(h)$$

we need to improve variance - does going from point to set help us?

$$\text{var}[E_{\text{val}}(h)] = \frac{1}{k^2} \sum_{k=1}^k \text{var}[e(h(x_k), y_k)] = \frac{\sigma^2}{k} \leftarrow \text{better!}$$

← covariants drop out, just use diagonals

estimate reliability depends -

$$\text{on } k \quad E_{\text{val}}(h) = E_{\text{out}}(h) \pm O\left(\frac{1}{\sqrt{k}}\right)$$

k is not free - we separate it out of the training set, so given data set

$$D = (x_1, y_1), \dots, (x_n, y_n)$$

D_{val} { k points $\rightarrow$ validation, chosen randomly

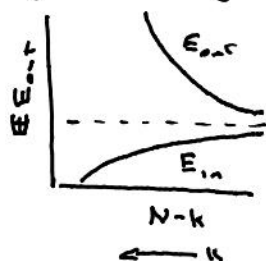
$\underbrace{N - k \text{ points}} \rightarrow \text{training}$

D_{train}

$$D_{val} \cup D_{train} = D$$

$O\left(\frac{1}{\sqrt{k}}\right)$: small $k \Rightarrow$ bad estimate

large $k \Rightarrow$ generalisation cost



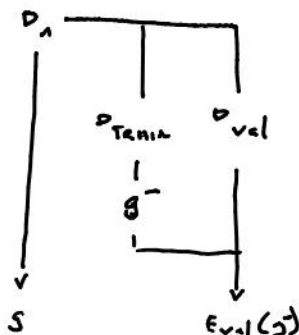
getting a better estimate of some quantity!

IS THERE A TRICK TO AVOID THIS?

TAKE k , THEN PUT BACK INTO D_{train}

$$D \rightarrow D_{train} + D_{val}$$

$$D \Rightarrow g \quad D_{train} \Rightarrow g^-$$



is $E_{val}(g^-)$
 $\sim E_{val}(g)$?

$E_{val} = E_{val}(\hat{f})$ Large $K \Rightarrow$ bad estimate!

so how do we pick K ?

rule of thumb: $K = N/5$

WHY IS THIS CALLED VALIDATION?

• we use it to make choices

• if an estimate of E_{out} affects learning the set is no longer a test set! it's a validation set

WHAT'S THE DIFFERENCE?

TEST SET IS UNBIASED; VALIDATION SET HAS OPTIMISTIC BIAS. TWO HYPOTHESES, h_1 & h_2 w/ $E_{out} = 0.5$

estimate error e_1, e_2 uniform on $[0, 1]$
(expected error 0.5)

pick $h \in \{h_1, h_2\}$ with $c = \min(e_1, e_2)$

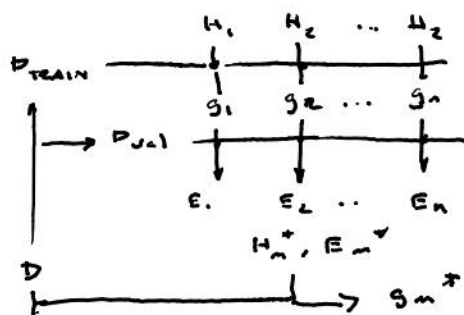
$E(c) < 0.5 \leftarrow$ optimistic bias

USED D_{VAL} MORE THAN ONCE

M models $h_1, \dots, h_M$

use D_{TRAIN} to learn g_m^-

EVALUATE g_m^- USING D_{VAL}



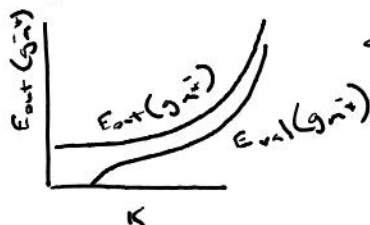
THIS ISN'T DATA SNOOPING BECAUSE IT'S PART OF THE LEARNING PROCESS? AND HOW DO WE CONTROL BIAS?

$$E_m = E_{val}(\bar{g}_m); \quad m=1, \dots, M$$

pick model $m = m^*$

THE BIAS

WE SELECTED H_n^* USING D_{val}



WHY IS E_{out} GOING UP?? BIGGER $K \rightarrow$ SMALLER N

WHY ARE THE CURVES CONVERGING? BIGGER $K \rightarrow$ BETTER ESTIMATE

FOR M models: $H_1, \dots, H_M$ D_{val} IS USED FOR TRAINING ON THE FINALISTS MODEL:

$$H_{val} = \{\bar{g}_{val1}, \bar{g}_{val2}, \dots, \bar{g}_{valM}\}$$

BACK TO HOEFFDING AND VC

$$E_{out}(\bar{g}_n^*) \leq E_{val}(\bar{g}_n^*) + O\left(\sqrt{\frac{\ln(M)}{K}}\right)$$

CAN M BE INFINITE? YES!

IS SIMPLE PARAMETER w/ VC dimension = 1!

regularisation parameter λ ; early-stopping T

THE MORE PARAMS WE ADD TO VALIDATION THE LESS ACCURATE

DATA CONTAMINATION

ERROR ESTIMATES: $E_{in}, E_{test}, E_{val}$

CONTAMINATION: OPTIMISTIC (DECEPTIVE) BIAS IN ESTIMATING E_{out}

TRAINING SET: totally contaminated

TEST " : totally 'clean'

VALIDATION " : slightly contaminated :>
(multiple sets!?)

CROSS - VALIDATION

$$E_{out}(g) \approx E_{in}(g^-) \approx E_{val}(g^-)$$

(small k) (large k)

HOW DO WE DO BOTH?

"LEAVE ONE OUT"

USE $N-1$ POINTS FOR TRAINING; 1 POINT FOR VALIDATION

$$D_n = (x_1, y_1), \dots, (x_{n-1}, y_{n-1}), (\overline{x_n}, \overline{y_n}), (x_{n+1}, y_{n+1}), \dots, (x_N, y_N)$$

final hypothesis learned from D_n is $\bar{g}_n$

$$e_n = E_{\text{val}}(\bar{g}_n) = c(\bar{g}_n(x_n), y_n)$$

WHAT IF WE DO THIS ACROSS ALL POINTS?

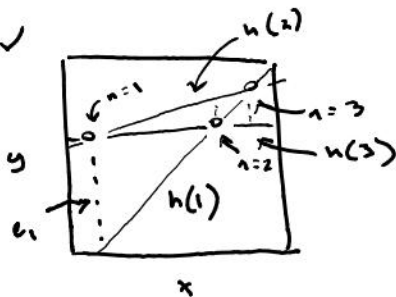
HYPOTHESES ARE ALL DIFFERENT BUT ALL TRAINED ON $N-1$ POINTS

WE KNOW THEY'RE RELATED w/ E_{in}

cross-validation error: $E_{\text{cv}} = \frac{1}{N} \sum_{n=1}^N e_n$ ← NOT INDEPENDENT BUT SEEMS NOT TO MATTER !

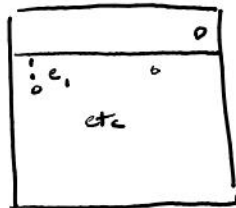
ILLUSTRATION

Linear:



$$E_{\text{cv}} = \frac{1}{3} (e_1 + e_2 + e_3)$$

constant



check E_{cv} and see which model is better

CROSS-VALIDATION IN ACTION

digit-classification requires non-linear transform

$(1, x_1, x_2) \rightarrow$ something giant

use cross-validation to choose which features to use!



w/o VALIDATION:

$$E_{in} = 0\%, E_{out} = 2.5\%$$

w/ VALIDATION

$$E_{in} = 0.8\%, E_{out} = \underline{\underline{1.5\%}}$$

GENERALLY... LEAVE MORE THAN ONE OUT



so we only run k TIMES: N/k or $N-k$ points each. in practice 10-fold cross validation works well

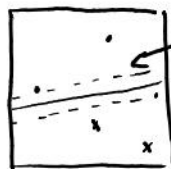
LECTURE 14: SUPPORT VECTOR MACHINES

• MOST SUCCESSFUL CLASSIFICATION TOOL IN ML

- MAXIMISING THE MARGIN
- THE SOLUTION
- NONLINEAR TRANSFORMS

BETTER LINEAR SEPARATION

DATA



many possible lines to separate the data. is there an advantage in choosing any particular one?

imagine $\frac{1}{2} \frac{1}{2}$ as the margin. this is a measure of robustness to noise! so we maximise margin. ①

which w maximises the margin? ②

① REMEMBER THE GROWTH FUNCTION?

- CHOOSING BIG MARGIN CONSTRAINS DICHOTOMIES!

$\therefore$ SEPARATING POINTS WITH LARGE MARGINS DROPS

D_{VC} !

③ FINDING w

LET x_n BE THE NEAREST DATA POINT TO THE PLANE $w^T x = 0$. DISTANCE IS MARGIN

2. PRELIMINARY TECHNICALITIES:

1. NORMALISE w :

$$|w^T x_n| > 0 \rightarrow |w^T x_n| = 1$$

i. w^T has a scale invariance but non-unit w will cause problems later

ii. expanding $w^T x_n$ currently gives a 'distance' but it's not Euclidean; normalising gives us a proper distance

2. Pull out w_0 ← no longer convenient to use, so $w = (w_1, \dots, w_d)$ apart from b

The plane is now $\boxed{w^T x + b = 0}$ (no x_0)

COMPUTING THE DISTANCE

THE DISTANCE BETWEEN x_n and plane $w^T x + b = 0$

WHERE $|w^T x_n + b| = 1$.

$w \perp$ TO THE PLANE IN X SPACE

TAKE $x' \notin x''$ on plane

$$w^T x' + b = 0, w^T x'' + b = 0 \Rightarrow w^T (x' - x'') = 0 \Rightarrow w^T \text{ orthogonal}$$

and the distance is...

take any point on the plane. take the projection of $x_n - x$ on w

$$\hat{w} = \frac{w}{\|w\|} \Rightarrow \text{distance} = |\hat{w}^T (x_n - x)|$$

$$\text{distance} = \frac{1}{\|w\|} |w^T x_n + b - w^T x - b| \quad \text{what are the } b\text{'s doing?}$$

making $w^T x$ terms go away!

$$\Downarrow$$
$$\frac{1}{\|w\|}$$

②

$$\text{subject to } \min_{n=1,2,\dots,N} |w^T x_n + b| = 1$$

notice: optimising mins = :c

$$|w^T x_n + b| = y_n (w^T x_n + b)$$

③ SIGNAL AGREES WITH
VALUE

EQUIVALENT TO

$$\text{minimise } \frac{1}{2} w^T w$$

$$\text{④ subject to } y_n (w^T x_n + b) \geq 1$$

for $n = 1, 2, \dots, N$

⑤ THESE ARE EQUIVALENT BUT IT'S MESSY!

CONSTRAINED OPTIMISATION

$$\text{MINIMISE } \frac{1}{2} w^T w$$

$$\text{SUBJECT TO } y_n (w^T x_n + b) \geq 1 \text{ for } n=1, 2, \dots, N$$

$$w \in \mathbb{R}^d, b \in \mathbb{R}$$

LAGRANGE? inequality constraints $\Rightarrow$ KKT

WE'VE SEEN THIS BEFORE WITH REGULARISATION

$$\text{MINIMISE } E_n(w) = \frac{1}{N} (Zw - y)^T (Zw - y)$$

$$\text{subject to } w^T w \leq c$$

∇E_n normal to constraint

	OPTIMISE	CONSTRAIN
REGULARISATION:	E_n	$w^T w$
SUM:	$w^T w$	$E_n \leftarrow \text{LOL}$

LAGRANGE FORMULATION

$$\min \rightarrow \mathcal{L}(w, b, \alpha) = \frac{1}{2} w^T w - \sum_{n=1}^N \alpha_n (y_n (w^T x_n + b) - 1) \quad \leftarrow \text{slack}$$

wrt w, b and maximise w.r.t. each $\alpha_n \geq 0$

$$\nabla_w \mathcal{L} = w - \sum_{n=1}^N \alpha_n y_n x_n = 0 \quad \leftarrow \text{conditions}$$

$$\frac{\partial \mathcal{L}}{\partial b} = - \sum_{n=1}^N \alpha_n y_n = 0$$

Substitute conditions into original Lagrangian

$$\mathcal{L}(w, b, \alpha) = \frac{1}{2} w^T w - \sum_{n=1}^N \alpha_n (y_n (w^T x_n + b) - 1)$$

$$\mathcal{L}(\alpha) = \sum_{n=1}^N \alpha_n - \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N y_n y_m \alpha_n \alpha_m x_n^T x_m$$

① is multiplied by $-\sum_{n=1}^N \alpha_n y_n = 0$, so goes away

maximize wrt α subject to $\alpha_n \geq 0$ for $n=1, \dots, N$
and $\sum_{n=1}^N \alpha_n y_n = 0$

SOLUTION: QUADRATIC PROGRAMMING

$$\max_{\alpha} \sum_{n=1}^N \alpha_n - \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N y_n y_m \alpha_n \alpha_m x_n^T x_m$$

$\Downarrow$

$$\min_{\alpha} \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N y_n y_m \alpha_n \alpha_m x_n^T x_m - \sum_{n=1}^N \alpha_n$$

$$\min_{\alpha} \quad \frac{1}{2} \alpha^T \underbrace{\begin{bmatrix} y_1 y_1 x_1^T x_1 & \dots & y_1 y_N x_1^T x_N \\ y_2 y_1 x_2^T x_1 & & \vdots \\ \dots & & \\ y_N y_1 x_N^T x_1 & \dots & y_N y_N x_N^T x_N \end{bmatrix}}_{\text{quadratic coefficients } Q} \alpha + (-1^T) \alpha$$

subject to $\underbrace{y^T \alpha = 0}_{\text{linear constraint}}$

$$0 \leq \alpha \leq \infty$$

WARNING: Q is $N \times N$ ← uh oh! many heuristics to deal with big N

QP hands us α

$$\alpha = \alpha_1, \dots, \alpha_N$$

$$\Rightarrow w = \sum_{n=1}^N \alpha_n y_n x_n$$

KKT condition: FOR $n=1, \dots, N$

$$\alpha_n (y_n (w^T x_n + b) - 1) = 0$$

EITHER SLACK IS ZERO OR α IS ZERO

FOR EACH POINT, INTERIOR POINTS HAVE $\alpha = 0$

$\alpha_n > 0 \Rightarrow x_n$ is a support vector

SUPPORT VECTORS

closest x_n 's to plane: achieve the margin

$$\Rightarrow y_n(w^T x_n + b) = 1$$

$$w = \sum_{x_n \text{ is SV}} \alpha_n y_n x_n$$

α is model parameters,
only taking support factors
drop params a lot!

solve for b using any SV:

$$y_n(w^T x_n + b) = 1$$

NONLINEAR TRANSFORMS

SVMs give linear boundaries

transform $x \rightarrow z$!

what happens to SVM formulation?

$$\mathcal{L}(\alpha) = \sum_{n=1}^N \alpha_n - \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N y_n y_m \alpha_n \alpha_m \underbrace{z_n^T z_m}_{\text{only change}}$$

SVM optimisation is dimension agnostic!
(apart from inner product of $z_n^T z_m$)

"SUPPORT VECTORS" in x space

support vectors live in space you're analysing (2-space)

in x space, "pre-images" of support vectors

THE MARGIN EXISTS ONLY IN z SPACE :)

GENERALISATION RESULT

$$\mathbb{E}[E_{\text{out}}] \leq \frac{\mathbb{E}[\#SV]}{N-1}$$

EXTRAS: DISTANCE $\hat{=}$ LAGRANGIAN

$$\frac{1}{\|w\|} |w^T(x_n - x)| \quad \textcircled{1}$$

$$w^T x + b = 0 \quad \textcircled{2} \quad (x \text{ is on hyperplane})$$

$$\therefore w^T x = -b$$

$$\text{EXPAND } \textcircled{1} \Rightarrow w^T(x_n - x) \Rightarrow w^T x_n - w^T x$$

$$\text{SUBSTITUTE } \textcircled{2} \Rightarrow w^T x_n + b$$

BUT WE KNOW THAT $w^T x + b$ ARE FREELY SCALABLE, SO WE CAN SIMPLY CHOOSE TO SCALE IT SUCH THAT $w^T x_n + b = 1$; $\therefore \text{DISTANCE} = \frac{1}{\|w\|}$

$$|w^T x_n + b| = y_n(w^T x_n + b) \leftarrow \text{for all correctly classified points}$$

replacing minimum over $n=1$ w/ all N quantities are at least 1 is saying that each support vector is 1 and rescaling until we get there

LAGRANGE

- we have a constrained function
- build a new unconstrained function that encodes penalty term w/ a multiplier. so

$$L(x, a) = f(x) - a \cdot g(x) \quad \text{w/ } a \geq 0$$

then minimising $f(x)$ subject to $g(x) \geq 0$ is equivalent to finding a saddle point of L

we need to maximise $L(x, a)$ over $a \rightarrow M(x)$

then minimise $M(x)$ over x

$$\Rightarrow \min_x (\max_a (L(x, a)))$$

FOR SVM LAGRANGIAN

$$L(w, b, a) = f(w, b) - \sum a_n (\text{constraint}_n) \quad a_n \geq 0$$

FIX SOME PARTICULAR (w, b) . WHAT IS

$$\max_a (L(w, b, a))?$$

CASE 1: (w, b) satisfies constraint

$\hookrightarrow -a_n \text{constraint}_n$. set a_n to 0

whenever constraint is satisfied to

maximise negative sum

CASE 2: (w, b) violates constraint

- a_n constraint_n $\Rightarrow$ positive, maximise by $a_n \Rightarrow \infty$
so $\max_a (L(w, b, a)) = f(w, b)$ when constraints
are satisfied and ∞ when not.

this builds an infinite wall when we come
to minimise x !

then back substitute as per your
lecture notes

NOTE: THE VALUE OF THE LAGRANGE
PENALTY, α/a , TELLS US SOMETHING ABOUT WHERE THE
SOLUTION WANTED TO GO. IF ZERO, NOT RELEVANT.
FOR SVMs, ONLY NON-ZERO AT SUPPORT VECTORS

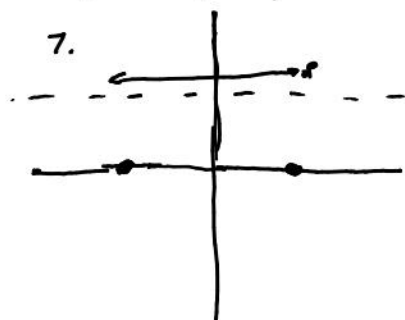
HOMEWORK #7

1. $k=6$ d ✓
2. $k=7$ c ✓
3. $k=6$ d ✓
4. $k=6$ d ✓
5. 0.1, 0.2 b ✓

6. e_1, e_2 independent random vars distributed over $[0,1]$ evenly. $\therefore E e_1, e_2 = 0.5$

but what is $E \min(e_1, e_2)$? $\Rightarrow 1/3$

$\therefore 0.5, 0.5, 0.4 \Rightarrow d$



what p is makes $h_0(x) = h_1(x)$?
 $\nearrow$ constant
 $\nearrow$ $x=p$

$h_0(x)$: 1. $-D_1 \rightarrow y = 1, 0, 1 \Rightarrow$ squared errors = $1/4$

2. $-D_2 \rightarrow y = 0, 0, 1 \Rightarrow$ " " 1

3. $-D_3 \Rightarrow D_1$

$$\therefore E_{cv}(h_0) = \frac{1}{3}(1.5) = 1/2$$

$h_1(x)$ 1. $-D_1 \rightarrow y = (x - 1/(p-1))$; $D_3 = 4/(1-p)^2$

$$x = -1; y = -2/(1-p)$$

$$e_1 = (0 - 2/(1-p))^2 = 4/(1-p)^2$$

2. $-D_2 \rightarrow y = 0$

$$x = p \Rightarrow y = 0$$

$$e_2 = 1$$

$$E_{cv}(h_1) = \frac{1}{3} \left(1 + \frac{4}{(p+1)^2} + \frac{4}{(p-1)^2} \right)$$

$$\text{SET } E_{cv}(u_i) = E_{cv}(u_o)$$

$$1/2 = 1/3 \left(1 + 4/(p+1)^2 + 4/(1-p)^2 \right)$$

$$1/2 = 4/(p+1)^2 + 4/(1-p)^2$$

$$1 = 8/(p+1)^2 + 8/(1-p)^2$$

$$= 8 \frac{(p+1)^2 + (1-p)^2}{(1-p^2)^2}$$

$$(1-p)(1+p) = 8[(p+1)^2 + (1-p)^2]$$

$$= 8[p^2 + 2p + 1 + 1 - 2p + p^2]$$

$$= 16(p^2 + 1)$$

$$u = p^2$$

$$(1-u)^2 = 16(u+1)$$

$$u^2 - 2u + 1 = 16u + 1$$

$$u^2 - 18u - 15 = 0$$

$$u = \frac{18 \pm \sqrt{324 - 60}}{2} = 9 \pm 4\sqrt{5}$$

$$\therefore p = \sqrt{9 + 4\sqrt{5}} < \checkmark$$

8. PLA vs SVM $N=10$

better 60% of the time $\rightarrow C \checkmark$

9. $N=1000$ better 50% $\rightarrow C \times$

10. 3 support vectors $\rightarrow b \checkmark$

LECTURE 15: KERNEL METHODS

remember $\rightarrow$ nonlinear transforms we only

pay for $z^T z$ inner product computation

extend SVMs further!

- KERNEL TRICK $\leftarrow$ lets us play in $\mathbb{R}^\infty$!!

- SOFT-MARGIN SVM $\leftarrow$ allow some errors

we will probably use both tricks on most SVM problems

What do we need from z -space?

dealing with inner product w/ kernel trick

$$\mathcal{L}(a) = \sum_{n=1}^N \alpha_n - \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N y_n y_m \alpha_n \alpha_m \underbrace{z_n^T z_m}_{\textcircled{1}}$$

① do we need anything other than this in z -space?

constraints: $\alpha_n \geq 0$ for $n=1, \dots, N$ & $\sum_{n=1}^N \alpha_n y_n = 0$

$g(x) = \text{sign}(w^T z + b)$ where $w = \sum_{n=1, n \neq SV}^N \alpha_n y_n z_n$
 $\uparrow$
needs $z^T z$!

b from $y_n(w^T z_n + b) = 1 \leftarrow$ needs $z^T z$!

can we get $z^T z$ without visiting z -space
or knowing what z -space is...

GENERALISED... INNER PRODUCT

given two points x and $x' \in \mathcal{X}$, we need $z^T z'$

$\swarrow$ z is $\phi(x)$; z' is $\phi(x')$

let $z^T z' = K(x, x')$ ← THE KERNEL "INNER PRODUCT" of x and x'

EXAMPLE $x = x_1, x_2 \rightarrow 2^{nd}$ ORDER ϕ

$$z = \phi(x) = (1, x_1, x_2, x_1^2, x_2^2, x_1 x_2)$$

$$K(x, x') = z^T z' = 1 + x_1 x_1' + x_2 x_2' + x_1^2 x_1'^2 + x_2^2 x_2'^2 + x_1 x_1' x_2 x_2'$$

THE TRICK - CAN WE COMPUTE $K(x, x')$ w/o

TRANSFORMING x & x' ?

Consider $K(x, x') = (1 + x^T x')^2$ ← not inner product in $\mathcal{X}$

$$= (1 + x_1 x_1' + x_2 x_2')^2 \quad \textcircled{1}$$

$$= 1 + x_1^2 x_1'^2 + x_2^2 x_2'^2 + 2x_1 x_1' + 2x_2 x_2' + 2x_1 x_1' x_2 x_2'$$

① THIS IS AN
INNER PRODUCT

$\swarrow$

$$z = \phi(x) = (1, x_1^2, x_2^2, \sqrt{2}x_1, \sqrt{2}x_2, \sqrt{2}x_1 x_2)$$

THE POLYNOMIAL KERNEL

$X = \mathbb{R}^d$ and $\Phi: X \rightarrow \mathbb{Z}$ is a polynomial of order Q
"equivalent" kernel $K(x, x') = (1 + x^T x')^Q$
 $= (1 + x_1 x'_1 + x_2 x'_2 + \dots + x_d x'_d)^Q$
 $\uparrow$ computationally easy

for $d=10$ and $Q=100$ explicitly is... big,
unpleasant, etc. so don't do the transform.

ok but we have $\sqrt{}$ terms $\subset \dots$ but we can
scale $K(x, x') = (a x^T x' + b)^Q$

WE JUST NEED $\mathbb{Z}$ TO EXIST!

if $K(x, x')$ is an inner product in some space
 $\mathbb{Z}$, this works!

$$K(x, x') = \exp(-\gamma \|x - x'\|^2)$$

$\uparrow$ infinite-dimensional $\mathbb{Z} \leftarrow$ wat.

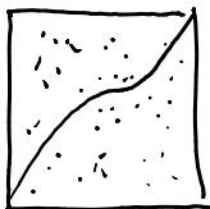
TAKE SIMPLE CASE

$$\begin{aligned} K(x, x') &= \exp(-(x - x')^2) \\ &= \exp(-x^2) \exp(-x'^2) \sum_{k=0}^{\infty} \frac{2^k (x)^k (x')^k}{k!} \quad (\gamma=1) \end{aligned}$$

split into two vectors and distribute
constants and it's an inner product!

IN ACTION

slightly non-separable case



Transform X into ∞ -dimensional Z . pass to
quad programming and we get support vectors
in this case 9 support vectors, 100 points
 $E_{\text{out}} \sim 9/91$

overkill! count support vectors :)

kernel formulation of SVM

take quadratic coefficients $[Q]$. $Q_{ij} = y_i y_j x_i^T x_j$

w/ kernel: $[Q'] = Q_{ij} = y_i y_j K(x_i, x_j)$

nothing else changes

FINAL HYPOTHESIS: $g(x) = \text{sign}(w^T z + b)$ in terms of $K(-, -)$

$$w = \sum_{z_n \text{ is SV}} \alpha_n y_n z_n \Rightarrow g(x) = \text{sign} \left(\sum_{\alpha_n > 0} \alpha_n y_n K(x_n, x) + b \right)$$

$$\text{where } b = y_n - \sum_{\alpha_n > 0} \alpha_n y_n K(x_n, x_n)$$

KERNEL CHANGES MODEL :)

SO HOW DO WE CHOOSE KERNELS? WHICH KERNELS ARE VALID?

$\exists$ EXISTS FOR GIVEN $K(x, x')$...

1. by construction
2. math properties of kernel $\leftarrow$ Mercer's condition
3. who cares? :) $\leftarrow$ probably don't do this but it works sometimes

DESIGN YOUR OWN KERNEL

$K(x, x')$ IS A VALID KERNEL IF & ONLY IF

1. it is symmetric (dot product works :))
2. the matrix $K_{ij} = [K(x_i, x_j)]$ is positive semi-definite for any $x_1, \dots, x_n$ (Mercer's condition)

SOFT-MARGIN SVM

TWO TYPES OF NON-SEPARABLE

-slightly

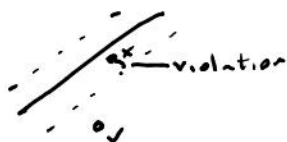
-very

KERNELS DEAL WITH VERY NON-SEPARABLE DATA.

SOFT-MARGIN DEALS WITH SLIGHTLY "

ERROR MEASURE

margin violation (# inviolated points is silly)
even if correctly classified!



QUANTIFY: $y_n(w^T x_n + b) \geq 1 - \xi_n \quad \xi_n \geq 0$

TOTAL VIOLATION: $\sum_{n=1}^N \xi_n$

NEW OPTIMISATION

minimise $\frac{1}{2} w^T w + C \sum_{n=1}^N \xi_n$

①

← augmented error for SVM, can set C as $1/2$

for big C , don't violate margin at all!

for small, can violate willy-nilly

① subject to $y_n(w^T x_n + b) \geq 1 - \xi_n$ for $n = 1..N$

← forces non-negative ξ ; no credit for non-violation

and $\xi_n \geq 0$ for $n = 1..N$

$w \in \mathbb{R}^d, b \in \mathbb{R}, \xi \in \mathbb{R}^N$

NOW WE HAVE TO GO THROUGH LAGRANGIAN AGAIN!

$$\mathcal{L}(w, b, \xi, \alpha, \beta) = \frac{1}{2} w^T w + C \sum_{n=1}^N \xi_n - \sum_{n=1}^N \alpha_n (y_n (w^T x_n + b) - 1 + \xi_n) - \sum_{n=1}^N \beta_n \xi_n$$

①

① constraint on $\xi \leftarrow$ must be ≥ 0

MINIMIZE w.r.t w, b, ξ ; MAXIMIZE w.r.t $\alpha_n \geq 0$ & $\beta_n \geq 0$

$$\nabla_w \mathcal{L} = w - \sum_{n=1}^N \alpha_n y_n x_n = 0 \quad (\text{same as before})$$

$$\frac{\partial \mathcal{L}}{\partial b} = - \sum_{n=1}^N \alpha_n y_n = 0$$

$$\frac{\partial \mathcal{L}}{\partial \xi_n} = C - \alpha_n - \beta_n = 0 \quad \leftarrow \text{forces terms to drop out to gives us back original SVM Lagrangian!!!}$$

only new constraint is that $\alpha_n \leq C$

$$\mathcal{L}(\alpha) = \sum_{n=1}^N \alpha_n - \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N y_n y_m \alpha_n \alpha_m x_n^T x_m \quad \text{w.r.t } \alpha$$

subject to $0 \leq \alpha_n \leq C$ for $n=1, \dots, N$

$$\text{& } \sum_{n=1}^N \alpha_n y_n = 0$$

minimises $\frac{1}{2} w^T w + C \sum_{n=1}^N \xi_n$

TYPES OF SUPPORT VECTORS

① margin support vectors ($0 < \alpha_n < C$)

$$y_n(w^T x_n + b) = 1 \quad (\xi_n = 0)$$

② non-margin support vectors ($\alpha_n = C$)

$$y_n(w^T x_n + b) < 1 \quad (\xi_n > 0) \leftarrow \text{margin violated}$$

all margin violators are support vectors!

get C from cross-validation

can we pick kernels similarly?

TECHNICAL OBSERVATION

1. HARD MARGIN - WHAT IF DATA NON-SEPARABLE?

"primal $\rightarrow$ dual" breaks down

check if solution satisfies data

2. Z: what if there is w_0 ?

All goes to b and $w_0 \rightarrow 0$

EXTRA: MERCER'S CONDITION: POSITIVE SEMI-DEFINITE MATRICES HAVE ALL EIGENVALUES ≥ 0 . WE'RE LOOKING FOR A SUM OF SQUARED LENGTHS FOR ALL VECTORS, SO THAT MUST BE POSITIVE. MERCER'S THEOREM STATES THAT THIS MATRIX CAN BE EXPRESSED AS A DECOMP OF EIGENFUNCTIONS - AN INNER PRODUCT

LECTURE 16: RADIAL BASIS FUNCTIONS

CAPTURES A PARTICULAR UNDERSTANDING ON INPUT SPACE. ALSO A GOOD SUMMARY OF THE WHOLE COURSE!

- NEAREST NEIGHBOURS
- NEURAL NETWORKS
- KERNEL METHODS
- REGULARISATION

BASIC RBF MODEL

EACH POINT IN YOUR DATA SET $(x_i, y_i) \in D$ influences $h(x)$ based on $\|x - x_n\|$ i.e. there is more influence on nearby data points than ones far away.

Standard form

$$h(x) = \sum_{n=1}^N w_n e^{(-\gamma \|x - x_n\|^2)}$$

"RADIAL" BECAUSE WE WORK IN A CIRCLE AROUND x ;
 $c^{\wedge}$ term is "BASIS FUNCTION"

THE LEARNING ALGORITHM

FINDING $w_1, \dots, w_n$:
$$h(x) = \sum_{n=1}^N w_n \exp(-\gamma \|x - x_n\|^2)$$

based on $D = (x_1, y_1), \dots, (x_n, y_n)$

DRIVING $E_n = 0$: $h(x_n) = y_n$ for $n=1, \dots, N$:

$$\sum_{m=1}^N w_m \exp(-\gamma \|x_n - x_m\|^2) = y_n$$

THE SOLUTION HAS N equations. AND UNKNOWN ARE w s $\leftarrow$ we have N of these. so we put in matrix form

$$\underbrace{\begin{bmatrix} \exp(-\gamma \|x_1 - x_1\|^2) & \dots & \exp(-\gamma \|x_1 - x_N\|^2) \\ \vdots & & \vdots \\ \exp(-\gamma \|x_N - x_1\|^2) & \dots & \exp(-\gamma \|x_N - x_N\|^2) \end{bmatrix}}_{\phi} \underbrace{\begin{bmatrix} w_1 \\ \vdots \\ w_N \end{bmatrix}}_w = \underbrace{\begin{bmatrix} y_1 \\ \vdots \\ y_N \end{bmatrix}}_y$$

if ϕ invertible $w = \phi^{-1} y$ 'exact interpolation'

WHAT DOES γ DO? WHAT HAPPENS WHEN γ IS SMALL?
GAUSSIAN IS WIDE/FLAT

WHEN γ LARGE, GAUSSIAN IS SPIKE. POINTS HAVE LESS INFLUENCE ON EACH OTHER

THIS IS A REGRESSION MODEL. WHAT ABOUT CLASSIFICATION?

$$h(x) = \text{sign} \left(\sum_{n=1}^N w_n \exp(-\gamma \|x - x_n\|^2) \right)$$

use the same technique as when we used linear regression for classification - focus on signal before sign; make the signal match ± 1 target

$$s = \sum_{n=1}^N w_n \exp(-\gamma \|x - x_n\|^2)$$

minimise $(s - y)^2$ on D ; $y \in \{-1, 1\}$

$$h(x) = \text{sign}(s)$$

RELATIONSHIP WITH NEAREST-NEIGHBOUR

adopt the y value of a nearby point: what label of closest point

tessellates decision plane into cells.

γ large might do something like this? smoothening is like applying k_m .

RBF w/ K PARAMETERS

N parameters w/ N data points. but generalisation?

can we pick fewer parameters? take k "centers"

$k \ll N$, and focus on $\mu_1, \dots, \mu_k$ instead of $x_1, \dots, x_N$

$$h(x) = \sum_{k=1}^K w_k \exp(-\gamma \|x - \mu_k\|^2)$$

but now μ_k is extremely complex - who do we manage? how do we choose centers? what about getting the weights?

choosing the centers

minimise the distance between x_n and closest center μ_k
clustering! k-means clustering algorithm

split $x_1, \dots, x_n$ into clusters $S_1, \dots, S_K$

$$\text{minimise } \sum_{x_n \in S_k} \|x_n - \mu_k\|^2$$

↑
sum over points in cluster $\mathcal{D}$

then sum over all clusters to get

$$\sum_{k=1}^K \sum_{\mathcal{D}} \|x_n - \mu_k\|^2$$

good news: unsupervised learning! no reference to labels.

bad news: np-hard! uh oh

so we need an iterative plan

LLOYD'S ALGORITHM

iteratively minimise $\sum_{k=1}^K \sum_{x_n \in S_k} \|x_n - \mu_k\|^2$ wrt μ_k, S_k

$$\textcircled{1} \mu_k \leftarrow \frac{1}{|S_k|} \sum_{x_n \in S_k} x_n$$

now freeze μ_k s and build new clusters

$$\textcircled{2} S_k \leftarrow \{x_n : \|x_n - \mu_k\| \leq \text{at } \|x_n - \mu_x\|\}$$

BOTH $\textcircled{1}$ & $\textcircled{2}$ REDUCE ERROR. WE'LL CONVERGE D/C
WE'RE ON A FINITE # OF POINTS, THEREFORE FINITE
CONFIGURATIONS. CONVERGENCE $\rightarrow$ LOCAL MINIMUM
TRY DIFFERENT STARTING POINTS AND SEE WHICH IS BEST.

LLOYD'S ALGORITHM IN ACTION

1. GET DATA POINTS
2. ONLY INPUTS. NO LABELS
3. INITIALISE CENTERS (CHOOSING 9 TO COMPARE w/ SVM)
(RANDOM PLACEMENT)
4. ITERATE
5. AFTER CONVERGENCE WE HAVE OUR μ_k

RBF CENTERS COMPARED TO SVM - THIS IS SUPERVISED VS
UNSUPERVISED LEARNING

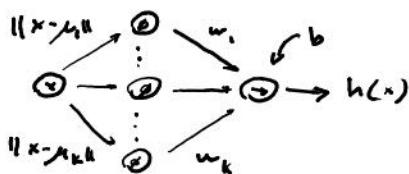
NOW WE NEED TO GO TO THE LABELS TO FIND WEIGHTS

$$\sum_{k=1}^K w_k \exp(-\gamma \|x_1 - \mu_k\|^2) \approx y_1 \quad - N \text{ equations in } K \ll N \text{ unknowns}$$

$$\phi w \approx y$$

if $\phi^T \phi$ invertible, $w = (\phi^T \phi)^{-1} \phi^T y$

RELATION TO NEURAL NETWORKS



you CAN IMAGINE AN SVM WORKS LIKE THIS - KERNEL SELECTION is 1st LAYER, THEN 2nd IS LINEAR COMBO. EVEN A KERNEL FOR NN. (TWO LAYERS)

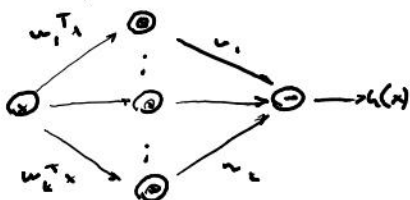
THE FEATURES ARE $\exp(-\gamma \|x - \mu_k\|^2)$

nonlinear transform depends on D

$\Rightarrow$ no longer a linear model

a bias term (b or w_0) is often added

compare w/ neural network



N.N. always contribute RBF depends on distance

ϕ is closer to linear because inputs not based on labels

CHOOSING γ

Treat γ as parameter to be learned

$$h(x) = \sum_{k=1}^K w_k \exp(-\gamma \|x - \mu_k\|^2)$$

CAN MINIMISE USING GRADIENT DESCENT, BUT w_k IS SIMPLE, SO MAYBE WE SPLIT w_k FROM γ ?

ITERATIVE APPROACH: ~ EM ALGORITHM

1. FIX γ , solve w_k using our pseudo mv
2. FIX $w_1, \dots, w_K$, minimise error
3. ITERATE !)

we can have a different γ_k for each center μ_k

RBFS & KERNEL METHODS

SVM KERNEL IMPLEMENTS

$$\text{sign} \left(\sum_{n=1}^N \alpha_n y_n \exp(-\gamma \|x - x_n\|^2) + b \right)$$

RBFS

$$\text{sign} \left(\sum_{k=1}^K w_k \exp(-\gamma_k \|x - \mu_k\|^2) + b \right)$$

SVM TRACKS TARGET FUNCTION BETTER % OF UNSUPERVISED LEARNING & BORDER CENTERS

RBFs AND REGULARISATION

RBF can be derived purely based on regularization

$$\sum_{n=1}^N (h(x_n) - y_n)^2$$

SMOOTHNESS CONSTRAINT
 $\int_{-\infty}^{\infty} \left(\frac{d^k h}{dx^k} \right)^2 dx$
big $\rightarrow$ not smooth

$\Downarrow$

$$\sum_{n=1}^N (h(x_n) - y_n)^2 + \lambda \sum_{k=0}^{\infty} a_k \int_{-\infty}^{\infty} \left(\frac{d^k h}{dx^k} \right)^2 dx$$

IF YOU SAVE THIS... YOU GET RADIAL BASIS
FUNCTION. YOU'RE LOOKING FOR INTERPOLATION; NOT THAT
SURPRISING IT'S GAUSSIAN

EXTRAS - GO BACK TO NEURAL NETWORK
COMPARISON $\frac{1}{2}$ INTERNALISE:

LINEAR MODELS, SUPPORT VECTOR MACHINES,
AND SHALLOW NEURAL NETWORKS ARE
INSTANCES OF ONE TEMPLATE: A LAYER OF
FEATURE DETECTORS FEEDING INTO A LINEAR
COMBINATION LAYER. THE DIFFERENCE IS HOW
THESE FEATURES ARE CHOSEN!

PROBLEM SET #8

1. PRIMAL VS DUAL PROBLEM

N is size of the data set, d is dimensionality of input space. The original formulation of the hard margin SVM problem w/o going through the Lagrangian dual formulation is ...
 $\frac{1}{2} w^T w$; $y(x^T w) \geq 1$.

$$w^T = \begin{bmatrix} w_1 \\ \vdots \\ w_d \end{bmatrix} + w_0 \quad \therefore d+1 \text{ variables } d \quad \checkmark$$

2. see polynomial-kernels.py in repo

max E_{in} from target: $\{0, 2, \dots, 8\}$

E_{in} target 0: $0.1059 \rightarrow a$ ✓

3. min E_{in} from target: $\{1, 3, \dots, 9\}$

E_{in} target 1: $0.0144 \rightarrow a$ ✓

4. # SVs for target 0: 2178

" 1: 386

$\Delta \sim 1800 \rightarrow c$ ✓

5. RANGING C FROM $\{0.001, 0.01, 0.1, 1\}$

MAX C ACHIEVES LOWEST $E_{in} \rightarrow d$ ✓

6. WHEN $C=0.001$, # support vectors is lower for $Q=5 \rightarrow b \checkmark$

7. see `crossval_soft_margin.py`

C selection over 100 runs:

$C: 0.0001 \rightarrow 0$

$C: 0.001 \rightarrow 47 \checkmark \quad b \checkmark$

$C: 0.01 \rightarrow 27$

$C: 0.1 \rightarrow 8$

$C: 1 \rightarrow 18$

8. average E_{cv} for $C=0.001 \rightarrow 0.0048 \rightarrow 0.5 \rightarrow C \checkmark$

9. see `rbf_kernel.py`

C	E_{in}	E_{out}
0.01	0.0038	0.0236
1	0.0045	0.0212
100	0.0032	0.0189
10^4	0.0026	0.0236
10^6	0.0006	"

$\uparrow$
 E_{in} min when $C=10^6 \rightarrow c \checkmark$

10. E_{out} min when $C=100 \rightarrow c \checkmark$

