

Physics 693: Modern ML for Physics & Astronomy

ML is revolutionizing nearly every field of science & society more broadly

powerful new tool — enabling new analyses previously impossible

— enhancing sensitivity & precision

— accelerating simulation & inference

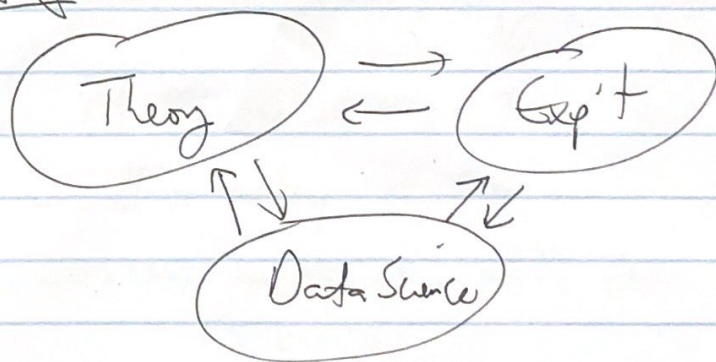
— unifying solutions to problems across domains

↓
made possible by

- Big data
- Computing (GPU)
- Sophisticated algorithms (NNs, ...)

ML like a telescope!
or microscope

~~This course:~~



~~This course:~~

↓ formal mathematics
novel mix of ~~theory~~, statistics, numerics
in ML same paper can have all 3!

This course: general ML concepts

~~lectures~~

popular & state of the art architectures

(I know more $\longrightarrow$ less)
applications to LHC, Astro, Cosm, Condensed Matter

^{required}
- No HWs but will be hands on tutorials & exercises
& no exams

- BB + demos/slides

- Will ask people to pitch applications to present in class
(either you can present, or we can learn it together & I will present)

need input!

- No domain knowledge of any subfield required

- Prior experience w/ python, numpy, matplotlib would be good

- Lectures M-Th 10:20-11:40

- poll for makeup slt?

- in Nov some travel - either Zoom or guest lectures (4 lectures)

- will provide refs & reading on course website; no textbook

- Will not require 568 (future will) but would help if you took it
will try to go fast in overlapping content (560, backprop, ... NN basics ...)

please interrupt
frequently &
ask me lots of Q's

What is ML?

"learning from data" "glorified curve fitting"

• Data: $\vec{x}_i \in \mathbb{R}^d \leftarrow d \text{ can be high!}$
e.g. $10^2, 10^3, \dots$

↳ a vector of features

features could include labels
 $\vec{y}_i$ e.g. cat vs dog, ...

"low level"
vs

e.g.

↳ particle momenta or $x, y, z, v_x, v_y, v_z, G, BR, \dots$
@ LHC @ Gaia

"high level"
features

or



pixel intensities in image

↳ "tabular data"

$i = 1, \dots, N$

of instances

(e.g. events @ LHC)

stars in Gaia

images ~~for~~ ~~sample~~ of galaxies

N could be large
 $10^6, 10^9, \dots$

• Assume in most applications: each $\vec{x}_i \sim p_{\text{data}}(x)$

drawn iid

parameters $\leftarrow$ fit $f(\vec{x}_i; \theta)$ to data

• Goal: minimize some "loss" or "objective" fn

$$L(\theta) = \sum_{i=1}^N \mathcal{L}(f(\vec{x}_i; \theta), \vec{y}_i)$$

wrt parameters θ .

Example: maybe want $f(x_i; \theta) = y_i$ ↙ target labels
- discrete "classification"

or could use $\mathcal{L}(\theta) = (f(x_i; \theta) - y_i)^2$ - continuous "regression"
"mean squared error (MSE)"

• many more examples to follow (& how to choose loss fn in principled way)

• Where do labels come from?

- if data is simulated have access to "ground truth"

- in actual data - human labeling? (common in real world app like Imagenet, also Astro citizen science)

- no labels or labels derived from data itself

DNN, RNN, CNN, Graph NNs
+ transformers, ...

→ building blocks

↑ Categories of ML

pretty much
solved problems
in ML

- fully supervised - all data labeled & used in training
- classifier, regression

- unsupervised - no labels!

- generative modeling, density estimation, ~~and~~ ^{some kinds of} anomaly detection

- weakly supervised - noisy labels

- semi supervised - mix of labeled & unlabeled data

- self supervised - labels generated from data e.g. contrastive loss

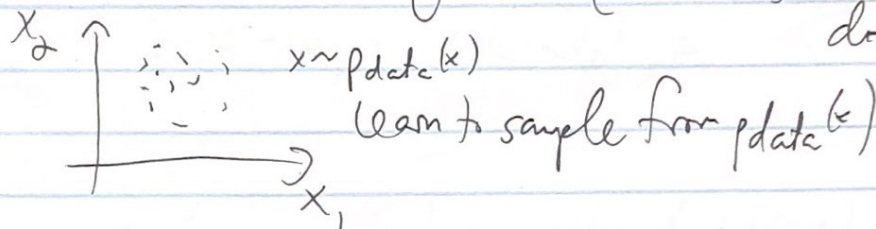
this is where
the active ML
development
is happening

less-than-sperseed

Also really important for science - strive to be as
data-driven as possible
(simulation-free)

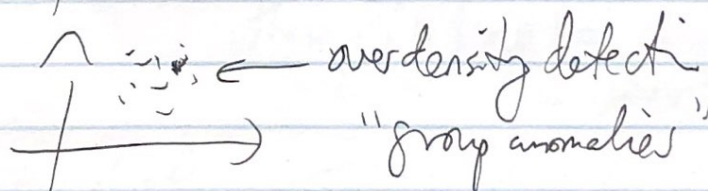
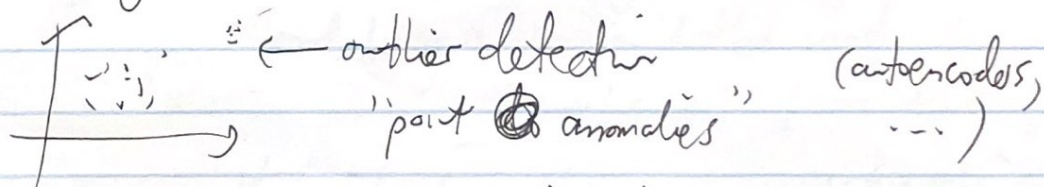
In more detail:

- Generative modeling (GANs, VAEs, normalizing flows, diffusion, ...)



- Density estimation - learn $p_{\text{data}}(x)$ (normalizing flows, ...)
(Gen modeling $\leftrightarrow$ Density est)
one doesn't guarantee the other

- Anomaly detection



General principle for loss fns: Maximum Likelihood Estimation (MLE)

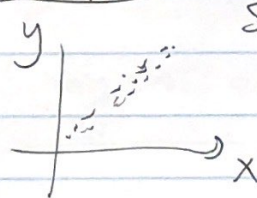
want to maximize $P(\text{data} | \text{model})$

if data iid $\prod_{i=1}^N P(x_i | \theta)$

$$L = -\log P(\text{data} | \text{model}) = -\sum_{i=1}^N \log P(x_i | \theta)$$

- MLE has properties in limit of large N :
 - consistency (as $N \rightarrow \infty$, estimated $\theta \rightarrow \text{true } \theta$)
 - efficiency (as $N \rightarrow \infty$, minimum variance estimator of θ)

Example:



Suppose want to predict y given x (regression)
model: y is gaussian dist'd around $f(x; \theta)$ w/ some width σ

$$\text{then } P(y | x, \theta) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(y - f(x; \theta))^2}{2\sigma^2}}$$

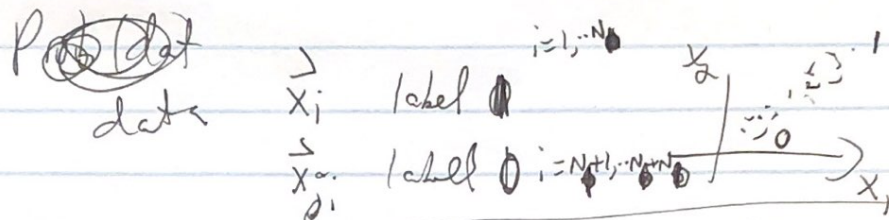
$$L = -\log P = \sum \frac{(y_i - f(x_i; \theta))^2}{2\sigma^2} + \log \sqrt{2\pi}\sigma$$

if σ known $\rightarrow$ MLE! (otherwise do fit for σ)

Ex: binary classification
~~what about bin~~

labels $y_i = 0$ or 1 from data x_i

model: $f(x_i; \theta) = \text{prob of } x_i \text{ being label } 1.$



$$P(\text{data} | \text{model}) = \prod_{i=1}^{N_0} f(x_i; \theta) \prod_{i=N_0+1}^{N_0+N_1} (1 - f(x_i; \theta))$$

$$L = -\log P = -\sum_{i \in 1} \log f(x_i; \theta) - \sum_{i \in 0} \log (1 - f(x_i; \theta))$$

$$= -\sum_i (y_i \log f(x_i; \theta) + (1 - y_i) \log (1 - f(x_i; \theta)))$$

"binary cross entropy loss" $\rightarrow$ best loss for classification

$\rightarrow$ could use MSE
 but it would be suboptimal

Lecture II

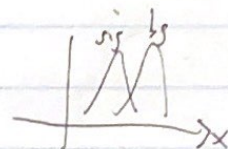
- email list
- slack

- pitches for applications

Neyman-Pearson Lemma

$$L = - \left(\sum_{i \in S} \log f(x_i) + \sum_{i \in B} \log (1 - f(x_i)) \right)$$

What f minimizes L in ideal case?



Continuum limit (∞ data)

probability of s & b

$$L = - \int dx \left[p_S(x) \log f(x) + p_B(x) \log (1 - f(x)) \right]$$

$\int dx p(x) \approx \sum_{x \sim p(x)}$ $f \rightarrow f + \delta f$, $\left. \frac{\partial L}{\partial \delta f} \right|_{f=f_*} = 0$
 $f_* \leftarrow$ optimal f

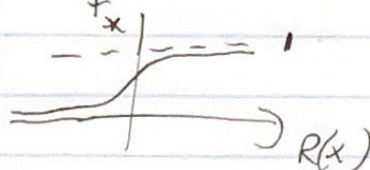
(assume $N_S = N_B$
WLOG)

$$0 = \frac{p_S(x)}{f_*} - \frac{p_B(x)}{1 - f_*} \Rightarrow \boxed{f_* = \frac{p_S(x)}{p_S(x) + p_B(x)}}$$

So optimal classifier $\odot$ $f_* = \frac{R(x)}{R(x) + 1}$ $\odot$

$$R(x) = \frac{p_S(x)}{p_B(x)}$$

f_x is monotonic w/ $R(x)$



$$\Rightarrow \boxed{f_x > f_c \Leftrightarrow R(x) > R_c} \quad \begin{array}{l} \text{cut on } f \text{ equals} \\ \text{cut on } R \end{array}$$

So optimal classifier is likelihood ratio

"Neyman-Pearson Lemma"

- powerful result! fundamental result in statistics
(LR uniformly most powerful test for simple hypothesis)

later will see:

- Can turn around $\rightarrow$ learn ^{ML} classifier from samples
 $\rightarrow$ approach LR.
"likelihood ratio trick"

- Another perspective: Bayes Thm

$$f_x = \frac{p_S(x)}{p_S(x) + p_B(x)} = \frac{p(x|S)}{p(x|S) + p(x|B)} = \frac{p(x|S) p(S)^{\frac{1}{2}}}{p(x)} = \frac{p(x|S) + p(x|B)}{2} = p(S|x)$$

So f_x is the prob of signal given x which is where we started!

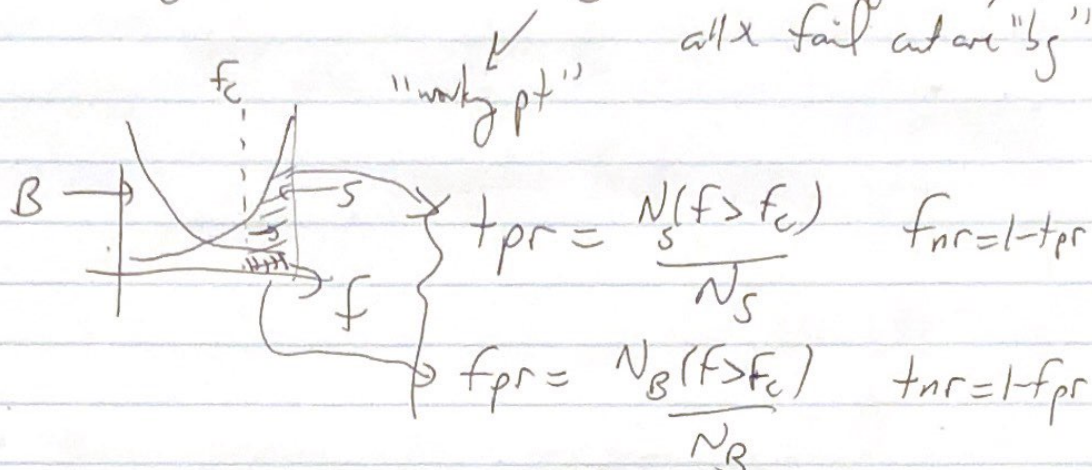
i.e. PCE is minimized when our model for this prob $\checkmark$
= true prob.

Back to NP lemma: what does "most powerful" test mean?

$\rightarrow$ metrics for binary classification

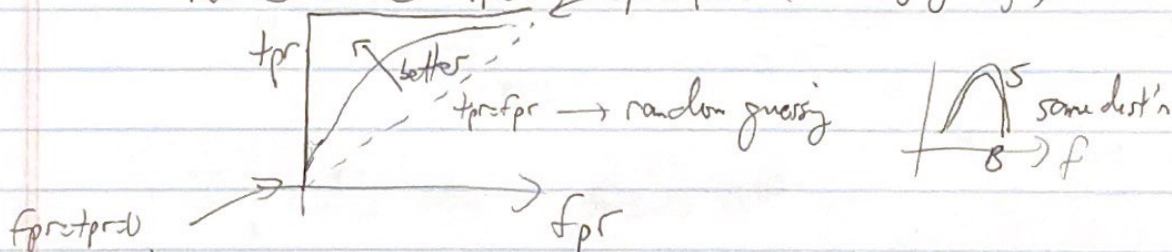
How do we use binary classifiers?

Usually: cut $f(x) > f_c$ classified as
 all x sets f_c are "good"
 all x fail cut are "bad"



• accuracy = $\max_{f_c} \frac{tpr + (1 - fpr)}{2}$

• ROC curve $tpr=1 \leftarrow fpr=tpr=1$ (let everything through)



$fpr=tpr=1$
 reject everything

Vary f_c , sweep out "ROC curve" ($fpr(f_c), tpr(f_c)$)

→ AUC "Area under the curve"

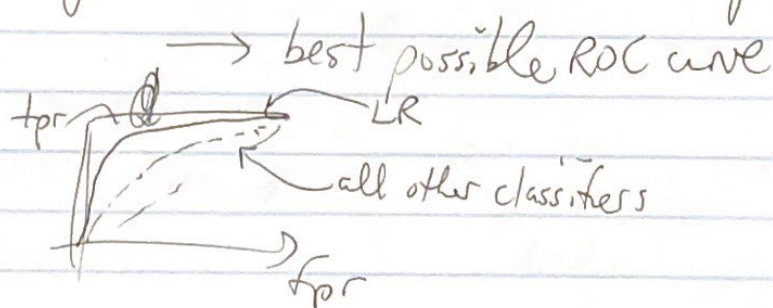
$\begin{cases} AUC = 1 \text{ perfect} \\ AUC = 0.5 \text{ random} \end{cases}$
(same for acc but $AUC \neq ACC$)

• ~~after~~ AUC, ACC mostly sensitive to $tpr, fpr \sim 0/1$ part of ROC curve

When $N_{bg} \gg N_{sig}$ (as in many cases, e.g. LHC)
care more about $fpr \ll 1$ ~~work~~ part of ROC curve

→ often report fpr @ fixed tpr e.g. $tpr=50\%$
or 30%
or rejection factor $R_{50} = \frac{1}{fpr @ tpr=50\%}$
 R_{30} etc

- Neyman-Pearson LR classifier is optimal



- In practice cannot achieve NP optimality

- exact likelihood unknown

- finite training data → ~~loss~~ of Big data

- limited model capacity (expressivity)

↳ NNs very ~~powerful~~ expressive

→ can maybe get very close to optimal!

Then possibly biased on val set

→ final metrics on another held out dataset
"test set"

Example 1: Gaussian toy model (10d)

$$\begin{cases} s_g(0) & p_g(x) = \mathcal{N}(0, 1)^{10} \\ s_g(1) & p_{s_g}(x) = \mathcal{N}(0.5, 1)^{10} \end{cases}$$

• Can compare w/ NP lemma

$$p_{NP}(s_g|x) = \frac{p_{s_g}(x)}{p_g(x) + p_{s_g}(x)}$$

• train simple ONN
on samples

- plot losses

- ROC curve

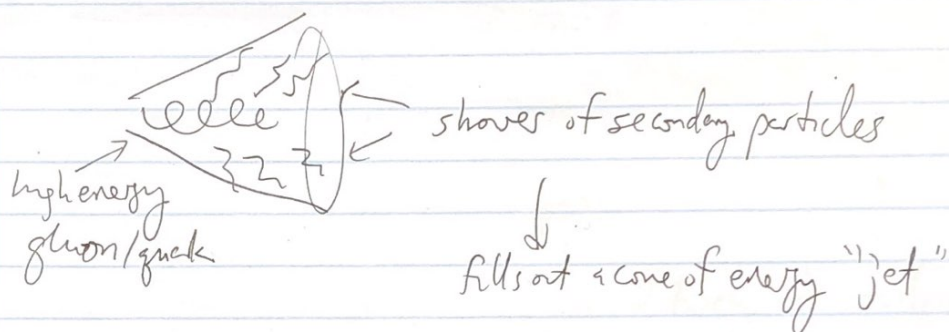
- AUC

- classifier output

• Jupyter nb demo

Example 2: Jet classification @ LHC

- Jets are essential objects @ LHC
Product of strong interaction (QCD) $\rightarrow$ showers of hadrons

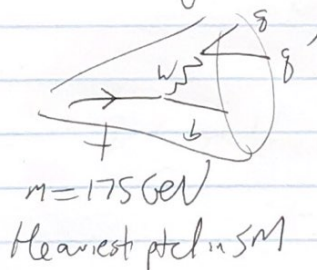


- Many different flavors of jets — important to tell them apart!
 $g, g, b, W, Z, \text{top}, \dots$

Typically QCD jets (g, g) are background & everything else is signal

SM measurements BSM searches

- Especially nice example: boosted top jets



- Decays to bW
W can decay to quarks
- If top is very energetic $\rightarrow$ boosted decay products collected "fat jet"

Note: QCD
xsec thousands
times more than
top (or more)
depend on
energy
(~100x
@ $m_j \sim 1 \text{ TeV}$)

- Top jet has substructure $\rightarrow$ expect 3 "prongs"
of energy $\approx b, g, g'$
 - expect $m_{jet} \approx m_{top}$
 - expect g, g' ~~form~~ mass $\approx m_W$
 - expect b_{jet} (won't use here)
w/ ^{its own} special properties

- Meanwhile QCD jets are relatively structureless.

$\rightarrow$ can use these properties to classify top vs QCD!

(unlike Gaussian toy model, here we do not know
true likelihoods, so do not know NP optimal classifier)

- Jupyter nb example demo
 - $(m_j, m_W, t_{32}) \rightarrow$ cuts
 - $\rightarrow$ ONN
 - jet constituents $\rightarrow$ ONN

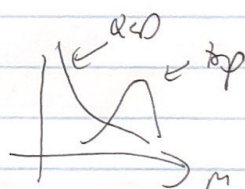
Metric: R_{30}
for @ $tpr \approx 30\%$
more informative than
AUC/ACC

A jet = collection of $p_{T,i}$ $\vec{p}$ -vectors "LLFs"

$$J_i = \begin{pmatrix} p_x^1, p_y^1, p_z^1, E^1 \\ \vdots \\ p_x^{N_i}, p_y^{N_i}, p_z^{N_i}, E^{N_i} \end{pmatrix} \quad \begin{array}{l} N_i \text{ diff from jet to jet} \\ (E^a)^2 = (\vec{p}^a)^2 \end{array}$$

HLFs like jet mass $m_i^2 = \left(\sum_{a=1}^{N_i} E^a \right)^2 - \left(\sum_{a=1}^{N_i} \vec{p}^a \right)^2$

↓ and N-subjectiness τ_{32}
physics motivated fns of LLFs



We have seen:

cut based classifier on HLFs $\approx$ DNN on HLFs

$\approx$ DNN on LLFs (30 highest p_T constituents)

"Automated feature engineering"

• DNN can construct HLFs from LLFs that perform as well as physics HLFs!

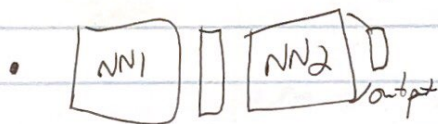
• Q: is it learning the physics HLFs or something else?
How would we know?

A: Could add physics HLFs to inputs - perf. improves?



last layer → that or like NN engineered HLFs!
 ↓
 outbanded the HLF engineer

not necessarily very last layer



↳ could interpret these as HLFs

- Is this all or can we surpass physics HLF performance?

→ yes but need more powerful NN architectures!

General Lesson: → leverage symmetry / structure of data!

our DNN p_T ordered constituents
 selected 30 hardest

↓
 like getting more data for free
 "inductive bias" → helps machine learn more effectively

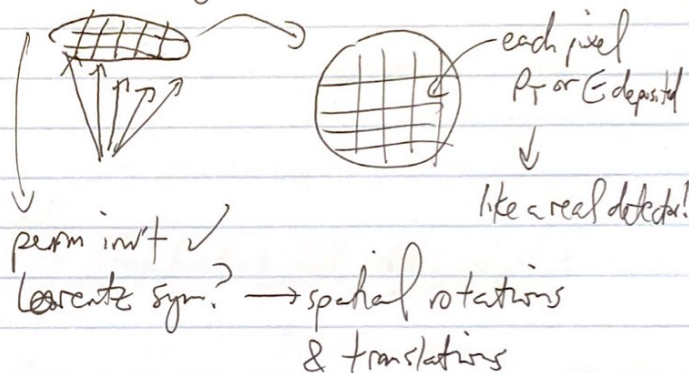
But jets have permutation invariance!

also: Lorentz symmetry → spatial rotations
 & boosts

"point cloud"

• A rich playground for different ML architectures

• CNNs — jets as images



• RNNs, LSTMs, ... — jets as sequences
— not perm inv't! what ordering?

• Graph NNs — perm inv't!
(can add Lorentz) $\xrightarrow{\quad} \rightarrow \dots$

• Deepsets — a perm inv't DNN

• Transformers — a perm inv't sequence NN.

• Plan for next few lectures + (detour into Astro w/ CNNs)