

Fudan Summer School Lectures on ML4HEP

August 2026

Motivations: AI/ML having a huge breakthrough moment right now!

LLMs, built from NNs, and trained on enormous amounts of (text) data, are exponentially expanding in capabilities $\rightarrow$ approaching AGI?

Modern ML:

Smart Algorithms/Arch. / lectures / techniques

NNs

Data

This has given us AI (LLMs, agents) but it can also do much, much more!

HEP, fundamental physics $\rightarrow$ enormous data $\rightarrow$ LHC
advantage $\rightarrow$ also simulations $\rightarrow$ neutrinos
 $\rightarrow$ well-understood (SM, GUT, ...) $\rightarrow$ astro/cosmo

$\downarrow$
ML4HEP: an unprecedented opportunity!

like a telescope/microscope for data!

see deeper & farther into data
& new better measurements, more discovery potential
faster ~~new~~ simulations

~~For~~ ~~theorists~~

for theorists: new methods, proof of concept work
lots of room for creativity & innovation!

- NOT "just" for data analysis — also for theory too increasingly
 eg. ML for simplification, Feynman loop, integral reduction
 my recent work — also theoretical "data"!
 scrambled/unscrambled expressions
- ML: an amazing new toolkit (not just tool)
~~benefit from~~ interdisciplinary (benefit from CS research)
 on par with calculus or statistics...

ML Basics: Rough Outline

- I. ML basics
- Tasks
 - NNs
 - Training

II. Data Representations — Top Tagging @ LHC example

- Images — CNNs
- Sequences — RNNs
- point clouds — GNNs & Transformers

III. Generative Models

- IV. Case studies / Applications to HEP
- Fast sim / surrogate models
 - Anomaly Detection
 - DM density from Gaia

What is ML?

but super turbocharged
P

"Learning from data" $\rightarrow$ "Global curve fitting"

We have some data $x_i \in \mathbb{R}^d$ $i=1, \dots, N$

"features"

e.g. $p_1, y_1, p_2, y_2, p_3, y_3$

or pixels

$\nwarrow$ # of instances

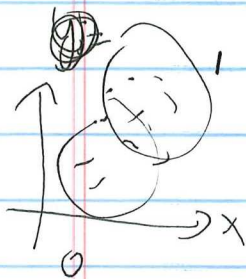
e.g. events @ LHC
or images of cats & dogs

$\circ$ generally iid from some dist'n

Want to fit a fn $f(x; \theta)$ to the data θ = parameters

Minimize some "loss" or "objective" $L(\theta) = \sum_{i=1}^N \ell(f(x_i; \theta))$ $\leftarrow$ ℓ = loss of fn.
wrt parameters θ .

- For example maybe want $f(x_i; \theta) = y_i$ $\nwarrow$ target labels
e.g. 0 for tops, cats, 1 for QCD, dogs.



$$L(\theta) = (f(x_i; \theta) - y_i)^2 \quad \text{"mean squared error"}$$

$\left\{ \begin{array}{l} y_i = \text{discrete "classification"} \\ \text{= continuous "regression"} \end{array} \right.$

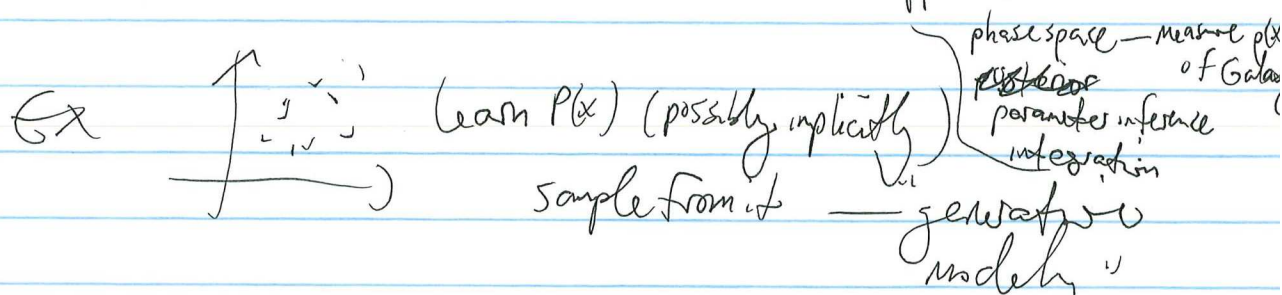
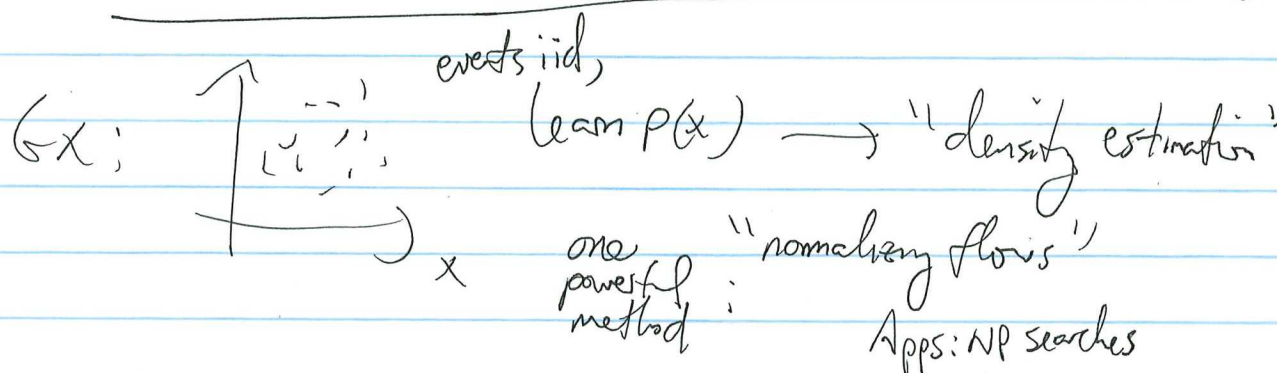
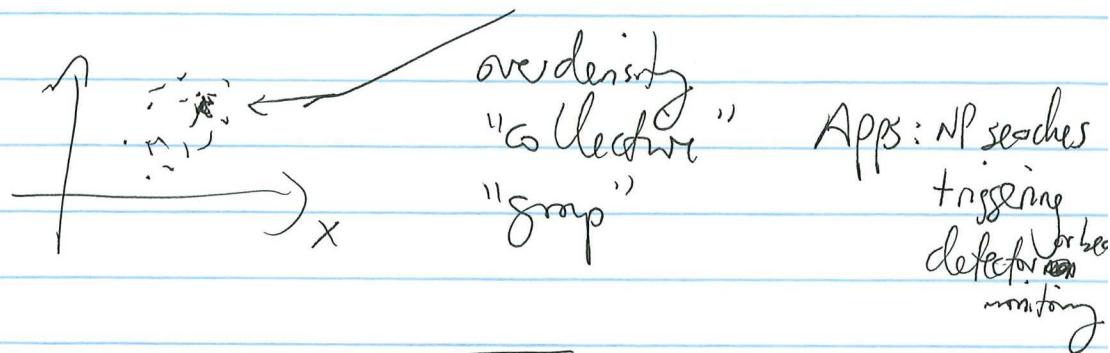
ML w/ target labels: "supervised ML"

$\left. \begin{array}{l} \text{(or any exp?)} \\ \text{in HEP, truth labels} \\ \text{generally only possible} \\ \text{w/ simulation} \end{array} \right\}$

Other kinds of ML ~~are~~ more data driven,

$\left\{ \begin{array}{l} \text{no labels - unsupervised} \\ \text{noisy labels - weakly supervised} \end{array} \right.$

$\left. \right\}$ less than supervised



Loss terms
for less than supervised
not as obvious — more later!

— "GANs", "VAEs", flows

Apps: NP searches
surrogate models (e.g. GGAN4
calo shape)

General principle for loss terms / objective fns:

Maximum Likelihood Estimation (MLE)

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

$\{x_1, x_2, \dots, x_N\}^\theta$

usually data iid, then

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

work w/

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

MLE has nice properties in lot of large data (bias min var)

Ex: regression $f(x, \theta) \rightarrow y$

suppose y is gaussian distributed around true $\hat{y}(x)$

$$P(y|x) = \mathcal{N}e^{-\frac{(y - \hat{y}(x))^2}{2\sigma^2}}$$

our model

$$P(y|x, \theta) = \mathcal{N}e^{-\frac{(y - f(x, \theta))^2}{2\sigma^2}}$$

Given data

$$\{x_i, y_i\} \quad L = -\sum_i \log P(y_i | x_i, \theta) \rightarrow + \sum_i (y_i - f(x_i, \theta))^2$$

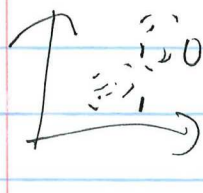
MSE loss!

Another Ex: binary classification

$$f(x; \theta) = \text{output of model}$$

data: $\{x_1, \dots, x_n\}$ class 1 $\{x_{n+1}, \dots, x_{n+m}\}$ class 0

$$P(\text{data}|\text{model}) = \prod_{i \in \text{class 1}} f(x_i; \theta) \prod_{j \in \text{class 0}} (1 - f(x_j; \theta))$$



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

$$= - \sum_{i \in \text{data}} y_i \log f(x_i; \theta) + (1 - y_i) \log(1 - f(x_i; \theta))$$

"binary cross entropy" — used for binary classification

Side note about bce: what minimizes it?

prob density of signal



prob density of y_0

$$L = \int dx (P_1(x) \log f(x; \theta) + P_0(x) \log(1 - f(x; \theta)))$$

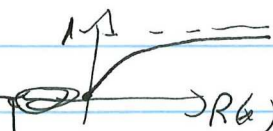
$$\delta L = \frac{P_1(x)}{f(x; \theta)} - \frac{P_0(x)}{1 - f(x; \theta)}$$

$$\Rightarrow f(x; \hat{\theta}) = \frac{P_1(x)}{P_1(x) + P_0(x)}$$

$$R(x) \equiv \frac{P_1(x)}{P_0(x)}$$

"Likelihood ratio"

$$= \frac{R(x)}{1 + R(x)}$$



So optimal classifier is (monotonic with) LR

"Neyman-Pearson lemma"