

Lecture 2 (?)

Now just need a family of fit fns $f(x; \theta)$

→ this is where ^{deep} NNs come in!

→ extremely expressive family of fns "universal approximators"

→ scale well to very high dims

→ generalize unreasonably well beyond training data $x \in \mathbb{R}^d$
 $d \sim 10^3, 10^4, \dots$

~~deep learning~~

CAN CUT

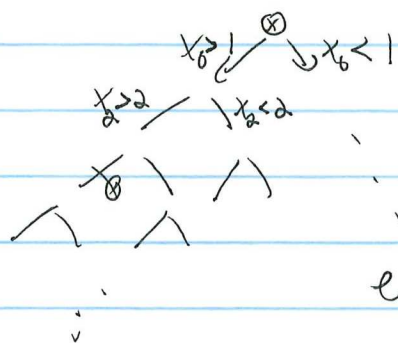
→ previous standard for ML now field: BOTS

↓
primary for classification

NNs much broader range of applications

Aside

"shallow ML"



ensemble & "boost" many weak classifiers in the form of decision trees

- limited to small # of features

- not as expressive as NNs so generally subopt

- more interpretable (just ifs)

NNs → "deep learning" vs BOTS "shallow ML"

CAN CUT

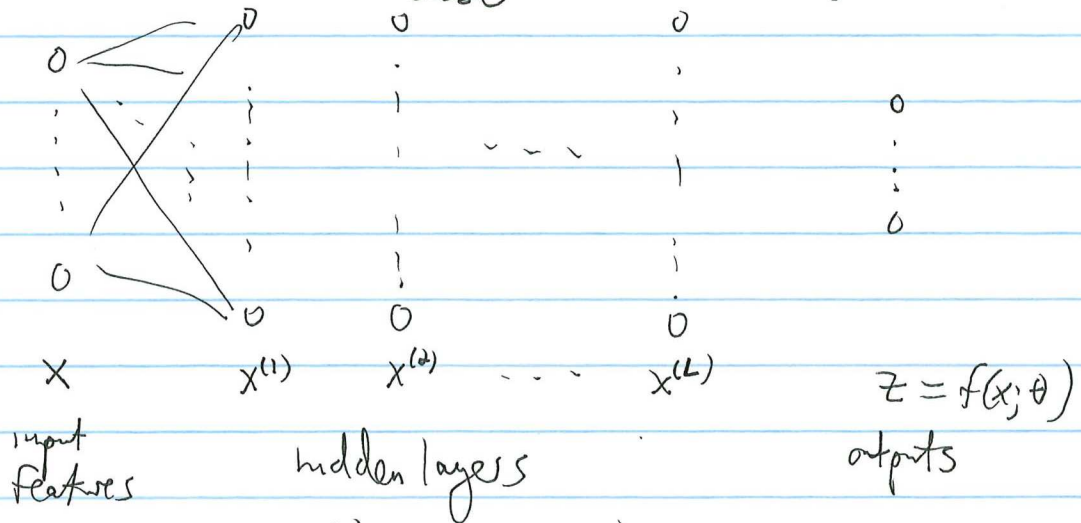
Simplest NN: "Feed-forward"

"MLP"

"DNN"

dense

"fully connected"

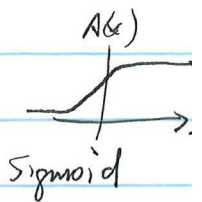


$$x^{(1)} = A^{(1)}(w^{(1)}x + b^{(1)})$$

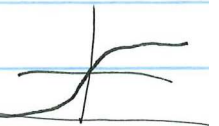
activation $\rightarrow$ nonlinearity

weights $\uparrow$ biases

$$A(x) = \frac{e^x}{1+e^x}$$



vanishing gradients $\rightarrow$ tanh

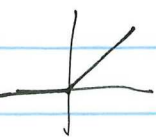


$$\{w, b\} \rightarrow \theta$$

NN trainable pars

most popular, "best" choice

$$\text{ReLU} \begin{cases} 0 & x < 0 \\ x & x > 0 \end{cases}$$



$$x^{(n)} = A^{(n)}(w^{(n)}x^{(n-1)} + b^{(n)})$$

[note: A acts elementwise]

$$z = A^{(L+1)}(w^{(L+1)}x^{(L)} + b^{(L+1)})$$

usually pick $A^{(1)}, \dots, A^{(L)}$ ReLU

$\rightarrow A^{(L+1)}$ diff. depen on application

Ex: binary classifier

$(L+1) \dots$ for known tasks

This defines a NN architecture

Many "hyperparameters"

• L : # of hidden layers

• $n_h^{(k)} = \text{dim } x^{(k)} \times \# \text{ of hidden nodes in layer } k$

• choice of activations

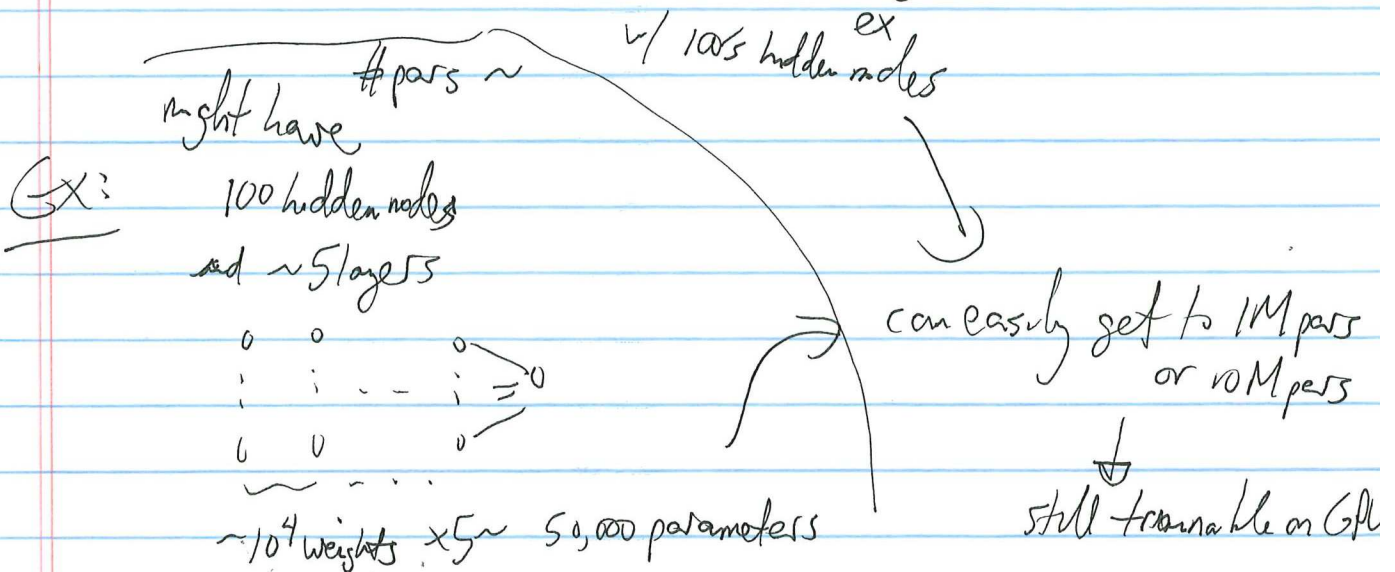
→ need to optimize w/ hyperparam

• DNN basically mtx mult + nonlinearity

→ GPUs optimized for this!

Deep learning on GPU → major breakthrough ~2010?

• typical DNN for classif. could start w/ 100 or 1000 inputs



Training NNs

- How to fit $f(x; \theta)$ w/ 50k or 1M θ 's to data??

No time to go into any detail. Just focus on rough ideas, concepts

— Most obvious approach: Newton method

$$\theta = \theta_0 + \delta\theta$$

initial guess

$$L(\theta) = \sum_{i \in \text{data}} \mathcal{L}(f(x_i; \theta))$$

$$\approx L(\theta_0) + (\theta - \theta_0) \frac{\partial L(\theta_0)}{\partial \theta} + \frac{1}{2} (\theta - \theta_0)^2 \frac{\partial^2 L(\theta_0)}{\partial \theta^2}$$

problem is θ is 1M dim!

→ cannot invert $1M \times 1M$ Hessian mtrx!

$$\frac{\partial L}{\partial \theta} + \delta\theta \frac{\partial^2 L}{\partial \theta^2} = 0$$

$$\delta\theta = - \frac{\frac{\partial L}{\partial \theta}}{\frac{\partial^2 L}{\partial \theta^2}}$$

- Instead NNs are optimized w/ 1st order methods

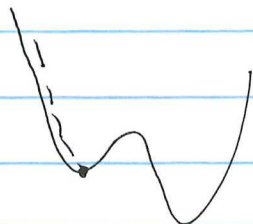
"Gradient descent"

"learning rate" — hyperparameter

"Gradient update"

$$\delta\theta = -\eta \frac{\partial L}{\partial \theta}$$

Hope that $\frac{\partial L}{\partial \theta}$ at least points in right direction



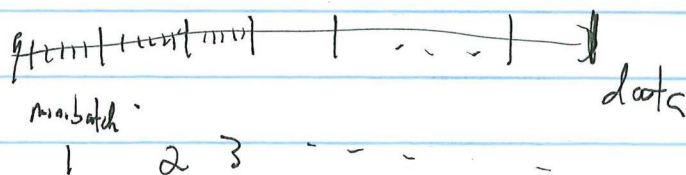
- Problem: can get stuck in local minima

Sol'n: Stochastic gradient descent: introduce noisy gradients

- don't calc $\frac{\partial L}{\partial \theta}$ w/ full data set

$$\frac{\partial L}{\partial \theta} = \frac{1}{N} \sum_{i=1}^N \frac{\partial \mathcal{L}(f(x_i; \theta))}{\partial \theta} = \left\langle \frac{\partial \mathcal{L}}{\partial \theta} \right\rangle$$

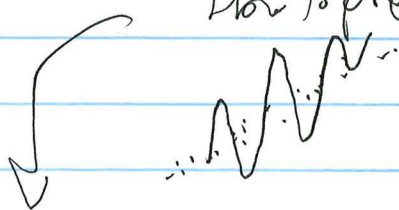
✓
 $\frac{\partial \mathcal{L}}{\partial \theta}(f(x_i; \theta))$ evaluate on single data instance — extreme version of SGD
~~loop over each~~
 evaluate on subset of data — "minibatch"



Trang: loop over ~~data~~ minibatches → 1 epoch
 shuffle and repeat → N epochs.

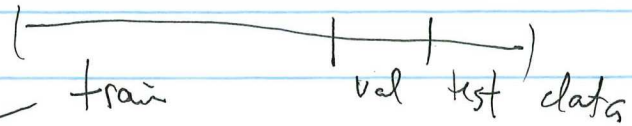
- In practice people use fancier versions of SGD
 Adam, Adadelta → think of as SGD w/
 adaptive learning rate γ .

- Overfitting: NNs have so many parameters
 How to prevent overfitting — or "memory" data.

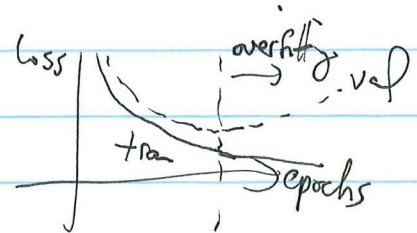


main tool: train/val/test split

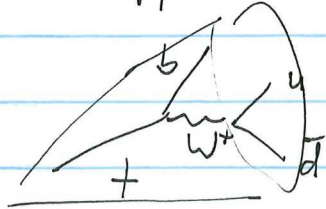
Split data into



- perform SGD on train data
- check loss on val data
stop if val loss increases
- Final evaluate on test data.

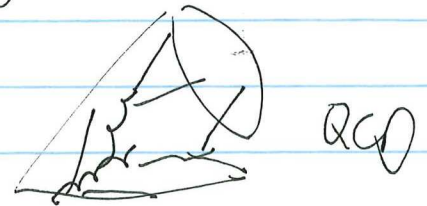


Application to HEP: Top Tagging

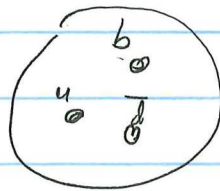


"fat jet"

vs

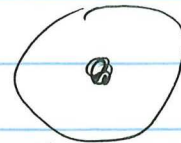


boosted



tops

3 prong substructure



QCD

no substructure (1-prong)

Good setting for developing ML classifier methods

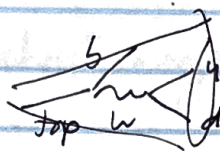
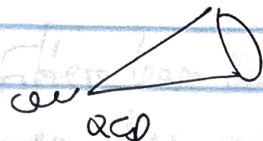
— many handles, rich substructure

— accurate simulations

Traditional approach: physics features vs "high level features"

II. Data representations

- Worky example: Boosted Top Tagging



high PT
both are "fat" jets
but tops have a lot of
substructure.

Distinguishing them is classic MLHEP task
"top tagging"

"Raw" data (unprocessed)

jet = $\{ (p_{Ti}, \eta_i, \phi_i) | i=1, \dots, N \}$
all the constituent
4 vectors

(could also include
b-tagging, particle flow, ...)

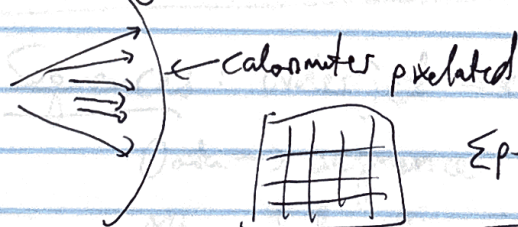
Walking through diff. approaches - great overview of standard
data representations

i) Early days: High level features $m_T, m_W, T_{3d}, \dots$

Data $\rightarrow$ "summary statistics"

Classify w/ BDT, ~~MLP~~ MLP, simple cuts

ii) Jet images



$\sum p_{Ti}$ in each 4×4 pixel
 $\rightarrow$ jet image!

Data $\rightarrow$ image

ML method: Convolutional NN (CNN)

instead of fully connected MLP:



image: X_{ij}
 $i=1, \dots, N_{\text{pixels}}$

filter W_{ab}

$a, b=1, \dots, 2$

(2x2 for example)

$$X_{ij} \rightarrow A \left(\sum_{a,b} X_{i+a, j+b} W_{ab} + b \right) = X'_{ij}$$

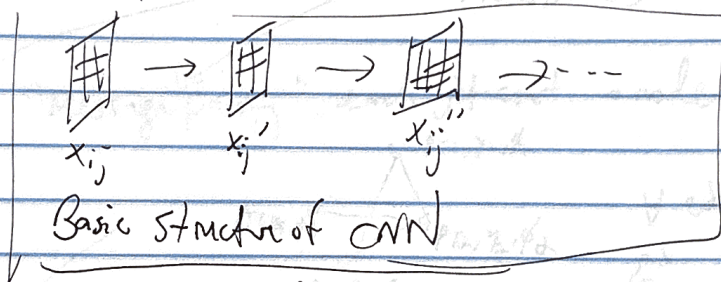
Produces hidden layer that's also an image!

Fully connected weights $\rightarrow$ 2d w/o weights of filter
 $N_{pixel} \times N_{hidden}$ "weight sharing"

Idea: Filters learn 2d "features" in image
eg edges, points, ...

These features can appear anywhere in image

Can repeat!



Also: channels "color"
padding

iii) Sequences: order just contributes somehow eg $p_{t1} > p_{t2} > \dots$

Data $\rightarrow$ sequence

ML methods: Recursive NNs (RNNs) } idea:
LSTMs ...

Stop for lack of time

weight sharing
along sequence
"in time" not 2d