

Texas A&M Institute of Data Science Tutorial Workshop Series

Introduction to Deep Learning

by Boris Hanin

June 12, 2020 

Deep Learning Tutorial

Goal: ① What is a NN?

② How NNs are used?

③ What kinds of choice are made by engineers?

④ Main use cases and failure modes?

Supervised Learning

① Data Acquisition: obtain dataset

$D = \{(x_k, f(x_k))\}_{k=1}^K$ (f unknown)
• e.g. $x_k = \text{image}$; $f(x_k) = \begin{cases} 1, & x_k \text{ has STOP} \\ 0, & \text{o/w} \end{cases}$

② Model Selection: Choose model

$$x \mapsto \mathcal{M}(x; \Theta)$$

where $\Theta = \text{vector of params}$

• e.g. $\mathcal{M}(x; \Theta) = a_1 x_1 + \dots + a_n x_n$, $\Theta = (a_i)$

③ Optimization: Use D to obtain setting Θ_* s.t.:

$$f(x_k) \approx \mathcal{M}(x_k; \Theta_*)$$

• e.g. linear regression

④ Testing: See how well Θ_* performs on unseen data:

$$f(x) \approx \mathcal{M}(x; \Theta_*)?$$

Salient Features:

- input dim $\gg 1$ ($10^2 \approx 10^6$)
- #params $\gg$ #data $\gg 1$
- NNs both interpolate and extrapolate

Neural Nets:

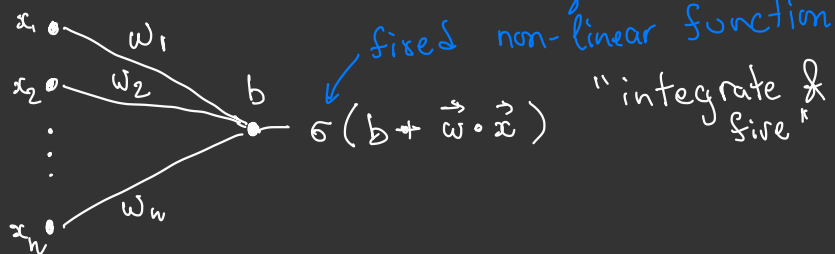
• Neural nets are built of "neurons"

$$\mathbf{x} = (x_1, \dots, x_n) \mapsto z(\mathbf{x}; \Theta) = \sigma(b + x_1 w_1 + \dots + x_n w_n)$$

$$\Theta = (b, w_1, \dots, w_n)$$

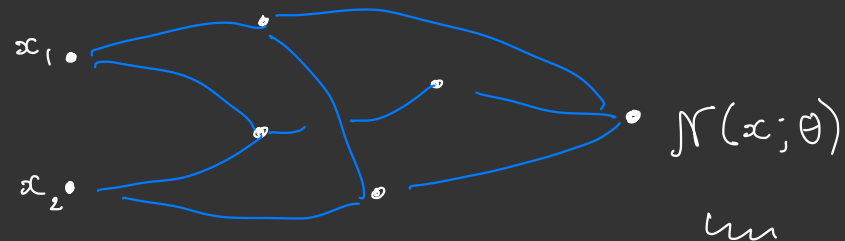
bias weights

fixed non-linear function

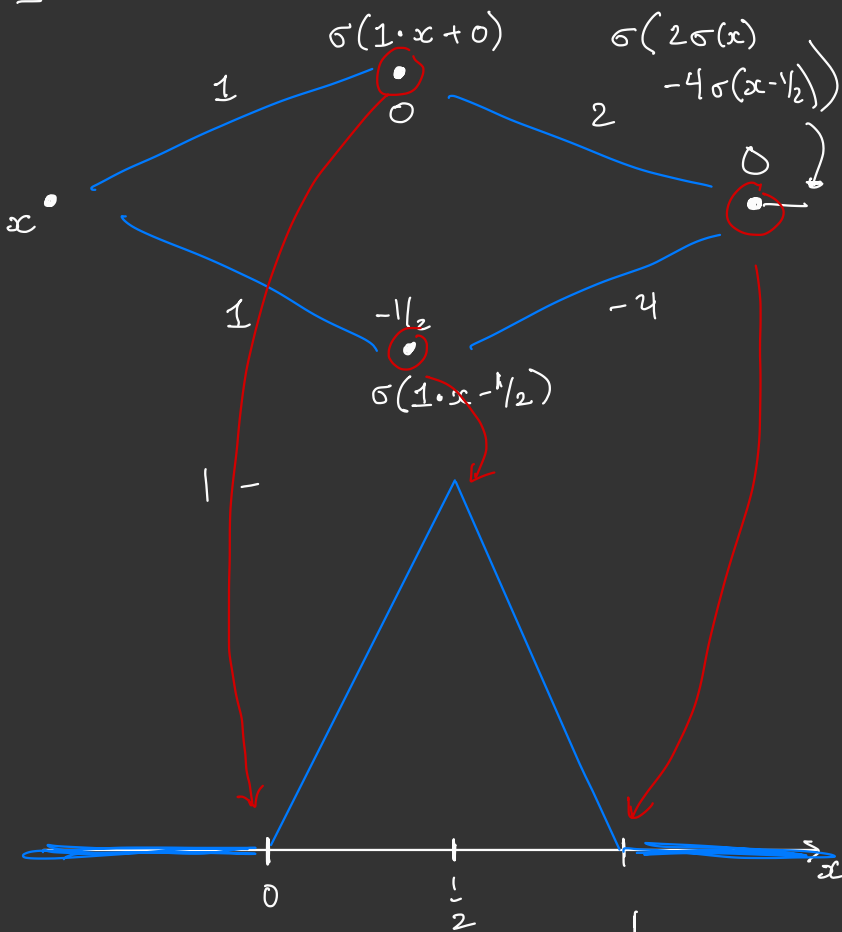


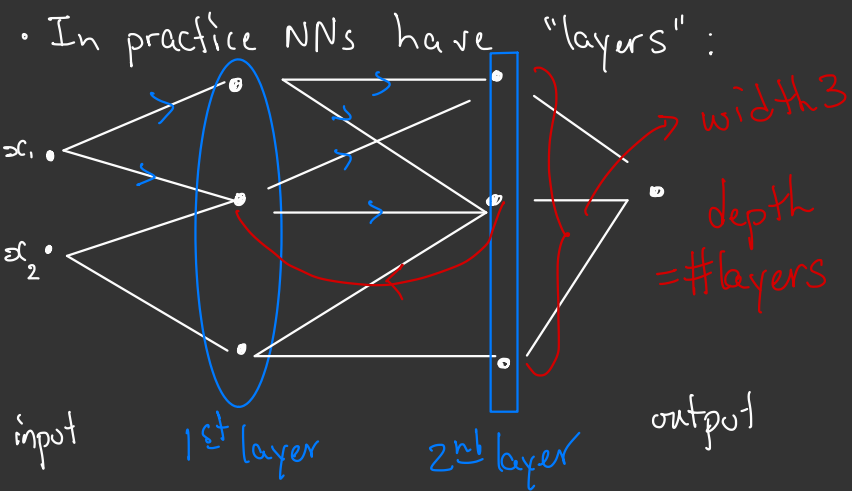
• Def A neural network is a collection of neurons and wiring diagram

"architecture" $\Theta = (\text{all } b\text{'s, } w\text{'s})$



$$\underline{\xi}_x: \sigma(t) = \text{ReLU}(t) = \max\{0, t\}$$





$$x \mapsto x^{(1)} \mapsto x^{(2)} \mapsto \mathcal{N}(x; \theta)$$

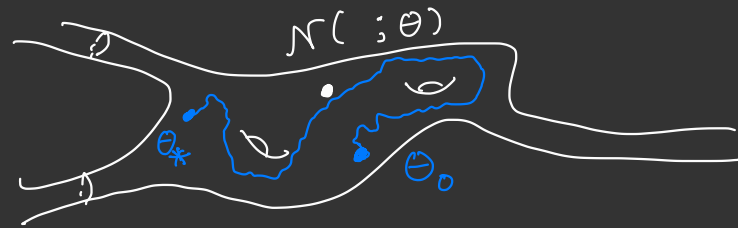
• layers $\leftrightarrow$ hierarchical reps

Typical Use: $\mathbb{R}^{n_0}$ $n_0 \gg 1$

① Data Acquisition: $\mathcal{D} = \{(\underline{x}_k, f(\underline{x}_k))\}$

② Architecture Selection: choose σ 's, wiring diagram, depth, width

③ Randomly initialize $\Theta = \{W, b\}$



③' Optimize θ by gradient descent on $\mathcal{L}_{\mathcal{D}}(\theta) = \frac{1}{|\mathcal{D}|} \sum_{x \in \mathcal{D}} l_x(\theta)$; $l_x(\theta) = |f(x) - \mathcal{N}(x; \theta)|^2$

④ Testing: Draw new $(x, f(x))$ and check whether

$$\mathcal{N}(x; \theta_*) \approx f(x)$$

Empirically: deeper is better*

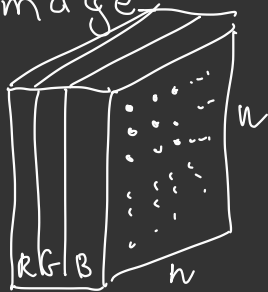
Main Uses:

① NLP

- x_k = "the cat is big" *
- $f(x_k)$ = "le chat est grand"
- Google Translate
- Siri
- Chat Bots

② Computer Vision

- x_k = image



$$f(x_w) = \begin{cases} 1 & , \quad x_w \text{ has human} \\ 0 & , \quad o/w \end{cases}$$

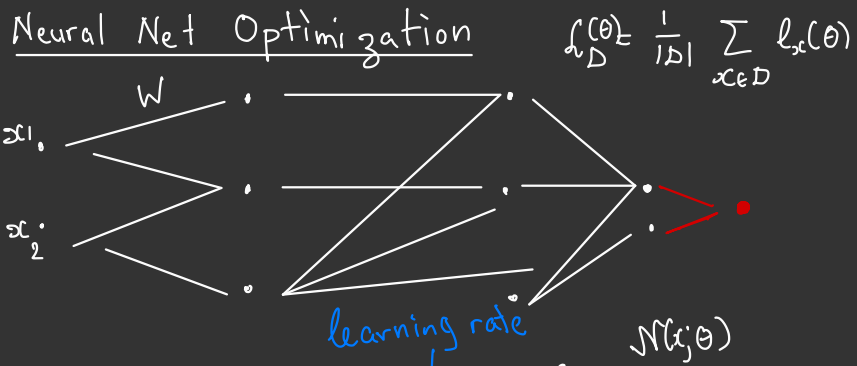
- Self Driving Cars
- Facial Rec

③ Reinforcement Learning

- x_k = state of system
(e.g. position of chess board)
- $f(x_k)$ = reward-max action
(e.g. best next move)

- Alpha Go
- Exploration by AU

Neural Net Optimization



• GD: $\Delta W = -\lambda_W \frac{\partial L_D}{\partial W}$

• Compute $\frac{\partial L_D}{\partial W}$ using "backpropagation" (i.e. chain rule)

• SGD: • randomly select "batch" $B \subseteq D$

• replace L_D by

$L_B \xrightarrow{\text{batch size}} \frac{1}{|B|} \sum_{x \in B} l_x(\theta) \approx L_D$

• good: • small batches mean less computation

• Key: How to choose λ , $|B|$?

Intuition: ① small $\lambda \leftrightarrow$ slow but accurate
large $\lambda \leftrightarrow$ fast but noisy

② Why choose same λ for all params?
 $N(x; \theta)$ might be very sensitive to θ_1 but not to θ_2 .

③ In practice: find λ by grid search on log-space

④ Why keep λ constant during train?

Schedules

- $\lambda \equiv \text{const}$
- λ decays ($\lambda \leftarrow .9\lambda$) every fixed # steps
- λ -periodic $\lambda(t) = |\cos(2\pi kt)|$

⑤ small $|B| \leftrightarrow$ noisy but fast
large $|B| \leftrightarrow$ accurate but slow

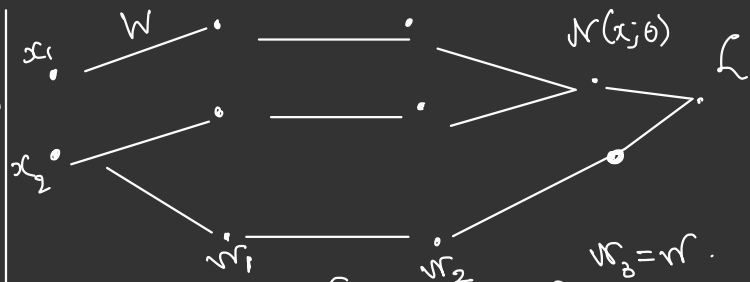
$[\lambda \cdot |B|^{\frac{1}{2}} \equiv \text{const}]$

⑥ λ , $|B|$ are inversely related

Architecture Selection:

- Best architecture is always data-dependent:
 - NLP $\leftrightarrow$ $\begin{cases} \text{Transformer-based} \\ \text{Recurrent (LSTM + Attention)} \end{cases}$
- CV $\leftrightarrow$ Convolutional, Residual
- Still leaves many choices:
 - details of wiring (width/depth...)
 - choice of σ (= ReLU)
 - how to initialize and to optimize
- Empirical: deep is good*

* = but often less stable

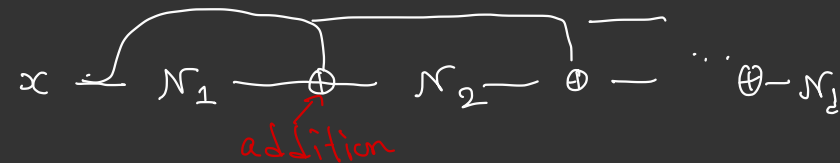


$$\frac{\partial L}{\partial W} = \frac{\partial L}{\partial r} \underbrace{\frac{\partial r}{\partial r_2} \frac{\partial r_2}{\partial r_1}}_{\text{product} \approx \# \text{ layers}} \frac{\partial r_1}{\partial W}$$

product $\approx$ # layers Jacobians

- $d \gg 1$: $\left| \frac{\partial L}{\partial W} \right| \approx 0$ or $+\infty$
"exploding / vanishing gradients"

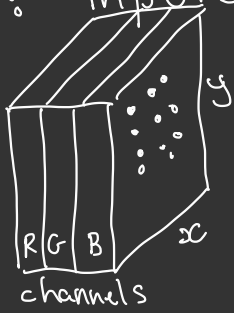
Residual Network:



$$\text{output} = x + N_1(x) + N_2(x + N_1(x)) + \dots$$

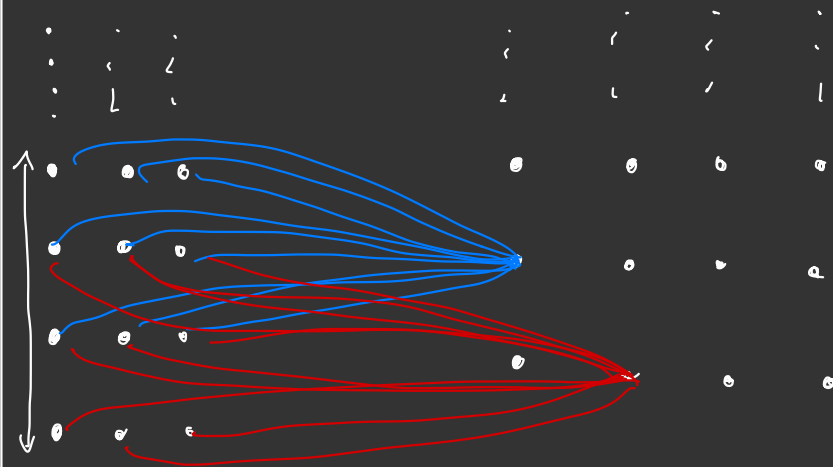
Intuition: N_j = j^{th} order correction to $x \mapsto x$

Conv Nets: inputs are $n \times n$ RGB



Key: Images are hierarchical and local

- In a ConvNet every layer has channels $x(y)$ structure
- Every neuron in layer shares weights



R G B all neurons $\rightarrow$ 1st look for same pattern

Challenges

① New Use Cases:

- PDE (fluids, physics, chemical...)
- Biology (genomics)

② Distribution Shift (nature of data change):

- change in hardware
- sunny vs. cloudy
- issue: NNs tend to be brittle

99.9% $\rightarrow$ 99.9% by brand

