

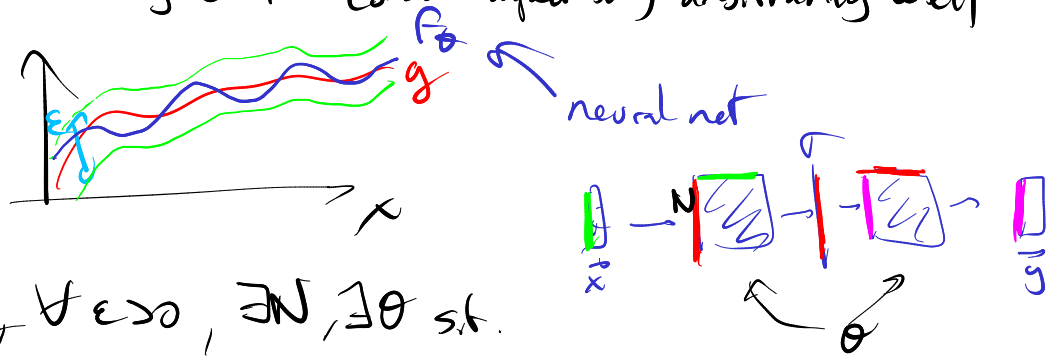
Deep Learning vs. classical ML & optimize

Properties (or lack of) of DL

- no guarantee

Universal approximation theorems (expressive power)

[Cybenko 89, Hornik 91] with just one hidden layer, one can approximate any $f \in C^0$ (on a compact set) arbitrarily well

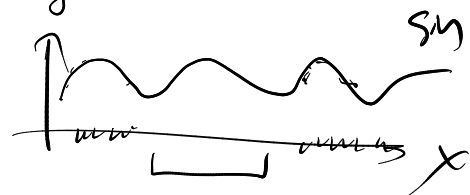


$\forall \epsilon, \forall f \in C^0, \forall \epsilon > 0, \exists N, \exists \theta$ s.t.

$$\forall x, \|f_\theta(x) - g(x)\| < \epsilon$$

Existence network

- doesn't provide the solution
- " tell whether it is easy to find it
- " " anything about generalization

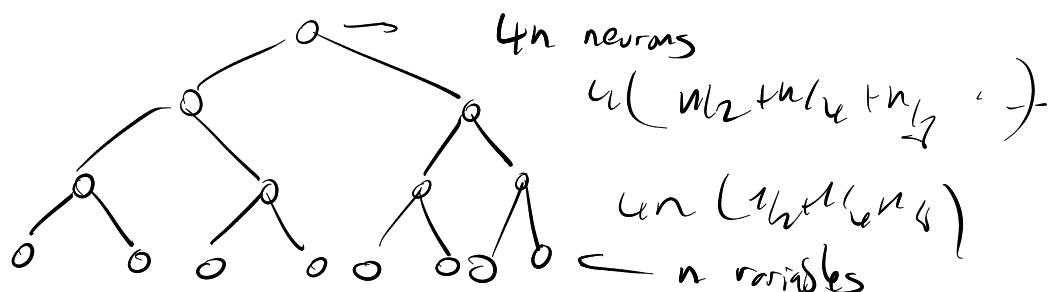


Depth simplifies the task

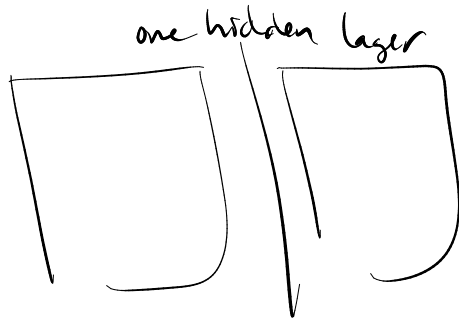
[Lin & co, 2017] multiplication of variables

- $\approx$ multiplication of 2 variables with 4 neurons

$$\sigma(\alpha \cdot \text{small}) \quad \sigma(\alpha \cdot \text{big}) \quad - \quad \alpha^2 \quad \text{big}$$



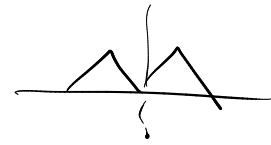
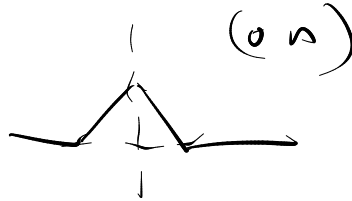
↑
n
variable



→ $\mathbb{R}$

need 2^n neurons
to be correct
on binary
inputs

[Talagrand, 2015]



→ 2^n triangles

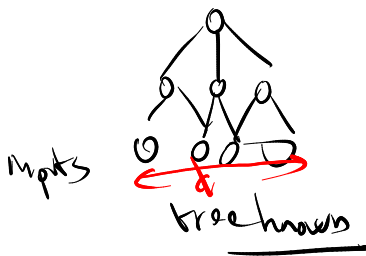
Plot network:

2 layers → $2^{n/2}$ neurons
depth $\sqrt{n}$ → $2^{\sqrt{n}}$ neurons
width

[Poggio & L., 2016]

th: if target P = computational tree
of input variables

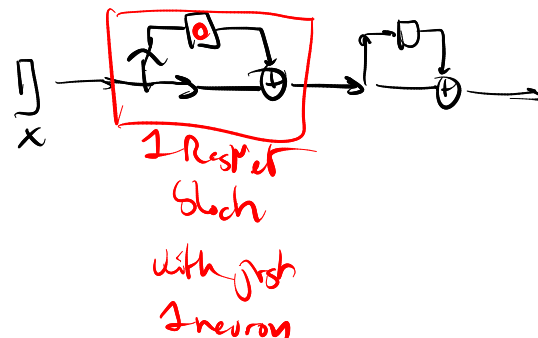
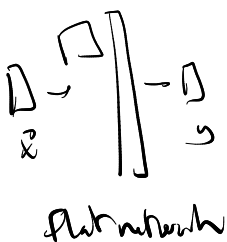
then #samples required to train a
similar-shaped network: $O(d)$



compare with e^d



[Lind & Jagelka, 2018] = depth is sufficient



universal
approx
property
(with $\sqrt{n}$ depth)

Does it work? when?

- computer vision
- NLP (text)

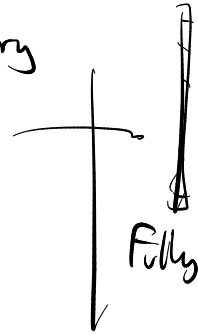
} huge success

eg: vision $\rightarrow$ CNN

- $\rightarrow$ equiv. invariant to translations
- $\rightarrow$ hierarchical models
- $\rightarrow$ handle geometry



conv



Fully-connected

classical ML

- on the opposite
data not structured

eg: medical file (age, cigarettes (days)) $\uparrow$ small

$\Rightarrow$ NN need many samples to be trained

$\Rightarrow$ dataset small $\Rightarrow$ classical ML = better

- Alpha Go : RL

$\hookrightarrow$ 2-0! - 46 millions of training games ($\gg$ one life)

$\approx$ total humanity?

- trained for only 4 hours on 5000 TPU

> 2 years $\approx$ very good GPU
if 1 TPU graphic card

≈ 1000 years of 1 GPU

- Computational power: 10^{12} parameters in GPT-4

Gap between classical ML & optim^o vs. DL

Reminder: classical ML

- dataset: samples (x_i, y_i)
- estimate a P : $F: x_i \rightarrow y_i$
- quantify "goodness" with a "loss" for training criterion

$$\Rightarrow \inf_{F \in \mathcal{F}} \mathbb{E}[\text{loss}(F(x_i), y_i)] + R(F)$$

predicted
parameterized
functional space

avoid overfit!

generalize

regularizer

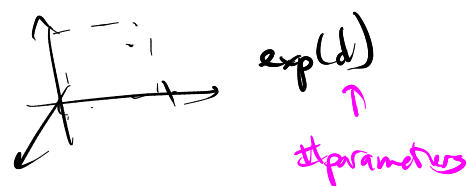
$\downarrow$
smooth

- Minimum Description Length principle (MDL):
- Occam's razor: simpler explanations are more probable

Potential issues with DL:

- 10^6 parameters $\rightarrow$ MDL?
- models: are able to overfit (easily)
- no regularizer
- possible to train huge models without overfitting?
- - - - - with fewer samples than parameters (estimator convergence)
- highly non-convex optimⁿ pb in a high-dim space

vs. classical optimⁿ:



- add noise to optimⁿ process $\Rightarrow$ better results
- we want to optim: accuracy
we train: cross-entropy } $\neq$ criterion!
- common recommendation: new task : check you are able to overfit!
architecture (small part of data)

A closer look at overfitting:

[Zheng & co, 2017] $\rightarrow$ huge models can do the job (might not overfit)
 $\rightarrow$ $\swarrow \searrow \swarrow \searrow$ "completely overfit also"

Th: dataset of n samples in $\mathbb{R}^d$

$\Rightarrow$ 2-layer network with $2n+d$ weights (with ReLU activⁿ P)
 that can represent any P on that dataset



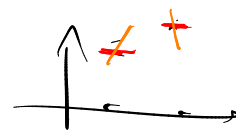
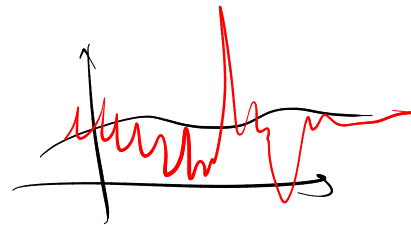
- capacity of a network = not an issue
 $\hookrightarrow$ Rademacher
 Vapnik-Ch. complexities

Palliatives for regularization

- About adding a regularizer:

$$\int_{x \in \Omega} \|\nabla_x F(x)\|^2 dx$$

$\underbrace{\int_{\Omega} p(x) dx}_{\text{prob}}$



$x \in \Omega$

$$\approx \frac{1}{N} \sum_{x \in \Omega} (\dots)$$

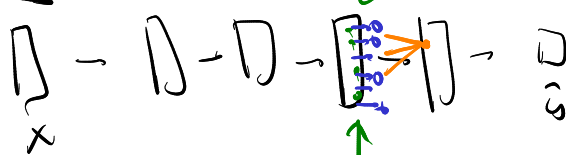
$\uparrow$ size N

trade-off
approx
of the
regularizer

th: just checking

that, over all possible binary inputs, a NN outputs always 0
is NP-hard

Dropout



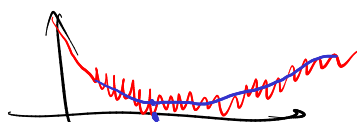
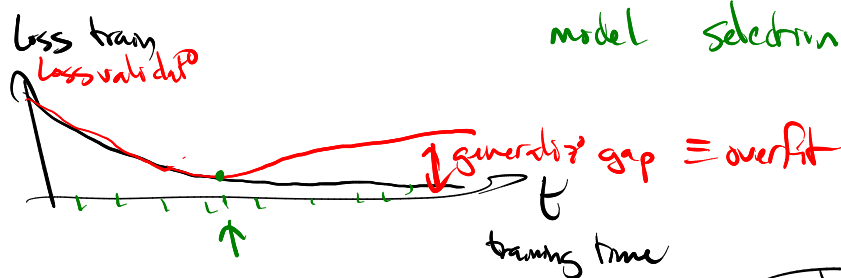
select randomly 50% of activities

ensemble methods
= robustify

activities
= outputs
of layer l

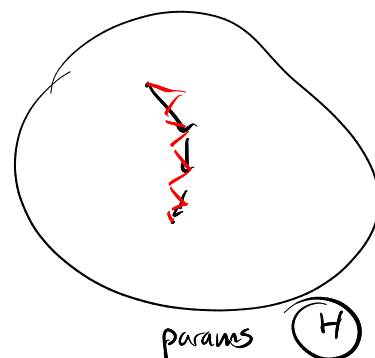
at test times: activities / 2

Early stopping

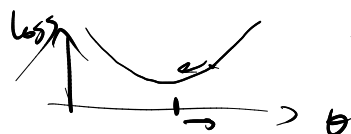


Optimize noise acts as a regularizer

source of noise: SGD : train on minibatches



[Passio] around convergence $\rightarrow$ local minimum



Loss $\propto$ Hessian matrix

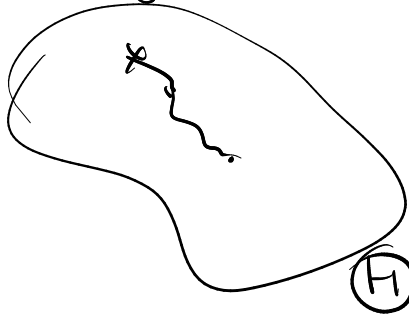
$\frac{d^2 \text{Loss}}{d\theta^2} \begin{cases} \text{if } > 0 : \text{then no pb} \\ \text{if not:} \end{cases}$

→ make the H > 0

↳ add a weight regularizer:

$$\lambda \sum \theta_i^2 \leftarrow \text{weight decay} = L^2 \text{ norm on parameters}$$

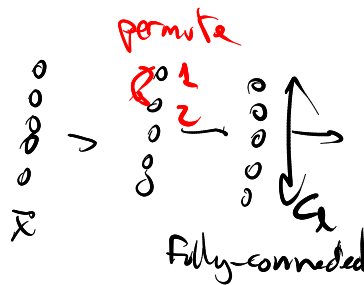
$$\theta_i \leftarrow \eta(-\nabla L - 2\lambda \theta_i)$$



≡ prior that favors smooth functions

Optimizable Landscape

local minima :



$$\underbrace{C_1! C_2! C_3! \dots C_L!}_{\text{huge}}$$

- So many parameters: → strong quality of local minima

1 Fixed ML task

= add more & more parameters

→ the more params, the easier

↳ $\nabla L = 0$ means none of $\frac{\partial L}{\partial \theta}$ improves the loss
billions of directions

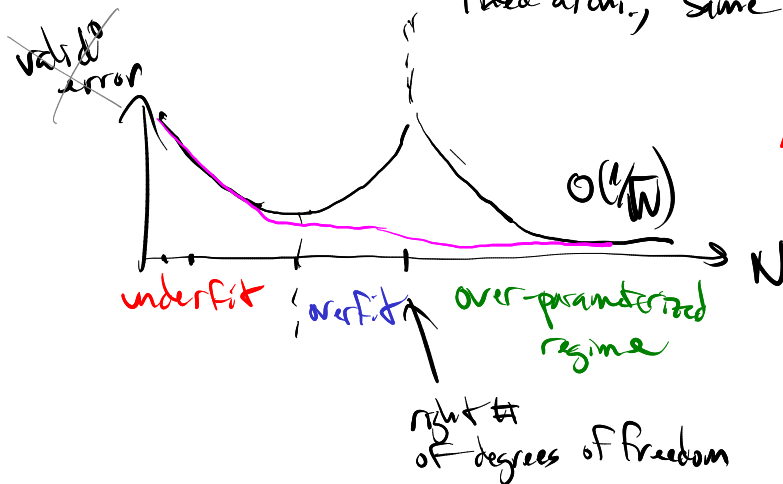
Over-parameterization helps

[Geiger et al, 2018]

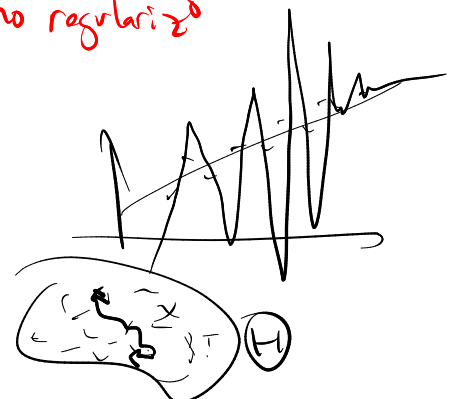
- Neural Tangent Kernel
- "double V descent"



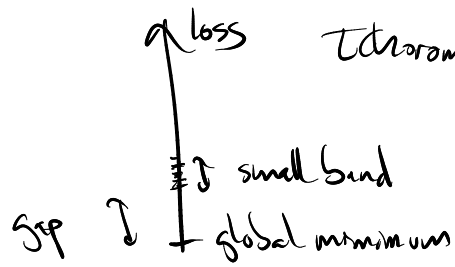
Fixed arch., same width N for all layers



no regularization



Quality of local minima



[Choromanska, 2015]

spin glasses

About MDZ

→ lost?

→ redundancies in parameters

[Welling] action recog. videos real-time on smartphone



→ 100 with ~ accuracy

→ tensoriz²

→ discr.

→ ~~distillate~~

→ pruning

) compression

