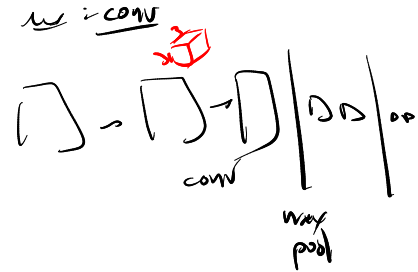
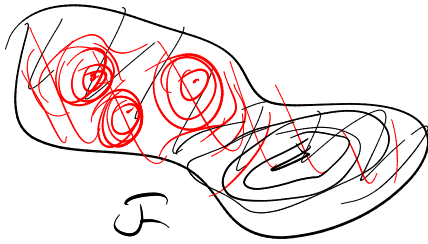


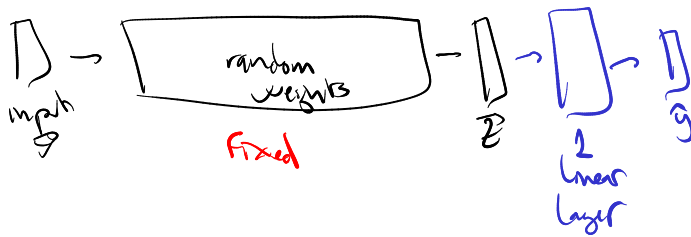
Architectures

I) Architectures as priors on function space

not just about expressivity, but probability
if random weights



"Extreme ML":

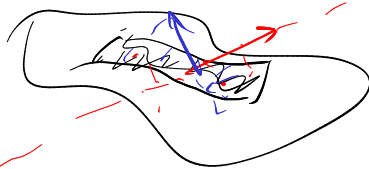


better accuracy (or $\approx$)
than with training
architecture bias

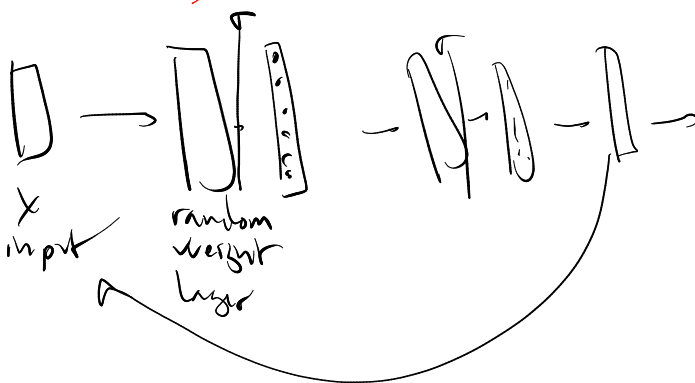
Random projections

high-dim space (embedding)

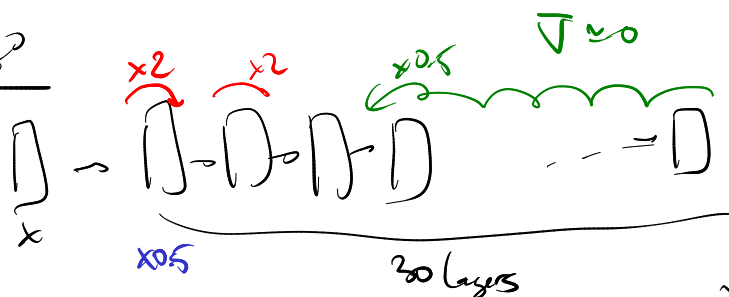
but intrinsic dim: not that much



Ch: keep all informat^o
if #projections $\geq \text{rank}(\text{manifold})$



Introspect



$\times 2^{30}$
 $\times 2^{30}$

$w_i \sim \mathcal{L}$

eg: Xavier Glorot init
Kaiming He



ReLU $\Rightarrow 1/2$ variance $\Rightarrow \times \sqrt{2}$

ex: linear layer: weights $w_i \sim \bar{w}=0, \sigma_w = \frac{1}{\sqrt{d}}$
inputs: x_i



$$x = \begin{pmatrix} x_1 \\ \vdots \\ x_d \end{pmatrix}$$

$$\bar{w}$$

$$\bar{w} \cdot x = \sum_i w_i x_i$$

$$\downarrow \mathbb{E}[\dots] \Rightarrow$$

$$\begin{aligned} \text{var: } \mathbb{E}_x \mathbb{E}_w \left(\sum_i w_i x_i \right)^2 &= \sum_i \left[\mathbb{E}_x x_i^2 \right] \left[\mathbb{E}_w w_i^2 \right] \\ &= d \times 1 \times \boxed{\frac{1}{d}} \\ &= 1 \end{aligned}$$

Prick architectures easy to train

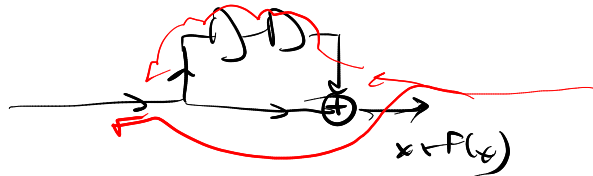
VGG

vs ResNets



$$F \sim \mathcal{F}$$

$$g \sim \text{Id} + \mathcal{F}$$



$$\text{output} = \sum_{\text{blocks}} F_{\text{block}}(x)$$

$$\frac{\partial L}{\partial \text{output}}$$

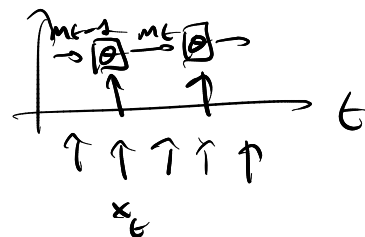
Zoo

- Fully-connected layers ("MLP")
multi-layer perceptron

- conv nets: 1D, 2D, 3D ... geometry = structured data $\Rightarrow$

- recurrent networks: LSTM, GRU

$\vec{x}$ = temporal signal
1 input



- attention block

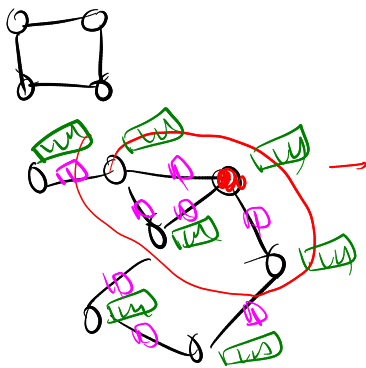
- graph NN

Graph-NW

1 input = 1 graph

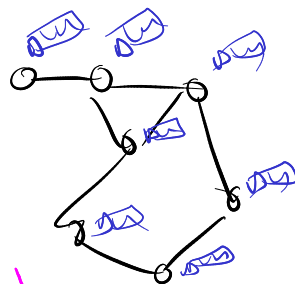


$$\sum_{p \in \mathcal{V}} w_p x_p$$



$$w_c x_p + \sum_{\substack{n \in \mathcal{V} \\ \text{neighbors}}} w_n x_n$$

average = $\frac{1}{N}$



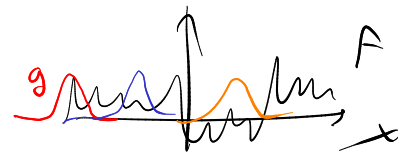
GAT
graph Attention Network
attention within neighborhood

$$w_n(e_{n,p})$$

$A \cdot A = A$
MVP

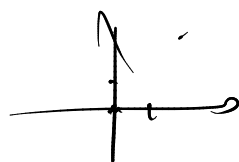
conv:

$$F * g(x) = \int F(x') g(x-x') dx'$$



Fourier

$$\widehat{F * g} = \widehat{F} * \widehat{g}$$



image

$$\sin(\omega x)$$



Graph Laplacian

eigenmodes

coordinate system

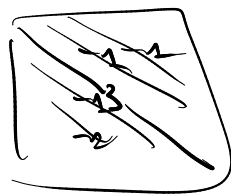
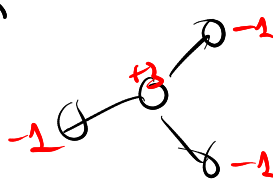
$$\text{each node} \leftrightarrow (i) \uparrow h$$

(if h eigenmodes)

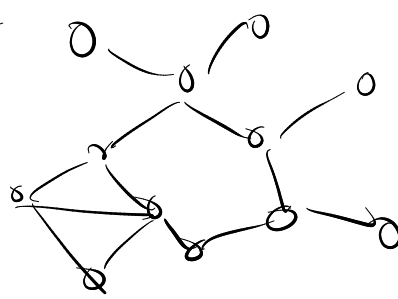
good to describe shapes (clouds of points)



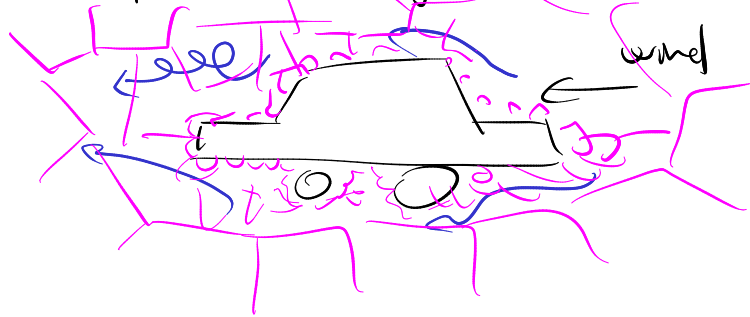
3D laser points on a surface



Other graphs: social networks



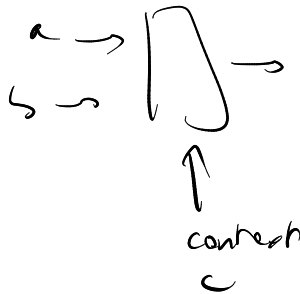
- mesh: computational fluid dynamics



dealing with resolution: graph pooling / unpooling
graph U-net
etc
ResNet structure ~

Transformers

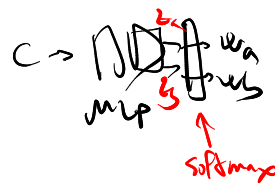
attention mechanism:



transformation

$$w_a a + w_b b$$

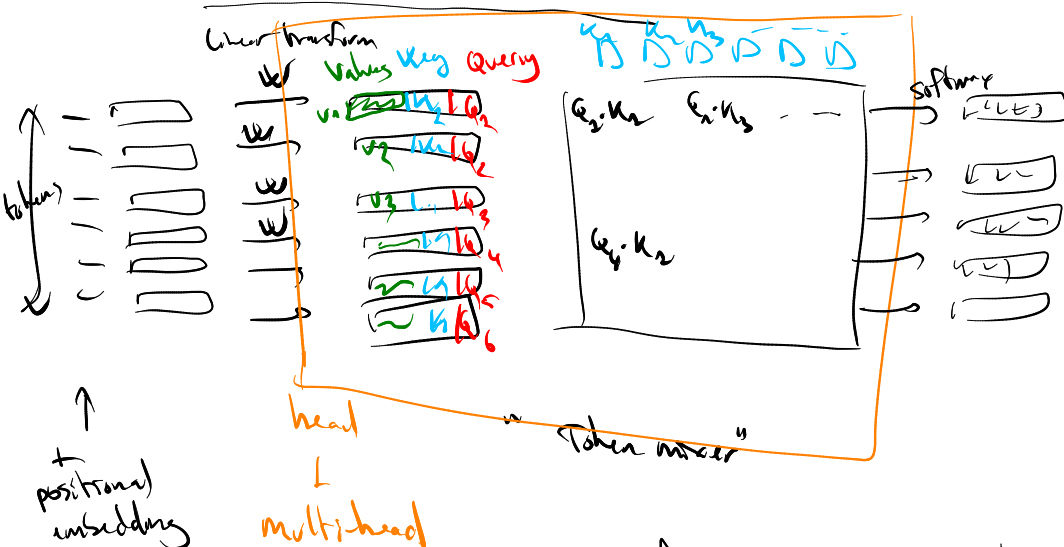
$$\begin{matrix} 1 & 0 \\ 0 & 1 \end{matrix}$$



$$w_a(c) a + w_b(c) b$$

$\sum = 1$

"Transformer" - self-attention



V

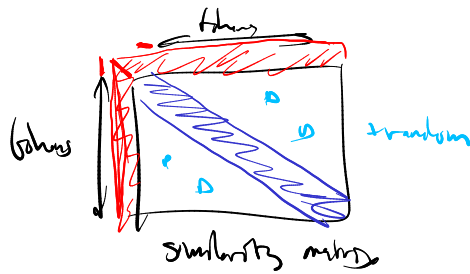
$$\begin{aligned} & \rightarrow \sum w_i^1 V_i \\ & \rightarrow \sum w_i^2 V_i \\ & \rightarrow \sum w_i^3 V_i \end{aligned} \quad \begin{matrix} \text{mip} \\ \text{mip} \\ \text{mip} \end{matrix}$$

properties:
- perm - equivariant

quadratic cost in #tokens

Variations

- Big Bird



2020

- kernelize:

$$\text{softmax}(Q^T K) V$$

$\approx \text{sim}$

$$k(\text{token}_i, \text{token}_j) = \phi(\text{token}_i) \cdot \phi(\text{token}_j) \quad \text{kernel trick} \\ \Rightarrow \text{linear!}$$

- reformer: hashes

→ discretize K & V into bins

& compute similarity only within the bin

- Linformer: low rank approx:

- self-transformer:



heads

L L L L

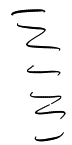
"substitutes over tokens"

tokens

L L L L L L

L L L

→



Neural Scaling Laws

- dataset size
- network size
- compute power

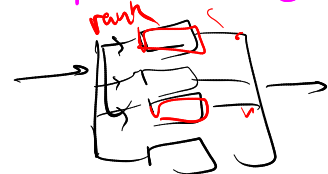
S^4 models → Mamba

HPC High Performing Computing

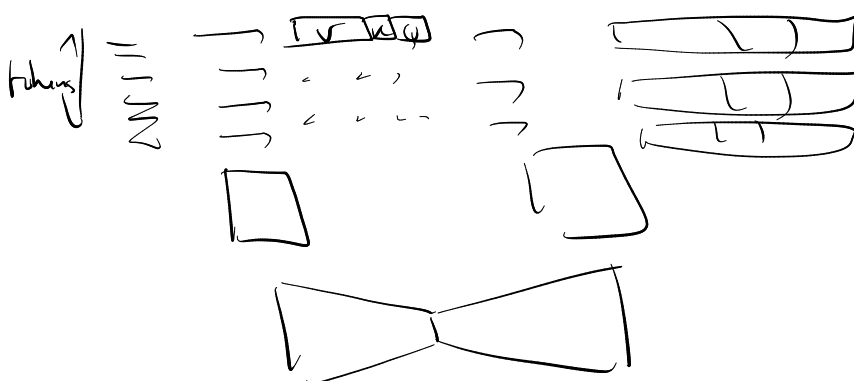
- Mistral & DeepSeek



"Mixture of Experts" $\equiv$ many MLPs in //



- DeepSeek:



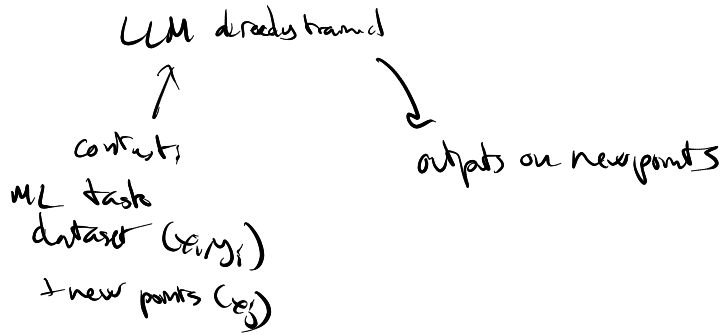
low precision

FP8

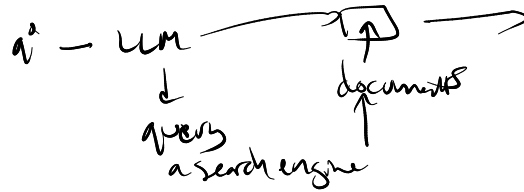
FP32

On top of LLMs

- prompting: alters significantly the output
- "In-Context Learning":



- RAG Retrieval-Augmented Gen.



- Agentic AI

Computer Vision Transformers ViT

convert

ViT $\rightarrow$ 10x

ResNets $\rightarrow$ ViT

14x	14x	14x
12x	12x	12x
8x	8x	8x
3x	3x	3x
14x	14x	14x

recipe to train $\neq$