

Learning

PAC Learning

Classification

Vc Dim Vc Thm

Time complexity

Regression

Linear OLS

Regularization

Probabilistic Classification

Logistic regression

NN

architec

SGD

Generation

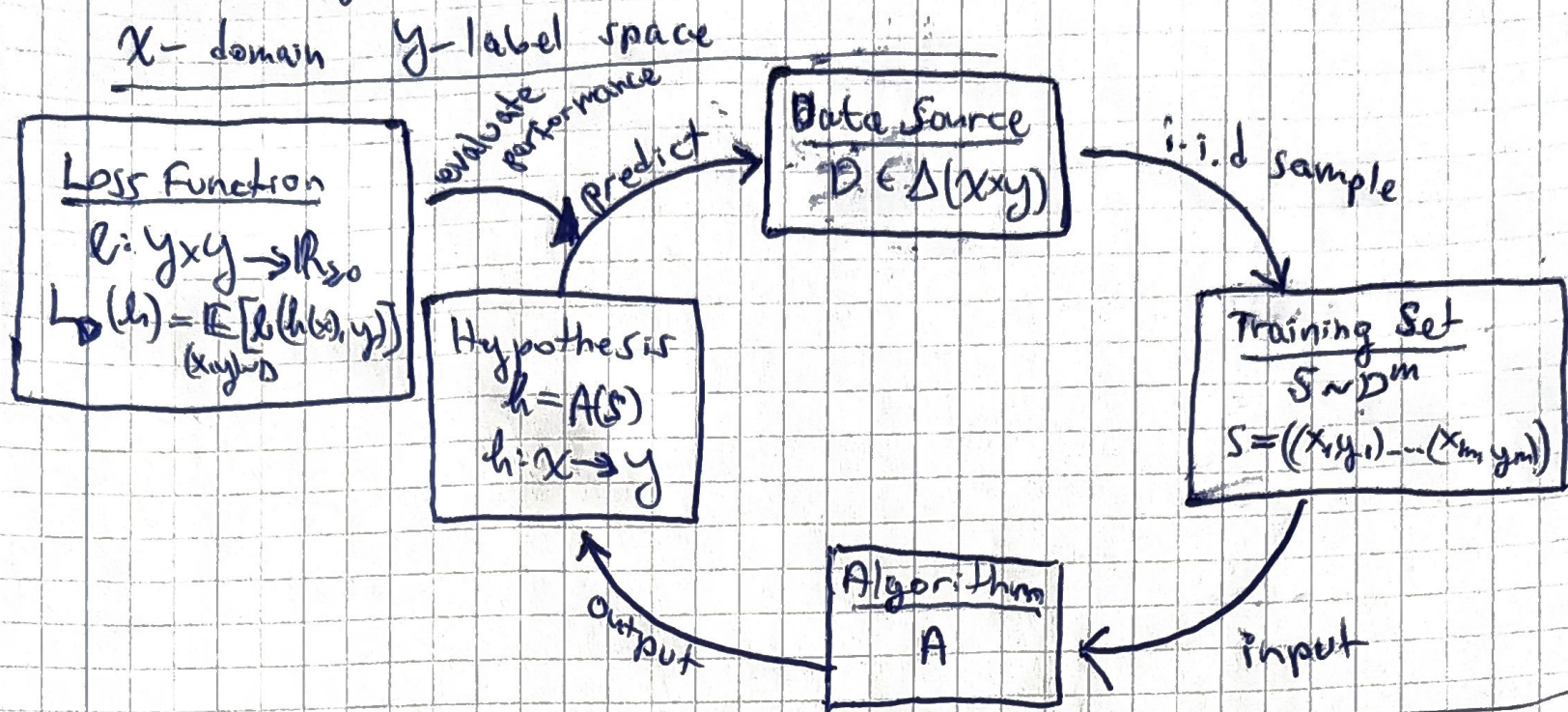
Vibe definition, indistinguishability

Next token prediction

RLHF

Diffusion

1. Learning



Instantiations

- (Binary) classification with '0-1' loss
- Regression with L_2 loss
- Probabilistic Classification with CE loss

Example

Will treatment cure patient?
How big will baby be?

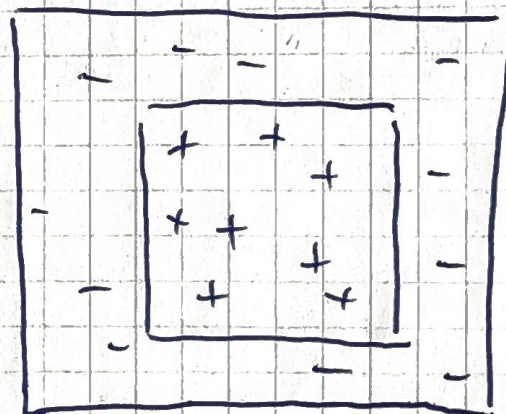
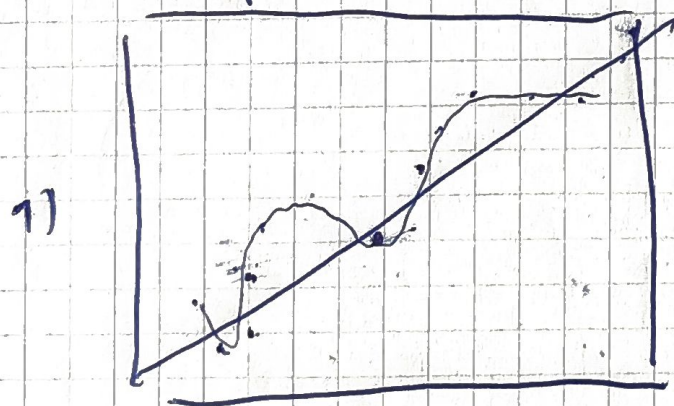
Want: low loss

$$L_S(h) = \frac{1}{m} \sum_{i=1}^m l(y_i, h(x_i))$$

Fitting to data - L_S small

overfitting - fitting too well to S can be bad

Examples:



VS memorization

Resources:

- data ("sample complexity")
- compute $\begin{cases} \text{train time} \\ \text{inference time} \end{cases}$
- memory
- randomness

1.1 PAC Learning

What is good enough?

In practice — based on use case (gambling, stock market vs. medical decision)

Ideally:- Learning = "require $L_D(h) = 0$ " (luck doesn't always exist)

Simple case where 0-loss hypothesis exists:

Def (Realizable dist). D is H -realizable if $\exists h \in H: L_D(h) = 0$

Q: Can we get 0 loss in realizable case?

A:

Limitation	Solution in Definition	Name
Limited training set	$L_D(h) \leq \epsilon$ "success"	"Approx Correct"
"Bad Luck" Bad samples	$P[\text{success}] \geq 1 - \delta$	"Probably"
Other resources time, memory...	ignore	

Def (PAC Learning - Realizable Case). Let $X, Y, \ell: Y^2 \rightarrow \mathbb{R}_{\geq 0}, H \subseteq Y^X$.

We say that $A: (X \times Y)^* \rightarrow Y^X$ PAC Learns H if:

$\forall \epsilon > 0 \forall \delta > 0 \exists m = m(\epsilon, \delta) \in \mathbb{N} \forall H\text{-realizable } D \in \Delta(X \times Y):$

$$P_{S \sim D^m} [L_D(A(S)) \leq \epsilon] \geq 1 - \delta.$$

Q: If 0-loss h doesn't exist? Don't want to assume anything about D ?

A: compare to benchmark / competitive with H

Def (Agnostic PAC Learning). Let X, Y, ℓ, H .

We say that $A: (X \times Y)^* \rightarrow Y^X$ agnostically PAC Learns H if

$\forall \epsilon > 0 \forall \delta > 0 \exists m = m(\epsilon, \delta) \in \mathbb{N} \forall D \in \Delta(X \times Y):$

$$P_{S \sim D^m} [L_D(A(S)) \leq \underbrace{\inf_{h \in H} L_D(h)}_{\text{OPT}} + \epsilon] \geq 1 - \delta.$$

Lemma. Let H finite, $L: Y^Z \rightarrow [0,1]$ (bounded). Then H is:

(1) PAC Learnable with samp comp $m(\epsilon, \delta) \leq \frac{\ln(H/\delta)}{\epsilon}$.

(2) Agnostic PAC Learnable — " — $m(\epsilon, \delta) \leq O\left(\frac{\ln(H/\delta)}{\epsilon^2}\right)$

ERM ("Empirical Risk Minimization")

choose arbitrary $h \in \arg\min_{h' \in H} L_S(h')$

return h

PF of Lemma.

[Def. Fix $h \in H$, $D \in \mathcal{M}(X \times Y)$. Let $S \in (X \times Y)^m$. We say that S is bad for h if $L_S(h) = 0$ and $L_D(h) > \epsilon$ ("misleading" S)

$$(1) \quad \mathbb{P}_{S \sim D^m} [L_D(\text{ERM}(S)) > \epsilon] \leq \mathbb{P}_{S \sim D^m} [\exists h \in H: S \text{ is bad for } h] \\ \leq \mathbb{P}_S \left[\bigcup_{h \in H} \{S \text{ is bad for } h\} \right] \leq \sum_{h \in H} \mathbb{P}_S [S \text{ is bad for } h]$$

$$\leq \sum_{h \in H} (1-\epsilon)^m \leq |H| e^{-\epsilon m} \stackrel{H \times \leq e^x}{\leq} \delta \quad \uparrow \quad m \geq m(\epsilon, \delta)$$

$$\mathbb{P}_{(x,y) \sim D} [L(h(x), y) > \epsilon] \geq \epsilon$$

(2) is similar

Thm. (Vapnik) For PAC Learning, optimal samp comp is

$$m(\epsilon, \delta) = O\left(\frac{VC(H) + \log(1/\delta)}{\epsilon}\right)$$

(ERM is optimal up to $\log(1/\epsilon)$ factor). optimal alg given by Hanneke (2015).

For agnostic PAC Learning; ERM is optimal with samp comp

$$m(\epsilon, \delta) = O\left(\frac{VC(H) + \log(1/\delta)}{\epsilon^2}\right)$$

2. Classification

Q: what is the best we can do? (samp comp)

A: VCDim (next page)

Def (Restriction). Let $\mathcal{H} \subseteq \{0,1\}^X$, $A \subseteq X$, $h \in \mathcal{H}$.

- $h|_A = g$ s.t. $g: A \rightarrow \{0,1\}$ and $\forall a \in A: g(a) = h(a)$
- $\mathcal{H}|_A = \{h|_A: h \in \mathcal{H}\}$

Def (Shattering). Let $\mathcal{H} \subseteq \{0,1\}^X$, $A \subseteq X$, $|A| < \infty$.

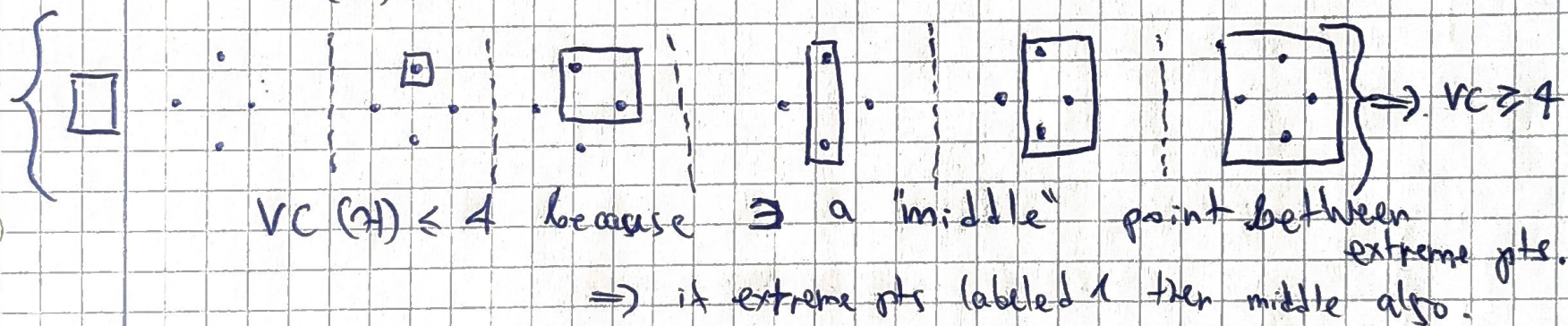
We say that $\mathcal{H}$ shatters A if $|\mathcal{H}|_A| = 2^{|A|}$

Def (VC Dimension). Let $\mathcal{H} \subseteq \{0,1\}^X$.

$$VC(\mathcal{H}) = \sup \{k \in \mathbb{N}: \exists A \subseteq X \quad |A| = k \wedge \mathcal{H} \text{ shatters } A\}.$$

Examples

- $|\mathcal{H}| \leq 1: VC(\mathcal{H}) = 0$
- $|\mathcal{H}| = 2: VC(\mathcal{H}) = 1$
- Thresholds: $\mathcal{H} = \{\mathbb{1}(x \geq t): t \in \mathbb{R}\}$. $VC(\mathcal{H}) = 1$, But $|\mathcal{H}| = \infty$.
- Rectangles: $h_{x_{\max}, x_{\min}, y_{\max}, y_{\min}}(x, y) = (\mathbb{1}_{x_{\min} \leq x \leq x_{\max}})(\mathbb{1}_{y_{\min} \leq y \leq y_{\max}})$
 $\mathcal{H} = \{h_{x,y}: (x,y) \in \mathbb{R}^2\}$
 $VC(\mathcal{H}) = 4$



Time Complexity

- ERM can be computationally hard, E.g.:

Thm (Ben-David, Elron, Long 2003). Let $X = \mathbb{R}^n$ and $\mathcal{H}$ be the class of axis-aligned rectangles in $\mathbb{R}^n$. If $P \neq NP$ then

$\forall \epsilon > 0$ $\nexists$ poly(n)-time algorithm that given input

$$S \in (X \times \{0,1\})^*, \text{ outputs } h \in \mathcal{H} \text{ s.t. } L_S(h) \leq \left(\frac{770}{767} - \epsilon\right) \min_{h \in \mathcal{H}} L_S(h).$$

- ERM can be undecidable. $h_{\text{halt}}(x) = \mathbb{1}(x \text{ is a TM that halts})$,

$$h_{\overline{\text{halt}}}(x) = 1 - h_{\text{halt}}(x). \quad \mathcal{H} = \{h_{\text{halt}}, h_{\overline{\text{halt}}}\}. \quad S = \{(x, 1)\}.$$

Def (LPN Search Problem). $X = \mathbb{F}_2^n$, $f_s(x) = \langle s, x \rangle \in \mathbb{F}_2$ ⑤
 $f_s: X \rightarrow \mathbb{F}_2$
 $\mathcal{H} = \{f_s: s \in \mathbb{F}_2^n\}$.

$\mathcal{D}_s \in \Delta(X \times \mathbb{F}_2)$ s.t.:

(1) $x \sim U(X)$; (2) $e \sim \text{Ber}(\frac{1}{q})$; (3) $y = f_s(x) \oplus e$
 output: (x, y)

~~Hardness Assumption~~

Problem: given $((x^{(1)}, y^{(1)}), \dots, (x^{(m)}, y^{(m)})) \sim (\mathcal{D}_s)^m$, find s .

Hardness assumption: LPN is hard. I.e. $\forall$ poly(n)-time

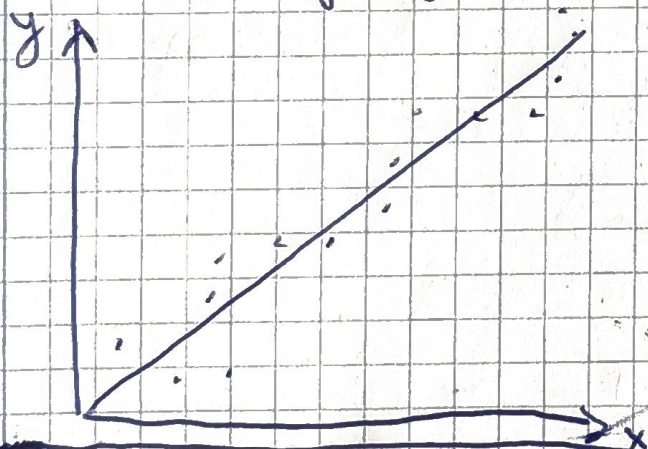
algorithm A : $\mathbb{P}_{s \sim U(X)} [A(X, X \cdot s \oplus e) = s] \leq \text{negl}(n)$.

$X \sim U(\mathbb{F}_2^{n \times n})$
 $e \sim (\text{Ber}(\frac{1}{q}))^m$

4. Regression

• Label space $\mathcal{Y}$ is ~~continuous~~ continuous, no longer hope/expect to predict perfectly, $\hat{y} = y$, enough to be close, $\hat{y} \approx y$.

• Concretely: $\mathcal{Y} = \mathbb{R}$, $l(y, \hat{y}) = (y - \hat{y})^2$ (l_2 loss)



$X = \mathbb{R}^n$, $\mathcal{H} = \{h_w: w \in \mathbb{R}^n\}$

$h_w(x) = \langle w, x \rangle$, $h_w: \mathbb{R}^n \rightarrow \mathbb{R}$

=ERM

OLS(S): ("Ordinary Least Squares")
 $S \in (\mathbb{R}^n \times \mathbb{R})^m$, $S = ((x^{(1)}, y^{(1)}), \dots, (x^{(m)}, y^{(m)}))$
 return $h_w \in \arg\min_w \sum_{i=1}^m (h_w(x^{(i)}) - y^{(i)})^2$

$f(w) = L_S(h_w)$ is convex (sum of squares of linear functions) in w , so solving

$$\nabla f(w) = 0$$

gives a closed-form solution for OLS.

4.1 Regularization

Often we get better generalization if we have inductive bias towards "simpler" functions using regularization:

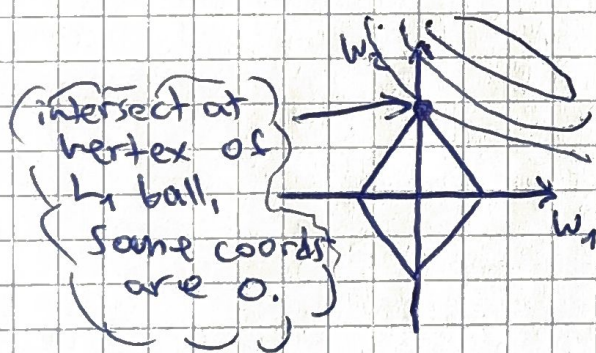
Select h_w that minimizes:

- L_2 / Ridge: $L_S(h_w) + \lambda \|w\|_2^2 = \frac{1}{n} \sum_i (\langle w, x_i \rangle - y_i)^2 + \lambda \|w\|_2^2$
- L_1 / Lasso: $L_S(h_w) + \lambda \|w\|_1 = \text{---} + \lambda \|w\|_1$

Intuition:

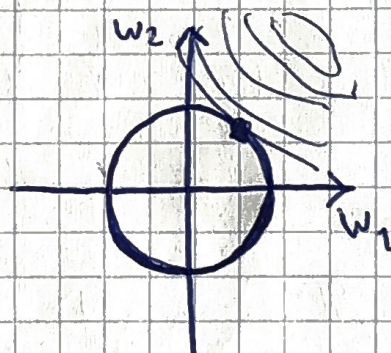
Regularization
limits fit to noise
($L_1 + L_2$)

- To fit outliers with large $|y_i|$, we need $\|w\|$ large. But if outlier is noise / not representative, then this large $\|w\|$ will give large population loss.
- For $x \in \mathbb{R}^n$, if there are some directions (covariates) x_i that are not informative (indep. of y_i), OLS fits w_i to noise in direction i .
- L_1 induces sparse solutions. $L_S(h_w)$ is quadratic function of w . Contours are ellipsoids. ★



L_1

sparse solutions



L_2

Non-sparse solution

5. Probabilistic Classification

$y = \{0, 1\}$ but $A \in [0, 1]^X$.

$\hat{p} = h(x)$ = "probability that label is 1"

Ex: Predict whether medication will cure patient

(We could replace $h(x)$ with $\text{eg. } \mathbb{1}(h(x) \geq \frac{1}{2})$ to make binary predictions, but $\hat{p} \in [0, 1]$ useful in many cases)

- conceptually, represents uncertainty
- good for optimization

★ another intuition:

Suppose x_1, x_2 have correlation with each other,
 $(x_1 \approx x_2)$, and with y (x_1 a bit better)

~~e.g. $w_1 \approx w_2$~~

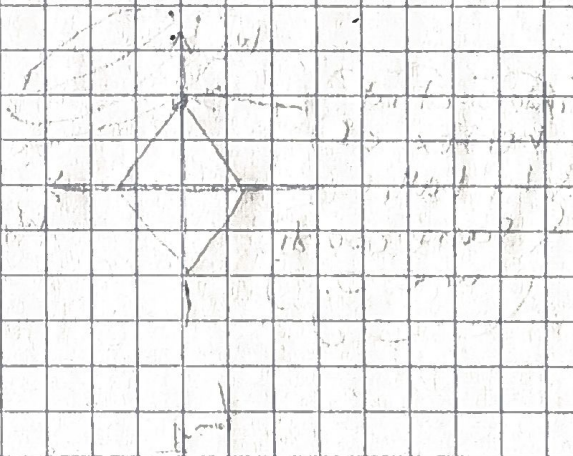
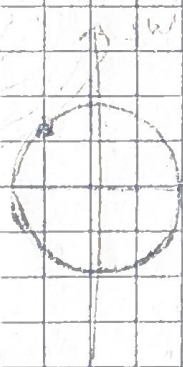
e.g. L_1 is minimized by any w with
 $w_1 + w_2 = 1$

then L_2 regularization would give $w_1 \approx \frac{1}{2}$

L_1

$w_1 = 1, w_2 = 0$

Because L_2 penalizes large abs value for each w_i
 L_1 doesn't, only sum of w_i

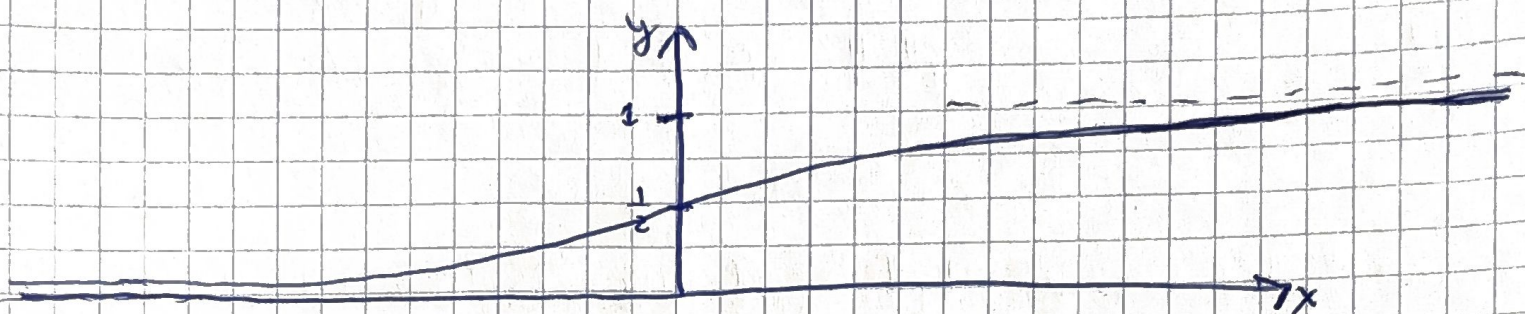


Loss function: Cross Entropy (CE) / Log Loss / Surprisal: (7)

$$l(y, \hat{p}) = \begin{cases} \ln \frac{1}{\hat{p}} & y=1 \\ \ln \left(\frac{1}{1-\hat{p}} \right) & y=0 \end{cases} = y \ln \frac{1}{\hat{p}} + (1-y) \ln \frac{1}{1-\hat{p}} = \mathbb{E}_y \left[\ln \frac{1}{\hat{p}_y} \right]$$

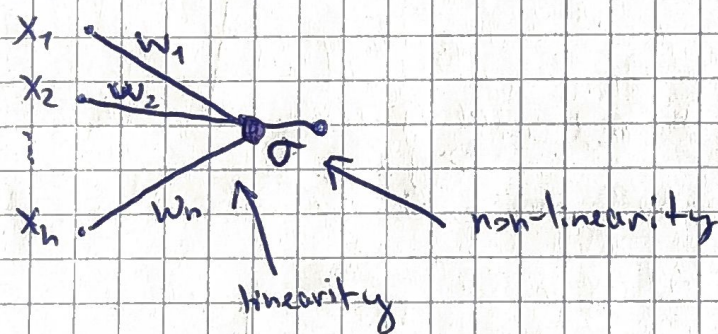
5.1 Logistic Regression

Logistic function: $\sigma(x) = \frac{1}{1+e^{-x}} = \frac{e^x}{e^x+1}$



Logistic Regression: a "generalized linear model":

- $X = \mathbb{R}^n$, $h_w = \sigma(\langle w, x \rangle + b)$, $\mathcal{H} = \{h_w : w \in \mathbb{R}^n\}$
- One of the most common models in science
- Is a single-neuron NN:



- no closed form solution

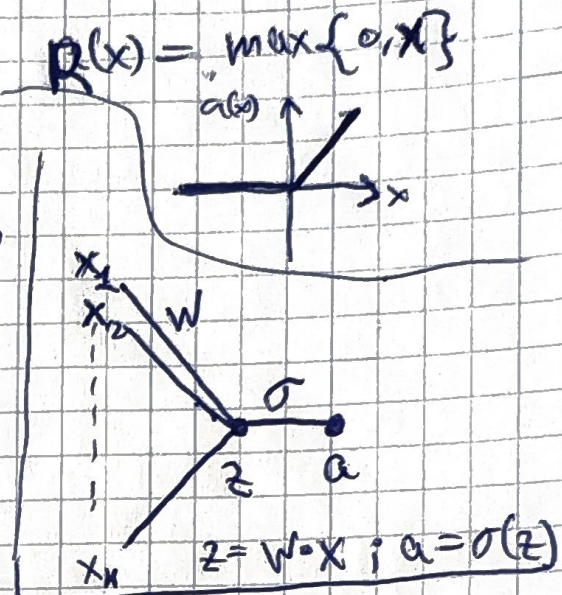
6. Neural Networks

Neuron: composition of linear function with non-linear "activation" function, e.g.:

- Rectified Linear Unit (ReLU): $R(x) = \max\{0, x\}$
- Logistic: $\sigma(x)$
- Softmax (generalized logistic,

maps $x_1, \dots, x_k \mapsto \Delta(L^k)$

$$\sigma(x_1, \dots, x_k) = \left(\frac{e^{x_i}}{\sum_{j=1}^k e^{x_j}} \right)_{i=1}^k$$



Neural Network: composition of many neurons

Architectures:

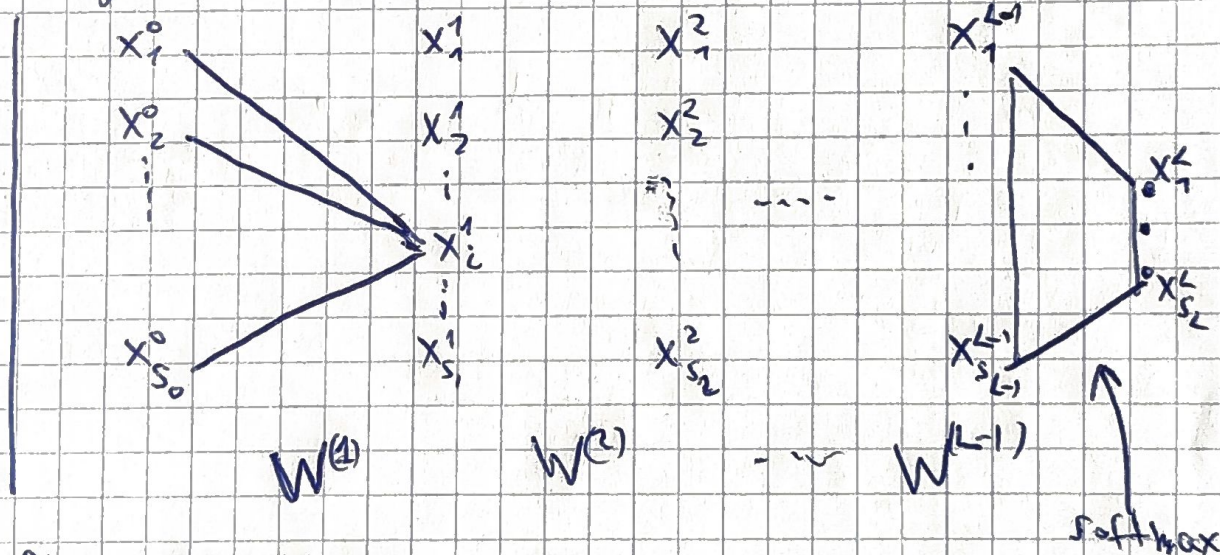
Forward Pass

- Fully-connected:

For $i = 1, 2, \dots, L$:

$$z_i^{(i)} \leftarrow W^{(i)} \cdot x^{(i-1)}$$

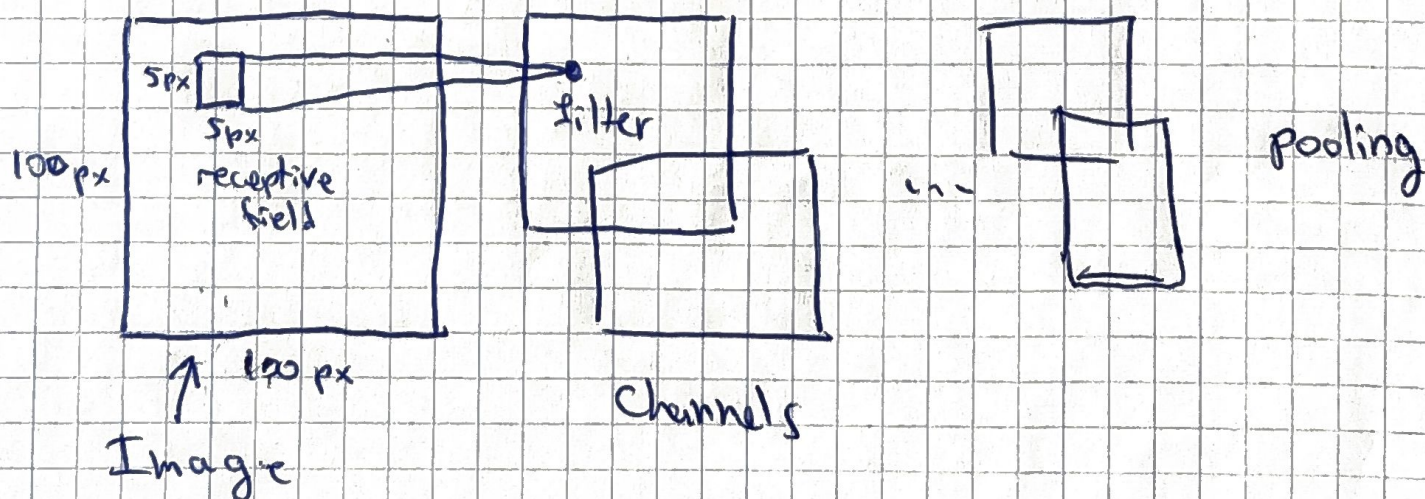
$$x^{(i)} \leftarrow \sigma_i(z_i^{(i)})$$



- Skip connections

- Recurrent

- Convolutional (CNN). Common in vision:



- Locality. Filters for local features like Gabor lines
- Translation ~~invariance~~ equivariance

(Minibatch) Stochastic Gradient Descent (SGD)

Initialize $W^{(0)}$ randomly
for $t = 1, 2, \dots, T$:

Select $v^{(t)} \in \arg\min_{v \in W^{(t-1)}} \ell(v, z)$

$$w^{(k)} \leftarrow w^{(k-1)} - \eta \cdot v^{(k)}$$

Sub gradient
w.r.t w



A dynamic-programming algorithm for computing gradients.

~~$$L = L(\sigma(w^{(1)}) \sigma(w^{(2)}) \dots \sigma(w^{(n)}), y)$$

~~By chain rule:~~

$$\frac{\partial L}{\partial w^{(k)}} = \frac{\partial L}{\partial a^{(k)}} \cdot \frac{\partial a^{(k)}}{\partial z^{(k)}} \cdot \frac{\partial z^{(k)}}{\partial a^{(k-1)}} \cdot \frac{\partial a^{(k-1)}}{\partial z^{(k-1)}} \dots \frac{\partial a^{(2)}}{\partial z^{(2)}} \cdot \frac{\partial z^{(2)}}{\partial w^{(k)}}$$~~

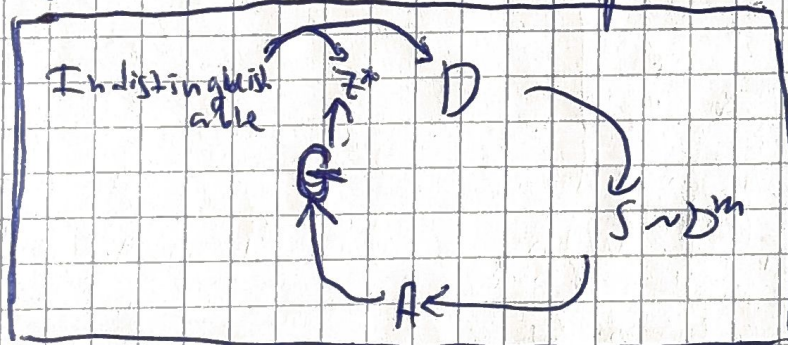
Thm (Backpropagation / auto-differentiation — informal)

Suppose $f: \mathbb{R}^d \rightarrow \mathbb{R}$ is computed by a differential circuit of size N . Then ∇f can be computed in time $O(N)$.

Idea:- compute forward pass $x^0, z^1, x^2, z^2, \dots, x^L, z^L$, loss $\ell = \ell(x^L, y)$ - Let $\delta^L = \frac{\partial \ell}{\partial z^L}$ - for $k = L-1, L-2, \dots, 1$:

$$\begin{aligned} \delta^k &:= \frac{\partial \ell}{\partial z^k} = \frac{\partial \ell}{\partial z^{k+1}} \cdot \frac{\partial z^{k+1}}{\partial z^k} = \delta^{k+1} \cdot \frac{\partial z^{k+1}}{\partial x^k} \cdot \frac{\partial x^k}{\partial z^k} \\ &= \delta^{k+1} \cdot W^{k+1} \cdot \sigma'(z^k) \\ \frac{\partial \ell}{\partial W^k} &= \frac{\partial \ell}{\partial z^k} \cdot \frac{\partial z^k}{\partial W^k} = \delta^k \cdot x^{k+1} \end{aligned}$$

7. Generation

Example: this person does not exist.comGeneral Idea: I show my learning alg. iid samples $z_1, \dots, z_m \sim \mathcal{D}^m$, and how it can generate new z^* that "look as if" they came from $\mathcal{D}$ Ex: After Seeing 100 Shakespeare sonnets, my system can generate new "Shakespeare sonnets"What is "indistinguishable"?Unclear, based on usecase
Not TV etc.

7.1 LLMs

can be ~~used~~ "prompted" w/ prefixes,
not directly conversational