



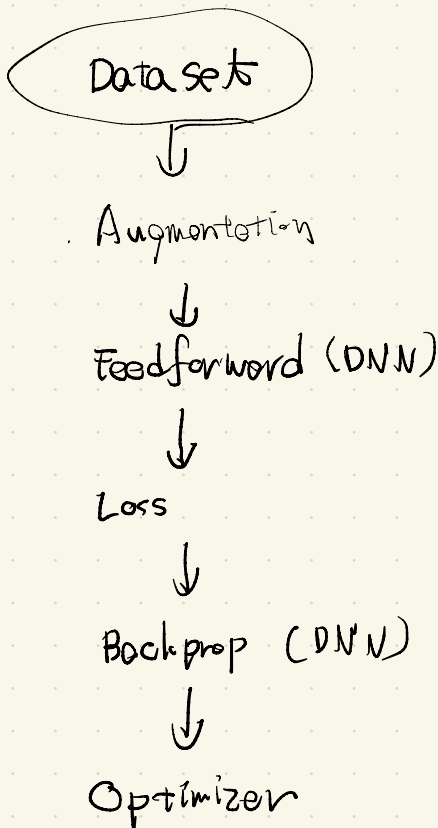
32ゲーム行列と深層学習

3/23, 24

Day 1

深層学習速習

- ✓ framework of deep learning
- ✓ examples.



Framework

LOS

$$\mathcal{D} \subset X \times Y \subset \mathbb{R}^n \times \mathbb{R}^{n_0}$$

finite

μ = prob. dist. on $X \times Y$

$$L: \mathbb{R}^{n_0} \times \mathbb{R}^{n_0} \rightarrow \mathbb{R}$$

$$\text{Find } f: \mathbb{R}^n \rightarrow \mathbb{R}^{n_0}$$

$$\text{s.t. minimize } \int_{(x,y) \in X \times Y} L(f(x), y) \mu(dx dy)$$

unknown

with information $\mathcal{D}$.

★ $\mathcal{D} \subset X \times Y \subset \mu \in \mathcal{D} \subset \text{approx}$

★ $f = f_\theta \quad (\theta \in \Theta \subset \mathbb{R}^p)$

parametrize

★ $\frac{1}{n} \in \text{approx}$

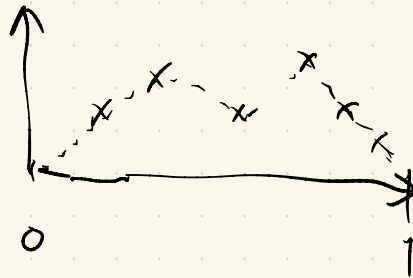
★ "minimize" $\in \text{approx}$

Dataset D

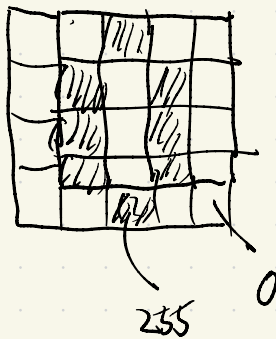
CV

1-dim

$$X, Y = [0, 1] \subset \mathbb{R}$$



Computer Vision



Pixel

RGB

$H \times W \times C$

↑
3

→ vectorize

HWC -dim
vector

output: $0 \rightsquigarrow \begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix} = e_1 \in \mathbb{R}^{10}$

$\nu \quad 0 \sim 9 \rightarrow e_1 \sim e_{10}$

NLP

MNIST: ~~手写字~~ 数字.

Natural Language Processing

(NLP)

I am eating pizza.

tokenize

$\rightarrow [0, 10, 201, 1018]$

vector

$\rightarrow x \in \mathcal{X}$.

y ?

mark

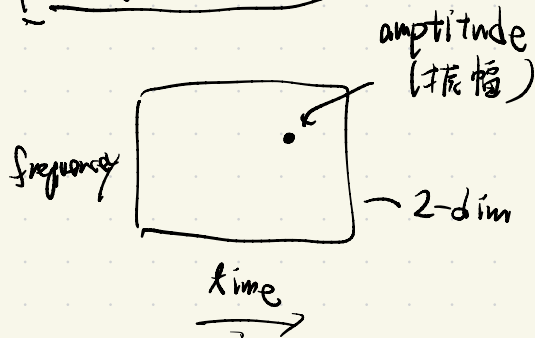
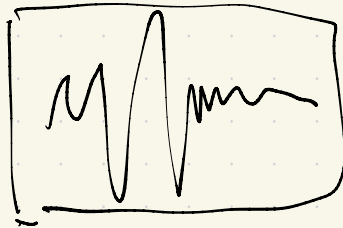


x : I am  pizza

y : I am eating pizza.

Audio

Audio

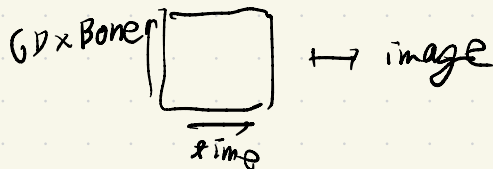
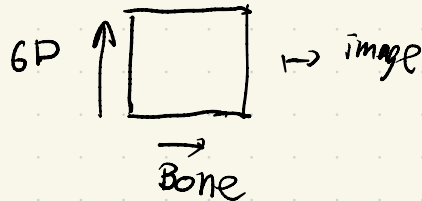


画像として扱う
→

3D Bone Data



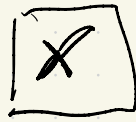
3D(+3D) x Bone



Augmentation

3D

Image



flip



Random
Flip



Random
Erasing

data に依存に増やす

Model

$$f: \mathbb{R}^n \rightarrow \mathbb{R}^{n_0}$$

$$f = f_{\theta}$$

$$f_{\theta}(x) = Wx + b$$

$$\theta = (W, b)$$

$$\begin{cases} W \in \mathbb{R}^{n_0 \times n} \\ b \in \mathbb{R}^{n_0} \end{cases}$$

Multi-layer-Perceptron

Aug 2

(MLP).

$$x_0 = x$$

$$\left\{ \begin{array}{l} h_l = w_l x_{l-1} + b_l \\ x_l = \varphi(h_l) \\ l = 1, \dots, L \end{array} \right.$$

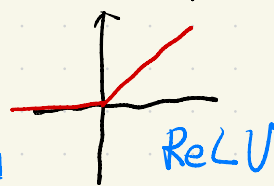
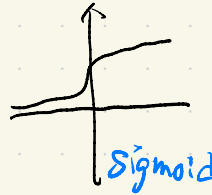
$L \in \mathbb{N}$

φ : activation ftn.

$$\varphi: \mathbb{R} \rightarrow \mathbb{R}$$

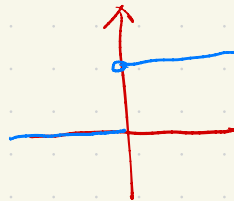
• conti.

• diff' except for finite pts.



$$f_{\Theta}(x) = h_L$$

$$\Theta = (w_1, b_1, \dots, w_L, b_L)$$



* L : Loss fn

FF

↳ L2-fn

$$\hookrightarrow L(f(x), y) = \|f(x) - y\|^2$$

↳ Softmax Cross Entropy --

↳ L1 - - -

Backpropagation

L_0

$$\left\{ \begin{array}{l} \text{minimize} \int \|f_\theta(x) - y\|^2 \\ \theta \in \Theta \mid (x, y) \in \mathcal{D} \end{array} \right\} + \lambda \|\theta\|^2$$

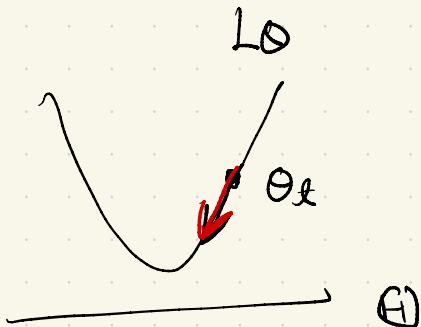
weight decay

Gradient Descent

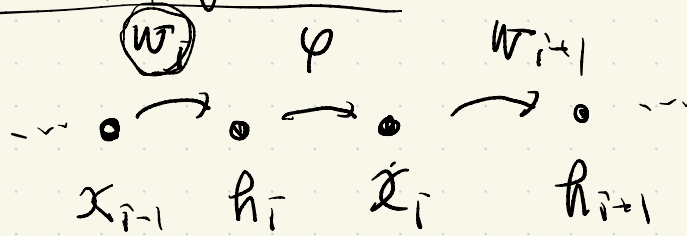
$$\theta_{t+1} = \theta_t - \gamma_t \nabla L \theta_t$$

$$\gamma_t > 0$$

: learning rate



Back propagation



$$\frac{\partial L_0}{\partial w_i} = \frac{\partial L_0}{\partial h_i} \cdot \frac{\partial h_i}{\partial w_i}$$

Diagram illustrating the backward pass (backpropagation) of the error gradient. The error gradient $\frac{\partial L_0}{\partial w_i}$ is calculated as the product of the error gradient $\frac{\partial L_0}{\partial h_i}$ and the partial derivative of the hidden state h_i with respect to the weight w_i .

$$\frac{\partial L_0}{\partial w_{e1}} \quad \frac{\partial L_0}{\partial w_L}$$

Diagram illustrating the backward pass (backpropagation) of the error gradient. The error gradient $\frac{\partial L_0}{\partial w_{e1}}$ is calculated as the product of the error gradient $\frac{\partial L_0}{\partial h_i}$ and the partial derivative of the hidden state h_i with respect to the weight w_{e1} .

Summary

✓ Dataset

✓ Augmentation

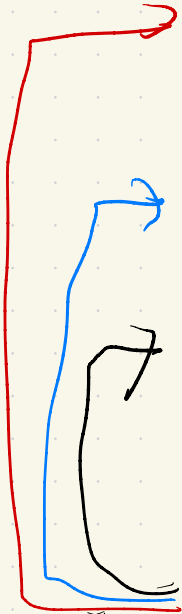
✓ NN

✓ Forward

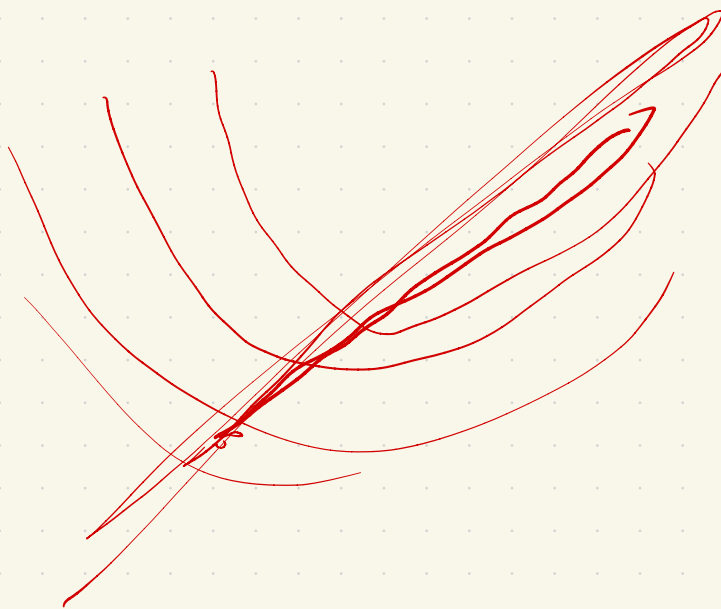
✓ Loss

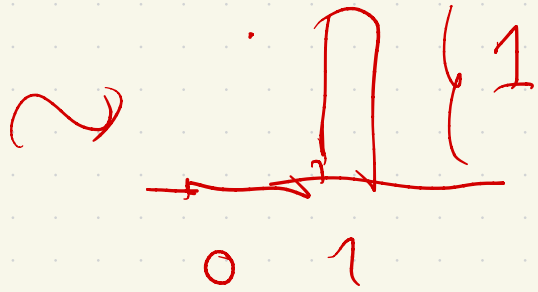
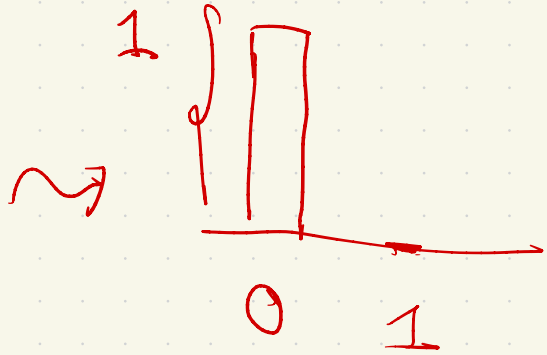
✓ Back propagation

optimizer



$$y = ABx$$

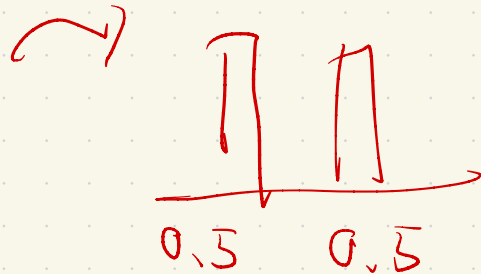




Mixup



$\alpha \approx 0.5$



Day 2 Signal Propagation

(1) ReLU $\hookrightarrow$ $\text{ReLU}'' \nabla L_0$

$$\Theta_{t+1} = \Theta_t - \eta_t \nabla_{\Theta} L_0$$

$$L_0 = \sum_{x,y} L(f_{\Theta}(x), y) + \boxed{\text{loss}}$$

$$\Theta_t = (w_1^{(t)}, \dots, w_L^{(t)}) \\ b_1^{(t)}, \dots, b_L^{(t)}$$

Initialization が 必要.

$$W_{\ell, \vec{i}j} \sim N(0, \frac{1}{n}) \quad \text{i.i.d.}$$

$$(\vec{i}j = 1, \dots, n)$$

$$(\ell = 1, \dots, L)$$

$$b_{\ell, \vec{i}j} \sim N(0, \frac{1}{n}) \quad \text{i.i.d.}$$

$$(\vec{i}j = 1, \dots, n)$$

$$(\ell = 1, \dots, L)$$

$$\Rightarrow \mathbb{E}[\|h_{\ell}\|^2] = \mathbb{E}[\|x_{\ell-1}\|^2] + 1$$



$x_i = \phi(h_i)$ 2倍ずつ増える

ランダム行列

: 行列値 (R-valued)

- 確率変数.

$$W_{ij} \sim N(0, \sigma^2) \quad (i, j = 1, \dots, n)$$

i.i.d

Ginibre

Random

Matrix.

$W \sim$ uniformly sample

from $O(n)$

Haar Orthogonal

Random Matrix

W : Gram matrix

$$W = U D V^T$$

$$U, V \in O(n)$$

Haar Orthogonal

$$D = \text{diag} \begin{pmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{pmatrix}$$

$\lambda_1 \geq \dots \geq \lambda_n$: singular value of W

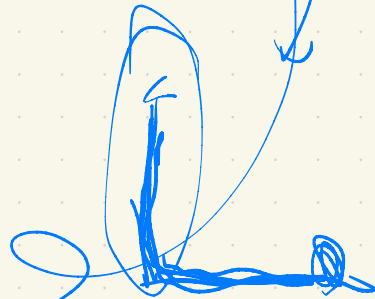
$(W^T W \text{ の固有値の square root})$

Back propagation

$$J = \frac{\partial h_L}{\partial x_0} = W_L D_{L-1} \underbrace{W_{L-1} D_{L-2}}_{\dots W_1}$$

J の singular value の値が小さくなる

→ 学習が進行する



Explosion / Vanishy gradient

→ J が orthogonal に近くなる

→ J の 特異値 が I 周辺

に 集中 している ような

状況 を つく いませんか?

↳ Dynamical Isotometry
(Saxe, Ponnigton)

$$J = W_L D_{L-1} W_{L-1} \dots W_1$$

$$D_e = \psi'(h_e)$$

$\Leftarrow$

$\uparrow$

Random
Matrix.

$\Leftarrow$

Random
vector.

RM の 特異 値 の singular value を controle
できます か?

→ free probability theory

prob. theory

random variable X

prob. distribution μ_X

expectation
/ variance

...

Independence

$$\mu_{X,Y} = \mu_X \otimes \mu_Y$$

Noncommutative

Probability

element of
algebra

$$a \in A$$

$$n \mapsto \text{tr}(a^n)$$

$$\text{tr}(a^n)$$

$$E[X^n]$$

free product

free probability theory

RM eigenvalue

$$X = W^T W$$

~~singular~~ value, of

$$\{ n \mapsto \text{tr}(X^n) \}$$

$\leadsto$ eigenvalue

固有空間 $U \sim O(n)$
uniform

$\Leftrightarrow$ non-commutative
with other RM

i.e. $U, V \sim \text{unif on } O(n)$

U, V : non-commutative

$\leadsto$ "asymptotically" free

3rd prob.

$$X \perp\!\!\!\perp Y \rightarrow \mu_{X+Y} = \mu_X \star \mu_Y$$

$$\mathbb{E}[e^{it(X+Y)}]$$

$$= \mathbb{E}[e^{itX}] \mathbb{E}[e^{itY}]$$

W_1, W_2 : asym. free

$$\mu_{W_1 + W_2} \approx \mu_{W_1} \boxplus \mu_{W_2}$$

free additive
convolution

$$\mu_{W_1 W_2} \approx \mu_{W_1} \boxtimes \mu_{W_2}$$

$$J = \cancel{D_1} D_2 \dots D_L \quad \cancel{W_1} W_2 \dots W_L$$

$$\Rightarrow \mu_{JJ} \approx \mu_{W_1^{\boxtimes(L-1)} W_L} \boxtimes \mu_{D_1^{\boxtimes(L-1)} D_L}$$

asym. freeness

B. Collins

Comm.

Math.

phys.

trivial
or



AISTATS

Pennington,

arXiv 2018
2019

Marchenko-Pastur

Emergence universality

in Neural Network

$\partial \mathcal{L}$

$$\Theta_{t+1} = \Theta_t - \boxed{\nabla_{\Theta} \mathcal{L}_{\Theta}}$$

$$(\nabla_{\Theta} \mathcal{L}_{\Theta})^T (\nabla_{\Theta} \mathcal{L}_{\Theta}) ?$$

$$\textcircled{H} := \sum_m \frac{\partial f(x(m))}{\partial \Theta}^T \frac{\partial f(x(m))}{\partial \Theta}$$

$$\mathcal{D} = \{ (x_m, y_m) \mid m=1, \dots, N \}$$

Jacot
2018

Neural Tangent Kernel

$$\frac{df_{\Theta}}{d\Theta} = \boxed{\textcircled{H}} (Y - f_{\Theta}(X))$$

$$\begin{cases} Y = (y_1, \dots, y_N) \\ X = (x_1, \dots, x_N) \end{cases}$$

arXiv

Conjugate Kernel

2022 F. Werg.