

Notation (I will try to be consistent!)

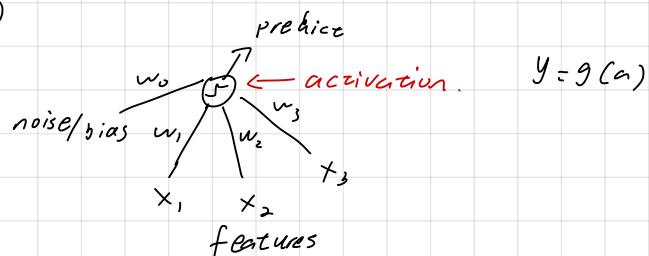
w_i : weight from input i	d : number of inputs
w_{ij} : weight from i to j	c : number of outputs
i : input index	C_k : k^{th} output category
j : hidden unit index	N : number of patterns
k : output unit index (not always!)	n : pattern label
a_j : weighted sum of inputs to unit j	t : target label
$g()$: activation function	x : input data
	y : output activation
	z : hidden unit activation

- Unsupervised. (A model of the data)
- Supervised. (A mapping from input to targets)
- Reinforcement. (learn from mistakes).
- Imitation.

Regression:

There is a underlying line corrupted by noise.

ex)



Linear regression:

$$\begin{bmatrix} 1 & x_1 & \dots & x_d^N \\ 1 & & & \end{bmatrix} W \approx \text{targets}$$

by minimizing the Sum Squared Error (SSE).

$$SSE = \frac{1}{2} \sum_{n=1}^N \left(\underbrace{\sum_{j=0}^d w_j x_j^n}_{\text{predicted}} - t^n \right)^2$$

number of examples/patterns.

Solving for analytical solution: ... $\vec{w} = (X^T X)^{-1} X^T \vec{t}$.

Gradient descent: ...

weighted sum of inputs $a^n = w^T x^n$.

$$y(x) = g(a)$$

↑
activation

Partial derivative of error wrt weights

$$\frac{\partial \text{MSE}}{\partial w_i} = \dots = \frac{1}{N} \sum_{n=1}^N (y^n - t^n) x_i$$

↓ δ^n

$$\text{so, } w_i = w_i + \alpha \frac{1}{N} \sum_{n=1}^N \delta^n x_i^n$$

update rule.

$$= w_i + \alpha \frac{1}{N} \sum_{n=1}^N (t^n - y^n) x_i^n$$

Batch Learning

Online Learning $\rightarrow$ SGD, change weight immediately, for each data point.
(train as data comes in stream).

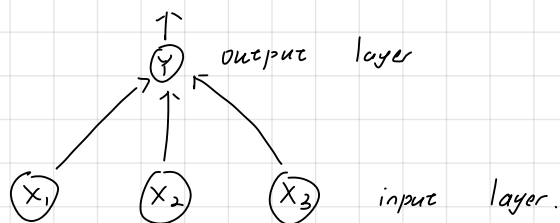
Perceptrons

Trying to fit a function of



(a classifier).

Binary threshold unit:



$$y = \begin{cases} 1 & \text{if } \sum w_i x_i \geq \theta \\ 0 & \text{otherwise.} \end{cases}$$

↑
sum of weights.
(regression).

To make a neurally-inspired machine that
could categorize inputs and learn to do this from example.

Perceptron Learning (perceptron convergence procedure).

ex)

Input	Target
0 0	0
0 1	1
1 0	1
1 1	1

Learn a mapping

Output activation:

$$y = \begin{cases} 1 & \text{if } \sum_{i=0}^d w_i x_i \geq 0 \\ 0 & \text{otherwise} \end{cases} \quad x_0 \triangleq 1$$

Error
Correction
Learning

(only correct it when it is wrong).

If output = 0 and should be 1,
lower weights
If output = 1 and should be 0,
raise weight

⇓

Form a learning rule:

$$w_i = w_i + \alpha (t - y) x_i$$

α is the learning rate

$$\text{or } w_i = w_i + \alpha \delta x_i$$

($\delta = t - y$).

(Delta Rule)

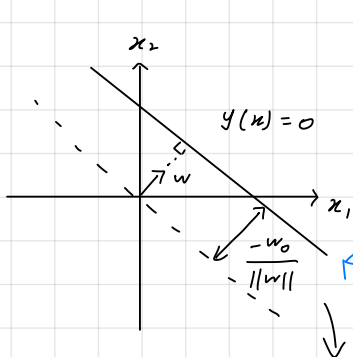
difference

Guarantee: anything a perceptron can compute, can learn to compute and will converge.

We call $y(x) = 0$ the decision boundary where output changes from 0 to 1.

↓

is a $(d-1)$ dimensional hyperplane in d -dimensional input space.



$$y(x) = 0$$

$$\Rightarrow w_1 x_1 + w_2 x_2 + w_0 = 0$$

(2-D case).

$$\Rightarrow x_2 = -(w_1/w_2) x_1 - w_0/w_2$$

line form

For two points on decision boundary, x^A, x^B .

$$y(x^A) = y(x^B) = 0$$

$$\Rightarrow y(x^A) - y(x^B) = 0$$

$$\Rightarrow w^T x^A + w_0 - w^T x^B - w_0 = 0$$

$$\Rightarrow w^T (x^A - x^B) = 0$$

$\Rightarrow \vec{w}$ and vector $(x^A - x^B)$ is orthogonal.

The distance d is then:

$$d = \frac{w^T x}{\|w\|} = \frac{-w_0}{\|w\|}, \text{ where } w = (w_1, \dots, w_d).$$

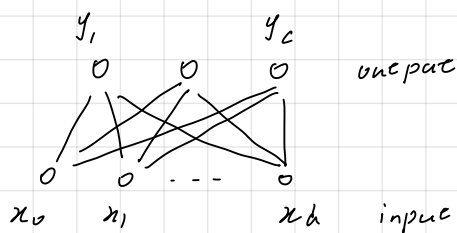
Perceptron as a linear discriminant:

- a two-class discriminant is one such that

$$y(x) = a = w^T x \quad \leftarrow \text{no activation.}$$

x in C_1 if $a \geq 0$
 else x in C_2 .

- For multiple output



x is assigned C_k if $k = \arg\max_j y_j(x)$

in particular,
 each region is convex set
 $\Rightarrow x^A, x^B$ in R_i

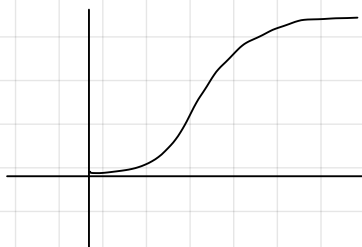


implies $\alpha x^A + (1-\alpha)x^B$ in R_i

Decision boundary is thus

$$y_i(x) = y_j(x) \text{ for } C_i \text{ and } C_j, \dots$$

Logistic Regression



sigmoid function

Estimate of prob. of C_1 or C_2 .

(still a linear classifier).

$$y(x) = g(w^T x + w_0)$$

$$g(x) = \frac{1}{1 + e^{-x}}, \text{ logistic function as activation.}$$

Derivation:

if we assume our data is gaussian distribution (actually not in many cases).

the data distribution $p(x|c_1) = \frac{1}{\sqrt{2\pi}\sigma^2} e^{-\frac{1}{2}\left(\frac{x-\mu_1}{\sigma}\right)^2}$

$$p(x|c_2) = \frac{1}{\sqrt{2\pi}\sigma^2} e^{-\frac{1}{2}\left(\frac{x-\mu_2}{\sigma}\right)^2}$$

By Bayes rule $p(c_1|x) = \frac{p(x|c_1)p(c_1)}{p(x)}$

$$p(c_2|x) = \frac{p(x|c_2)p(c_2)}{p(x)}$$

then, $p(c_1|x) = \frac{p(x|c_1)p(c_1)}{p(x|c_1)p(c_1) + p(x|c_2)p(c_2)}$

Denote $p(c_1|x) = \frac{A}{A+B} = \frac{1}{1+B/A} = \frac{1}{1+e^{\ln(B/A)}} = \frac{1}{1+e^{-\ln(A/B)}}$

$$= \frac{1}{1+e^{-\ln\left(\frac{p(x|c_1)p(c_1)}{p(x|c_2)p(c_2)}\right)}} = \frac{1}{1+e^{-a}} \quad (\text{sigmoid of } a)$$

↑
prob class 1 follows a sigmoid as a function of

the log ratio of the prob. of class c_1 to the prob. of class c_2 .
(the posterior)
called the logit.

So $p(c_1|x) = \frac{1}{1+e^{-a}} = g(a)$, where $a = \ln\left(\frac{p(x|c_1)p(c_1)}{p(x|c_2)p(c_2)}\right)$.

can be written as:

$$a = wx + w_0 \quad \text{where}$$

$$w = \frac{\mu_1 - \mu_2}{\sigma} \quad \text{and} \quad w_0 = -\frac{1}{2}\frac{\mu_1^2}{\sigma^2} + \frac{1}{2}\frac{\mu_2^2}{\sigma^2} + \ln\frac{p(c_1)}{p(c_2)}$$

Most data do not follow gaussian distribution, no analytical solution and we have to learn the weights

⇓
gradient descent

If we use MSE (bad objective function for demonstration):

$$\left\{ \begin{aligned} w_i &= w_i - \alpha \frac{\partial \text{MSE}}{\partial w_i} \\ &\downarrow = \frac{\partial \frac{1}{2N} \sum_{i=1}^N (t^i - y^i)^2}{\partial w_i} = \dots = \frac{1}{N} \sum_{i=1}^N (t - y) \left(-\frac{\partial y}{\partial w_i} \right) \\ &= \frac{1}{N} \sum_{i=1}^N (t - y) \left(-\frac{\partial g(a)}{\partial a} \frac{\partial a}{\partial w_i} \right) \end{aligned} \right. \quad \begin{array}{l} \text{chain rule} \\ \text{for activation.} \end{array}$$

$$= \frac{1}{N} \sum_{i=1}^N (t-y) \left(-g'(a) \frac{\partial a}{\partial w_i} \right)$$

$$= \frac{1}{N} \sum_{i=1}^N (t-y) (-g'(a)) x_i$$

$$w_i = w_i + \frac{\alpha}{N} \sum_{i=1}^N (t^n - y^n) g'(a) x_i^n$$

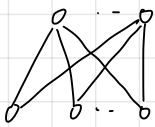
$$= w_i + \alpha \frac{1}{N} \sum_{i=1}^N (t^n - y^n) g(a^n) (1 - g(a^n)) x_i^n$$

since for $g(a) = \frac{1}{1 + e^{-a}}$

$$g'(a) = g(a)(1 - g(a))$$

Use cross-entropy as objective function.

Generalization of logistic regression to multiple output.



output

To turn output to prob.

↓

Softmax

$$a_k = \sum_{j=0}^d w_{jk} x_j = \vec{w}_k^T \vec{x}$$

where $x_0 = 1$,
 $w_{0k} = \text{bias}$.

$$\vec{a} = \vec{w}^T \vec{x}$$

$$y_k = g(a_k) = \frac{e^{a_k}}{\sum_{j=1}^c e^{a_j}} > 0$$

Output sum to 1

$$\sum_{k=1}^c y_k = \sum_{k=1}^c \frac{e^{a_k}}{\sum_{j=1}^c e^{a_j}} = 1$$

Softmax as regression: (can train with GD)

$$a_k = \sum_{j=0}^d w_{jk} x_j = \vec{w}_k^T \vec{x} \quad (x_0 = 1, w_{0k} = \text{bias})$$

$$p(c_k | x) = y_k = g(a_k) = \frac{e^{a_k}}{\sum_{j=1}^c e^{a_j}}$$

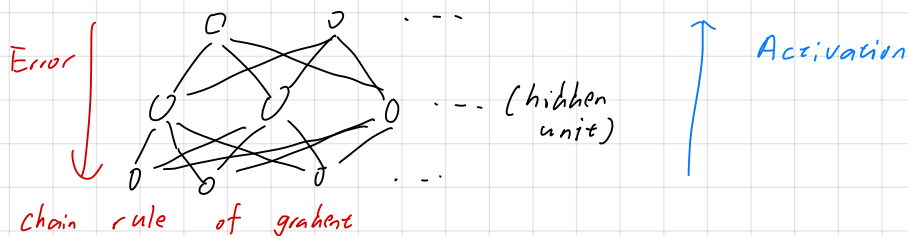
More Activation: $y = g(a) = \max(a, 0)$: ReLU

$y = g(a) = \tanh(a)$: tanh

...

- Create new feature from old feature for learning (old fashion).

Forward Propagation & Backward Propagation:



$$z_i = g(a_i) \xrightarrow{w_{ij}} a_j$$

$$z_i = x_i \xrightarrow{w_{ij}} a_j$$

Back prop. Learning: $w_{ij} = w_{ij} - \frac{\partial J}{\partial w_{ij}}$

(z_i mean output of hidden unit or input).

J is objective.

$$\frac{\partial J}{\partial w_{ij}} = \frac{\partial J}{\partial a_j} \frac{\partial a_j}{\partial w_{ij}} \rightarrow = \frac{\partial \sum_{l=0}^L w_{lj} z_l}{\partial w_{ij}} = \sum_{l=0}^L \frac{\partial w_{lj} z_l}{\partial w_{ij}} = z_l$$

every term is 0 except $i=l$

$$= \frac{\partial J}{\partial a_j} z_i$$

so $z_i = g(a_i) \xrightarrow{w_{ij}} a_j$

• Define $\delta_j = - \frac{\partial J}{\partial a_j} \left(= - \frac{\partial x_{enc}}{\partial a} = - \frac{\partial x_{enc}}{\partial y} \frac{\partial y}{\partial a} = (t - y) \right)$ logistic regression

$$w_{ij} \leftarrow w_{ij} - \frac{\partial J}{\partial w_{ij}} = w_{ij} - \frac{\partial J}{\partial a_j} \frac{\partial a_j}{\partial w_{ij}} \quad x_i = w_{ij} + (t_j - y_j) z_i$$

so, we have $\frac{\partial J}{\partial w_{ij}} = \frac{\partial J}{\partial a_j} \frac{\partial a_j}{\partial w_{ij}} = - \delta_j z_i$

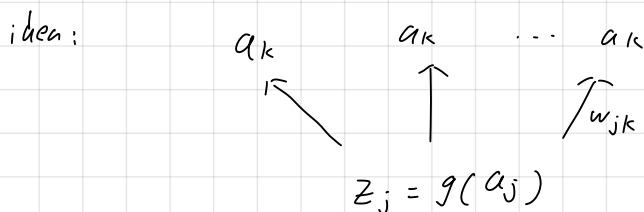
By delta rule, $w_{ij} \leftarrow w_{ij} - \frac{\partial J}{\partial w_{ij}} = w_{ij} + \delta_j z_i$

(weight changes in proportion to the input and the delta at that connection).

For output unit, delta rule is $w_{ij} \leftarrow w_{ij} + (t_j - y_j) z_i$

For hidden unit,

$$\frac{\partial J}{\partial a_j} = \sum_k \frac{\partial J}{\partial a_k} \frac{\partial a_k}{\partial a_j}$$



So, derivation: $\frac{\partial J}{\partial a_j} = \sum_k \frac{\partial J}{\partial a_k} \frac{\partial a_k}{\partial a_j}$

$$= - \sum_k \delta_k \frac{\partial a_k}{\partial a_j}$$

$$= - \sum_k \delta_k \frac{\partial a_k}{\partial z_j} \frac{\partial z_j}{\partial a_j}$$

j independent of $k = - \frac{\partial z_j}{\partial a_j} \sum_k \delta_k \frac{\partial a_k}{\partial z_j}$ def of a_k

$$= - \frac{\partial z_j}{\partial a_j} \sum_k \delta_k \sum_i \frac{\partial w_{ik} z_i}{\partial z_j}$$

every term is 0 except $i=j$

$$= - \frac{\partial z_j}{\partial a_j} \sum_k \delta_k w_{jk}$$

$$= - g'(a_j) \sum_k \delta_k w_{jk}$$

since $\frac{\partial z_j}{\partial a_j} = \frac{\partial g(a_j)}{\partial a_j} = g'(a_j)$
derivative of activation.

Thus, we derive that $\frac{\partial J}{\partial a_j} = -g'(a_j) \sum_k \delta_k w_{jk}$, if j is hidden unit.

Since delta is $-\frac{\partial J}{\partial a_j}$, $\delta_j = g'(a_j) \sum_k \delta_k w_{jk}$.

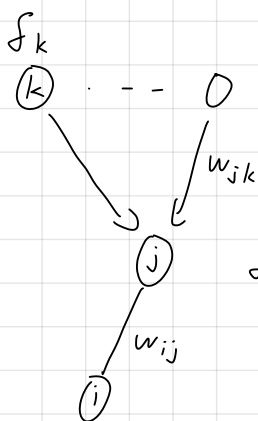
Backprop Learning:

$$w_{ij} = w_{ij} - \alpha \frac{\partial J}{\partial w_{ij}} = w_{ij} + \alpha \delta_j z_j$$

in which $\delta_j = (t_j - y_j)$ if j is an output unit.
 $\delta_j = g'(a_j) \sum_k \delta_k w_{jk}$ if j is a hidden unit

(recursive definition of delta).

For hidden unit, the idea is



$$\delta_j = g'(a_j) \sum_k \delta_k w_{jk}$$

$\uparrow$
activation
of self node

$\uparrow$
sum of previous
delta node times weights.

Back prop is linear operation.
 Vanishing/Exploding gradient
 As result of chain rule of activation.

(online) Back prop :

- present pattern to network, feed forward until output.
- compute δ (usually $t - y$)
- backprop δ through network. Every unit has a delta.
- update w_{ij}

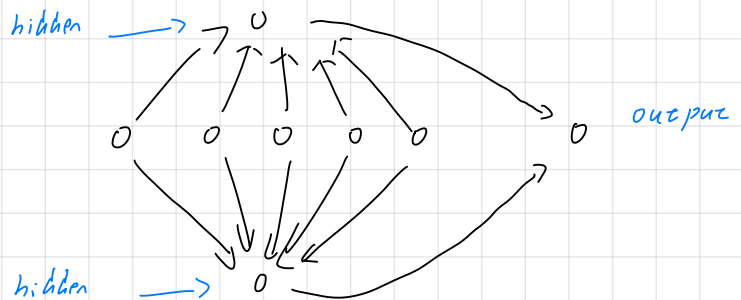
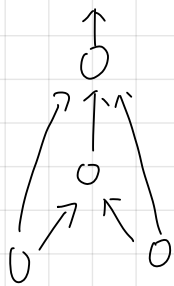
$$w_{ij} \leftarrow w_{ij} + \alpha \delta_j z_j$$

Depend on initial condition, backprop gives many solution.
(highly non-convex network).

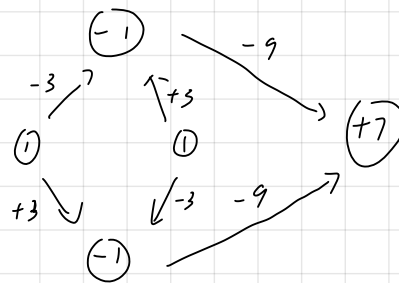
Learn internal representation through hidden units.

ex) XOR

symmetry around center



ex)



ex) NetTalk Architecture

- input uses a one-hot
- averaging vector in tree structure (hidden unit).

ex) Hinton's Family tree

Objective Function & MLE.

Maximum Likelihood.
 want w that maximize $P(w|D) = \frac{P(D|w)P(w)}{P(D)}$
 likelihood · prior
 normalizing constant

↓

$$\max P(w|D) = \max P(D|w)$$

↑ maximize likelihood
 ↓ How we model the data

$$L = \prod_{n=1}^N P(x^n)$$

Find parameter by

$\arg\max L$

(standard distribution example)

$$\Rightarrow \ln L = \ln \prod_{n=1}^N P(x^n)$$

↓

$$\arg\min - \ln \prod_{n=1}^N P(x^n)$$

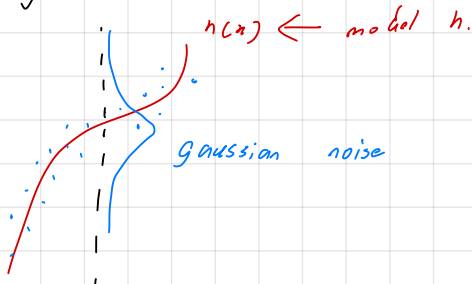
= ...
 = ...
 = leads to setting μ to empirical mean.

Modeling input-output data:

$$L = \prod_{n=1}^N p(x^n, t^n) = \prod_{n=1}^N p(t^n|x^n) p(x^n)$$

$$-\ln L = - \sum_{n=1}^N (\ln p(t^n|x^n) + \ln p(x^n))$$

• Assume gaussian noise



$$\Rightarrow p(t^n|x^n) = \frac{1}{\sqrt{2\pi}\sigma^2} e^{-\frac{(t^n - h(x^n))^2}{2\sigma^2}}$$

↓

$$L = \prod_{n=1}^N \frac{1}{\sqrt{2\pi}\sigma^2} e^{-\frac{(t^n - y(x^n; w))^2}{2\sigma^2}} \quad (\text{for all data})$$

The error (negative log likelihood).

$$E = - \ln \prod_{n=1}^N \frac{1}{\sqrt{2\pi}\sigma^2} e^{-\frac{(t^n - y(x^n; w))^2}{2\sigma^2}}$$

$$= \frac{1}{2\sigma^2} \sum_{n=1}^N (t^n - y(x^n, w))^2 + \ln \left(\sqrt{2\pi}\sigma^2 \right)^N$$

↓
 sum squared loss!
 constant

Cross-Entropy:

$$y(x^n) = P(c_i|x^n)$$

Modeling two categories is like modeling coin flip

↓
 Bernoulli distribution

↓

$$p(t^n | x^n) = (y^n)^{t^n} (1 - y^n)^{(1-t^n)}$$

if $t^n = 1$, $p(t^n | x^n) = y^n (1 - y^n)^{(1-1)} = y^n$.

so, the likelihood: $L = \prod_{n=1}^N p(t^n | x^n) = \prod_{n=1}^N (y^n)^{t^n} (1 - y^n)^{(1-t^n)}$

$$-\ln L = -\ln \prod_{n=1}^N (y^n)^{t^n} (1 - y^n)^{(1-t^n)}$$

$$= -\sum_{n=1}^N \ln (y^n)^{t^n} (1 - y^n)^{(1-t^n)}$$

$$X_{ent} = -\sum_{n=1}^N \underbrace{t^n \ln(y^n) + (1-t^n) \ln(1-y^n)}$$

(cross-entropy error)

Entropy: a measure of information in a message about a random variable.

$-p(x) \log p(x)$

Entropy of distribution: $-\sum_{n=1}^N p(x^n) \ln p(x^n)$.

↓
cross entropy: $-\sum_{n=1}^N p(t^n) \ln p(y^n)$.

Multinomial Regression:

$p(c_k | x^n) = y_{k^n}(x^n)$, $t_k = 1$ if from category k , 0 otherwise

$$\Rightarrow p(t^n | x^n) = \prod_{k=1}^C (y_{k^n})^{t_{k^n}}$$

$$\Rightarrow \prod_{n=1}^N p(t^n | x^n) = \prod_{n=1}^N \prod_{k=1}^C (y_{k^n})^{t_{k^n}} = \dots = -\sum_{n=1}^N \sum_{k=1}^C t_{k^n} \ln y_{k^n}$$

Siamese Neural Network:



"train to move output closer or farther apart."

if x_1, x_2 from same category, minimize dist y^1 and y^2 .

... difference ... maximize ...

The loss function: $L(w, y, \vec{x}_1, \vec{x}_2) = (1-y) \frac{1}{2} (D_w)^2 + (y) \frac{1}{2} \{\max(0, m - D_w)\}^2$.

- D is distance between outputs of two networks
- m is a margin.
- y is 1 for a "different" pair 0 for "same" pair.

↓ $D_w(\vec{x}_1, \vec{x}_2) = \|\mathbf{G}_w(\vec{x}_1) - \mathbf{G}_w(\vec{x}_2)\|_2$
↑
output vector
of network on stimulus x_i .

Generalization

Deal With Overfitting?

- More data  Data Augmentation:
ex) take different crops of the data or resize it.
- Regularization.
 - ↓
reduce model complexity ($J = E + \lambda C$)
 - $L_2 \rightarrow$ minimize $\|w\|_2^2$, weight smaller in proportion to its size.
 - $L_1 \rightarrow$ minimize $|w|$, weight smaller at a constant rate.
 - Minimize $C = \|w\|^2 / (\|w\|^2 + 1)$, penalize big weight less while penalizing small weights more.
- Dropout
 - $\Rightarrow$ probabilistically turn off some fraction of hidden units.
- Early stopping
 - ...
- Add noise to inputs/weight/hidden unit activations.
 - ↓
make model more robust to perturbations.

Neural Net Trick:

SGD:

1. get one example
2. compute gradient
3. change weight
4. If seen all example, reshuffle

Batch Learning

1. get one example
2. compute gradient
3. add to running average
4. If seen all example, change weight

mini-batch learning.
 $\rightarrow$ by running a small batch of data each time changing the weights.

↓
efficient due to memul.

Shuffling data before partition to mini-batch

↓
should have diverse data representation.

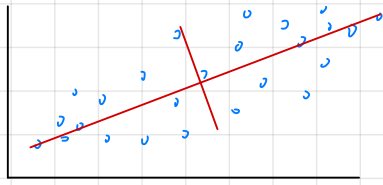
PCA on input:

idea: ex) if all positive input, $w_{ij} = w_{ij} + \alpha f_j x_i$

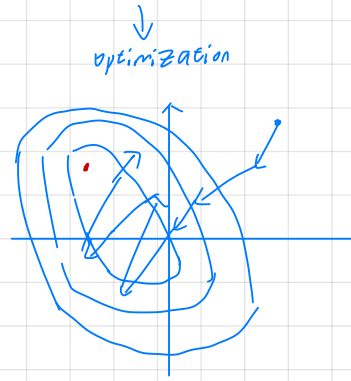
↑
same sign for all weights.

if x_i 's are all positive, then the weights changes are all + or -.

ex) highly correlated variables



$\Rightarrow$ get two different information.



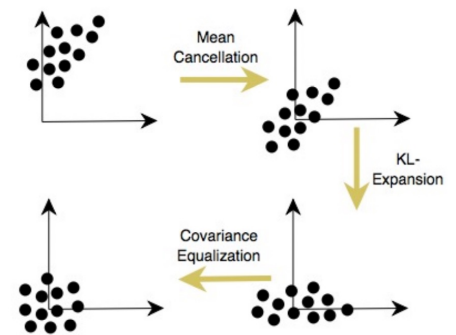
ex) different scale of input

PCA $\rightarrow$ shifts mean of input variable to 0.

- decorrelates the input.
- throw away dimension with smaller eigenvalues (dimensionality reduction).
- divide by standard deviation, make them all of equal size.

PCA vs. Z-scoring:

- Z-scoring shifts the mean of the input variables to be 0 – not all positive or negative!
- Z-scoring does not decorrelate the inputs – because it is applied to each variable independently.
- And...you cannot throw away the dimensions with the smallest eigenvalues :
–because you don't compute the eigenvectors.
- Z-scoring divides by the standard deviation, making all of the variables 0 mean and unit standard deviation.
- So the variables are all of the same size



1/18/24

Neural Networks/Deep Learning

Different activation function:

problem: standard sigmoid makes all input to next layer positive.

recommended sigmoid: $f(x) = 1.7159 * \tanh(0.667x)$.

• $f(+/-1) = +/-1$ (if input has unit variance, output has unit variance).

Weight Initialization:

Weight cannot initialize to 0

if initialize to 0 $\Rightarrow$ f's of the hidden will all be the same.

$\Downarrow$
weights from each input be the same.

Weight initialization coordinated with:

- Input normalization
- choice of sigmoid.

Suppose x_i mean 0, unit variance. w_{ij} mean 0, on average, STD of $a_j = xw$ will be.

$$\begin{aligned}\text{Var}(XW) &= \text{Var}(X) \text{Var}(W) + \text{Var}(X) (E(W))^2 + \text{Var}(W) (E(X))^2 \\ &= \text{Var}(W).\end{aligned}$$

$$\text{So, STD of } a_j = \sqrt{\text{Var}(W)} = \sqrt{\sum_i (w_{ij})^2}$$

To let $\sigma(a_j) = 1$, we set w 's to be $\left(\frac{1}{m}\right)^{1/2}$.

where m is the fan-in to node j .

assume input have been normalized, and sigmoid is $f(x) = 1.7159 \tanh(2/3 x)$.

Intuitively, if hidden unit has big fan-in, small changes on many of its incoming weights can cause learning to overshoot.

$\Rightarrow$ smaller incoming weights when fan-in is big, and vice-versa.

Batch Normalization

idea: normalizes all inputs and inputs to each hidden layer throughout the network.

(on a per unit basis, over each minibatch).

1. Z-score each input variable over mini-batch.

$$\hat{a}_i = \frac{a_i - \mu_i}{\sqrt{\sigma_i^2 + \epsilon}} \quad \text{where } \mu_i, \sigma_i^2 \text{ are mean and variance of } a_i$$

$\uparrow$
new input to a unit on any layer of the network.

2. Allow network learn to undo if necessary:

$$\tilde{a}_i = \gamma \hat{a}_i + \beta;$$

learnable parameter

Input: Values of a over a mini-batch: $\mathcal{B} = \{a_{1...m}\};$

Parameters to be learned: γ, β

Output: $\{\tilde{a}_i = \text{BN}_{\gamma, \beta}(a_i)\}$

$$\mu_{\mathcal{B}} \leftarrow \frac{1}{m} \sum_{i=1}^m a_i \quad // \text{ mini-batch mean}$$

$$\sigma_{\mathcal{B}}^2 \leftarrow \frac{1}{m} \sum_{i=1}^m (a_i - \mu_{\mathcal{B}})^2 \quad // \text{ mini-batch variance}$$

$$\hat{a}_i \leftarrow \frac{a_i - \mu_{\mathcal{B}}}{\sqrt{\sigma_{\mathcal{B}}^2 + \epsilon}} \quad // \text{ normalize}$$

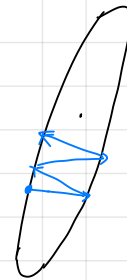
$$\tilde{a}_i \leftarrow \gamma \hat{a}_i + \beta \equiv \text{BN}_{\gamma, \beta}(a_i) \quad // \text{ scale and shift}$$

Momentum

weight change at time t .

$$\Delta w(t) = \underbrace{\gamma \Delta w(t-1)}_{\text{previous information}} - \alpha \frac{\partial E}{\partial w(t)}$$

(sum of previous weight change and the new gradient).



direction of steepest descent does not point at the minimum.

How it behaves:

$$\Delta w(t) = \frac{1}{1-\gamma} \left(-\alpha \frac{\partial E}{\partial w} \right) \quad (\text{by solving recurrence relation}).$$

Nesterov Method:

- first make a big jump in direction of previous accumulated gradient.
- then measure the gradient where you end up and make a correction.

Adaptive learning rate:

$$\Delta w_{ij} = - \sum g_{ij} \frac{\partial E}{\partial w_{ij}}$$

global
learning
rate

local
adjustment.

$$g_{ij}(t) = g_{ij}(t-1) + 0.05.$$

$$\left(\frac{\partial E}{\partial w_{ij}}(t) \frac{\partial E}{\partial w_{ij}}(t-1) \right) > 0$$

- increase if $\text{gradient}_{(ij)}$ does not change sign.
- decrease otherwise.

$$g_{ij}(t) = g_{ij}(t-1) \cdot 0.95.$$

$\Rightarrow$ big gain decay rapidly when oscillation starts.

Adaptive learning rate only deals with axis-aligned effects.

$\downarrow$
Adaptive learning rate affects individual weights, which change the learning rate along that axis.

Resilient Backprop (Rprop):

only for full-batch learning:

- increase step size for a weight multiplicatively (times 1.2) if the signs of its last two gradients agree
- otherwise, decrease the step size multiplicatively (times 0.5).

Does not work for mini-batch learning:

Rprop is to use the gradient but also divide by the size of the gradient $\left(\text{sign}(x) = \frac{x}{|x|} \right)$, problem is that we divide by a different number for each mini-batch.

$\downarrow$
Solution — RMSprop:

Keep a moving average of squared gradient for each weight.

$$\text{Mean Square}(w, t) = 0.9 \text{Mean Square}(w, t-1) + 0.1 \left(\frac{\partial E}{\partial w(t)} \right)^2$$

Then, divide the gradient by $\sqrt{\text{Mean Square}(w, t)}$.

Convolutional Network: (CNN)

Problem: • Large images size (1140 x 648)

Properties of Real World Image

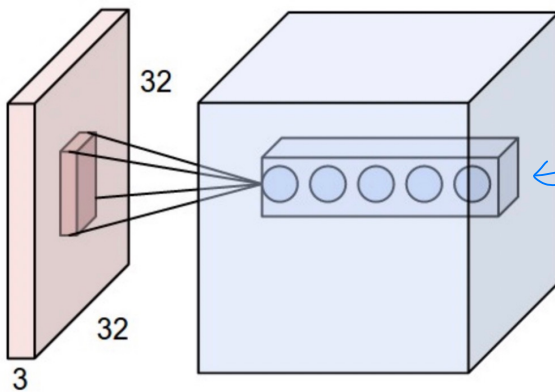
- Nearby pixel correlate with nearby pixel: locality
- Statistics of pixel are relatively uniform: stationary statistics
- object doesn't depend on its location in image: translation invariance
- object made of parts: Compositionality

Why CNN useful to solve above problem?

⇓

- "Locality": small, local receptive field and learned features (kernel).
- "Stationary statistics": replicated across the image
(if a feature is useful in one place, it's useful in others).
- "Translation invariance": spatial pooling.
- Objects are made of parts: receptive fields get larger deeper in the net.

Receptive field { size of visual world that innervates the neuron
visual feature in visual world that most stimulate neuron.



Convolutional Block.

5 neurons

each one compute a different features.

↓

computation is exactly the same

(inner product of weights and pink patch).

Activation Maps

pooling $\longrightarrow$  $\Rightarrow$.

2x2 set of response to 1 number (Translation Invariance).
(1/4 size)

convolution leads to shift equivariance.

(shifting number in different location does not affect the feature map).

Original image Feature map for 1 feature



If we move the 3 over... the features just shift over



This is shift equivariance

Image Net Data

Alex Net 8 layers



VGG 14 layers



Google Net 22 layers



ResNet 152 layers

Overview of convnets.

- Feed-forward

- Convolve input.

- Non-linearity (rectified linear)

- Pooling (local max).

- Supervised

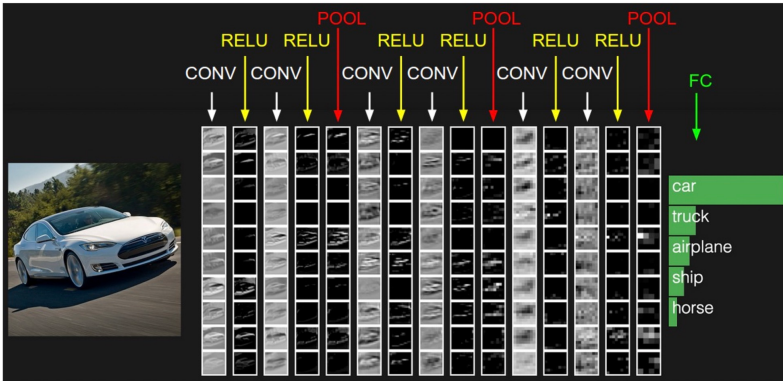
- Convolution filter trained by backprop.

Images $\longrightarrow$ Convolution $\longrightarrow$ Non-linear $\longrightarrow$ Pooling $\longrightarrow$ Feature maps
(filter)

Notes From Stanford CS231n. CNN

ex) (12 filters) $(softmax)$
INPUT $\rightarrow$ CONV $\rightarrow$ RELU $\rightarrow$ POOL $\rightarrow$ FC
 $[32 \times 32 \times 3]$ $[32 \times 32 \times 12]$ $[32 \times 32 \times 12]$ $[16 \times 16 \times 12]$ $[1 \times 1 \times 10]$

	CONV	RELU	POOL	FC
contain parameters	X			X
hyperparameter	X		X	X



Convolutional Layer

- each filter will produce a 2-D activation map
- stacking activation map along the depth dimension produce the output volume.
- connect each neuron to only a local region of the input volume (receptive field).
- the extend of connectivity along the depth axis is always equal to the depth of input volume.

ex) input volume $[32 \times 32 \times 3]$

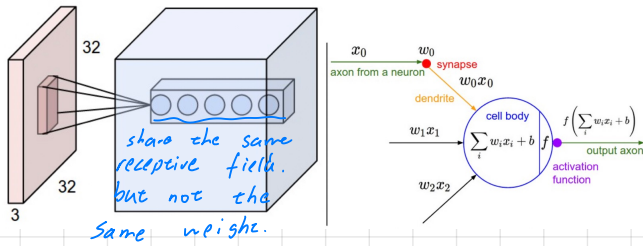
receptive field (filter size) $[5 \times 5 \times 3]$.

then a total of $5 \times 5 \times 3 = 75$ weights.

ex) input volume $[16 \times 16 \times 20]$.

receptive field $[3 \times 3 \times 20]$.

then a total of $3 \times 3 \times 20 = 180$ weights.



Three hyperparameters to control output volume

• Depth: # of filters (K)

• Stride: # of jumps when filter moves

• Zero-padding: whether to pad input volume with 0 around the border.

— — — — —

- W : input volume size
- F : receptive field size
- S : stride
- P : amount of zero padding

Then output volume size is

$$\frac{(W - F + 2P)}{S} + 1$$

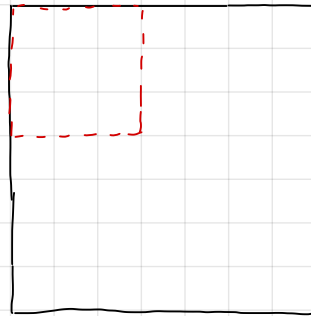
↑

choice of Stride is important to ensure that the above division yields an integer.

Thus, to have input and output the same size, we see zero padding

$$P = (F - 1) / 2$$

3x3 filter



7x7

stride = 1, pad = 0

$$\begin{aligned} \Rightarrow & \frac{7 - 3 + 2 \cdot 0}{1} + 1 \\ & = 4 + 1 = 5 \\ \Rightarrow & 5 \times 5 \text{ output.} \end{aligned}$$

Parameter sharing.

Since we are using filter striding over the images, the # of weights will be

$$F \times (\text{\# of filter}) \times (\text{depth of output}).$$

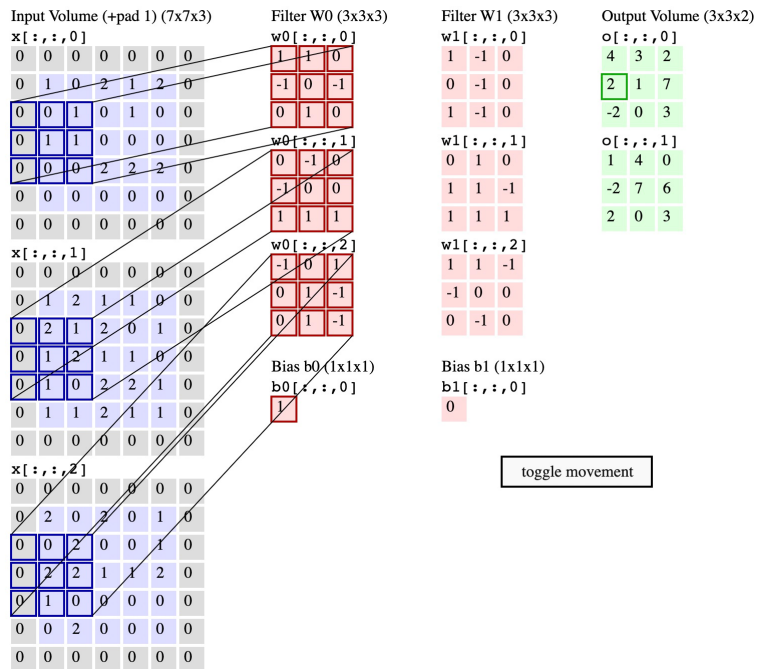
this is guaranteed by translationally invariance, if detecting a feature is important at some location, it should be useful at some other location as well.

If learning difference feature on different location of images is necessary, then it's common to relax parameter sharing schema



Locally-Connected Layer.

Conv a complete visualization



Pooling layer

- progressively reduce the spatial size of the representation to reduce the parameters.

• W_1, H_1, D_1 representing input volume size.

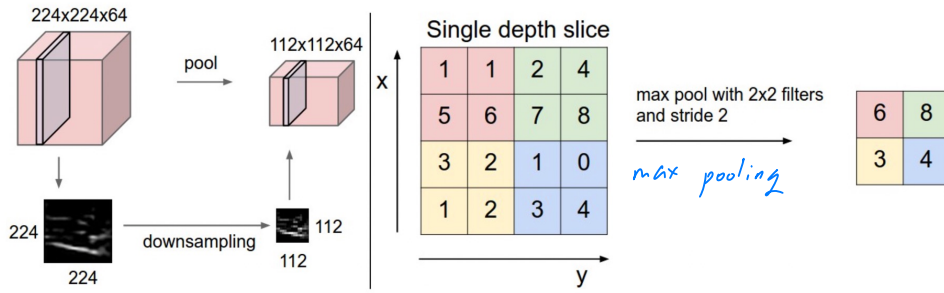
• For pooling layer, it requires two hyperparameters:

- spatial extent F
- stride S

• Then output volume: $W_2 = (W_1 - F) / S + 1$

$$H_2 = (H_1 - F) / S + 1$$

$$D_2 = D_1$$



Pooling layer downsamples the volume spatially, independently in each depth slice of the input volume. Left: In this example, the input volume of size [224x224x64] is pooled with filter size 2, stride 2 into output volume of size [112x112x64]. Notice that the volume depth is preserved. Right: The most common downsampling operation is max, giving rise to max pooling, here shown with a stride of 2. That is, each max is taken over 4 numbers (little 2x2 square).

Conv net Architectures

Most common form:

INPUT $\rightarrow$ $[(CONV \rightarrow RELU) \times N \rightarrow POOL ?] \times M \rightarrow [FC \rightarrow RELU] \times K \rightarrow FC$

↑
optional pooling

Filtering

- (Weights learned during training)
- each unit looks at local window.
- same feature computed every-where.

Non-linearity:

- ReLU (applied per-pixel).
- $\text{output} = \max(0, \text{input})$.

Spatial Pooling

- max is best
 - sum or average work less
 - applied to each filter layer separately.
 - Stride: how far to slide over each image
- role: invariance to small transformations.
larger receptive field.

Components of each layer:

pixel/feature $\Rightarrow$ filter with learned dictionary $\Rightarrow$ Non-linearity



Spatial local max pooling



Output features.

- Batch normalization.

How to choose architecture?

- Hyperparameter
 - ex) # layers, # feature maps, params of feature maps.

- Cross-validation
- Grid-search / random-search.

Best practice:

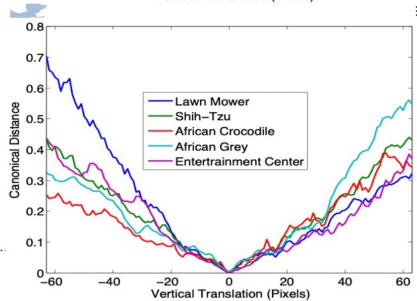
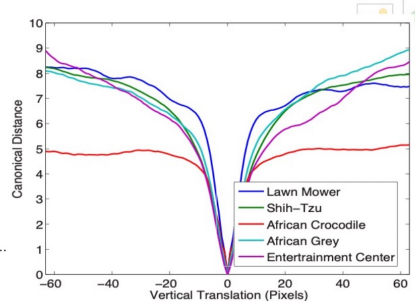
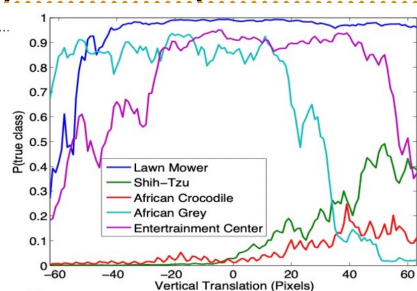
- width: 3
- stride: 1
- zero-padding: 1

Trick: Stack convolutions, then use max pooling.

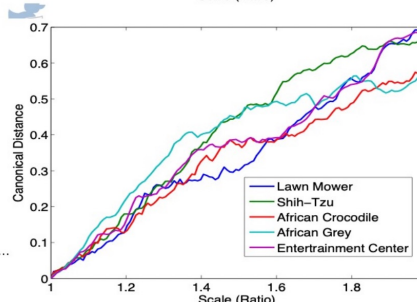
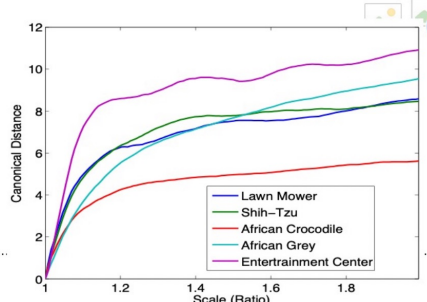
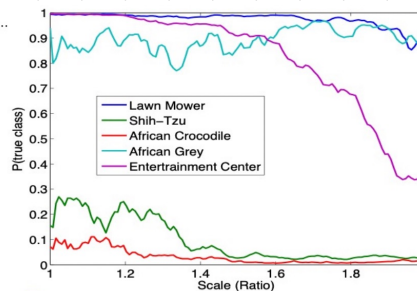
conv $\rightarrow$ conv $\rightarrow$ max pooling

Invariance Abilities

Translation (Vertical):



Scale Invariance:



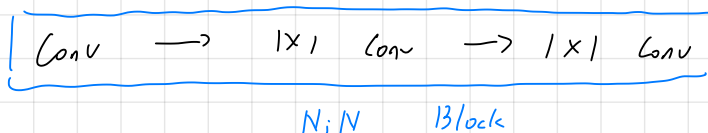
Deep Model:

idea: design one modular structure, then just add more of them

3x3 conv layer better than single 7x7 b/c non-linearizing.

• Network-in-network:

add two per-pixel fully connected layers before pooling.



• global average pooling.

replace final fully connected layer: one feature map per category.

• What is 1x1 conv?

$$[56 \times 56 \times 64] \xrightarrow{1 \times 1 \text{ conv}} [56 \times 56 \times 32]$$

with 32 filters

(each filter size $(1 \times 1 \times 64)$, performs a 64-dim dot product).

$\Rightarrow$ essentially a dot product of filter weights and 64 channels

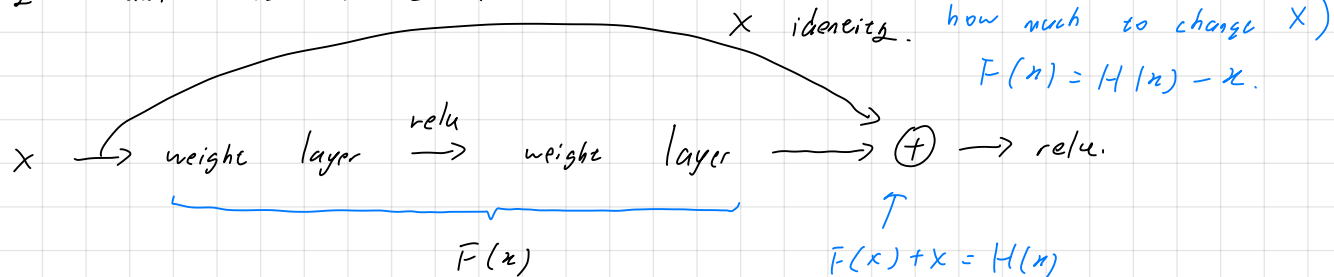
advantage: complete control of computation !!!

Residual Network

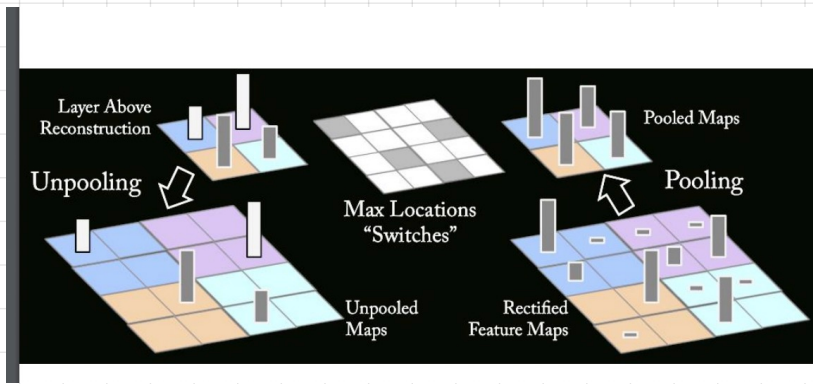
idea: introduce "pass through" into each layer.

$\Downarrow$

only residual needs to be learned.

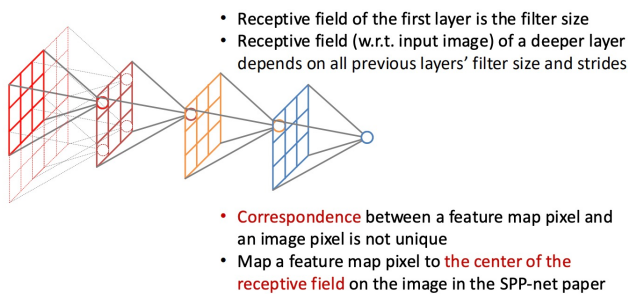


Visualization for higher layer network:



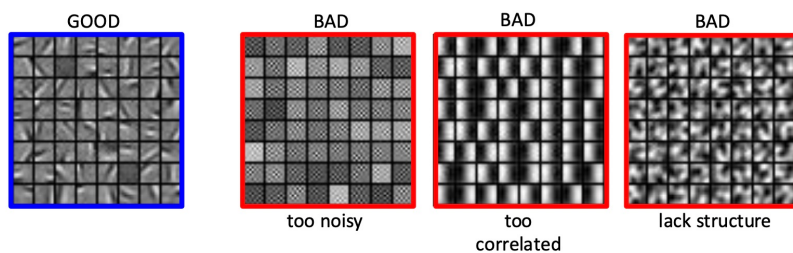
Max unpooling

- max activation from feature map associated with each filter.
- deconvolve to project back to pixel space.
- pooling switch.



Training Tips

- Annealing learning rate.
↳ start large, slowly reduce.
- Visualize feature map.



Good training: learned filters exhibit structure and are uncorrelated.

- Training diverges → decrease learning rate
- Network underperforming → make net larger.

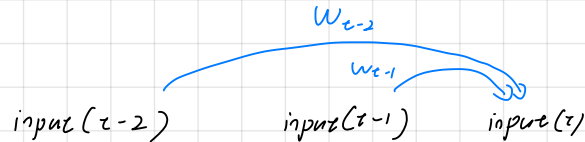
Recurrent Network:

Represent time in network?

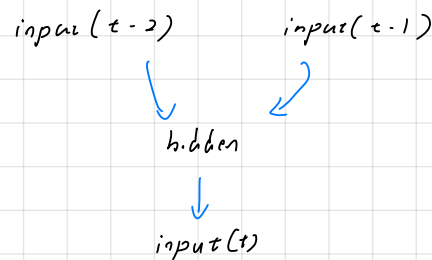
- map time into space
- map time into state of the network.

Time $\rightarrow$ Space.

• Auto regressive:



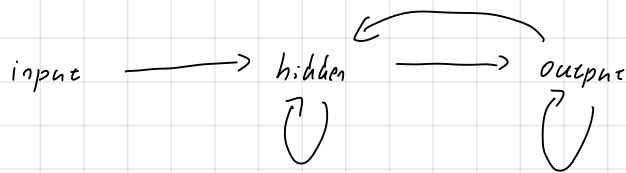
• Feed-forward net:



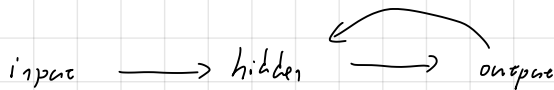
Time $\rightarrow$ State

• use activation memory:

new things are processed based upon what has come before.
(network state).



=> ① Jordan net

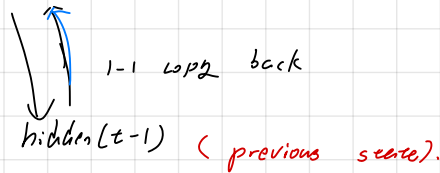


② Elman net (simple recurrent net).



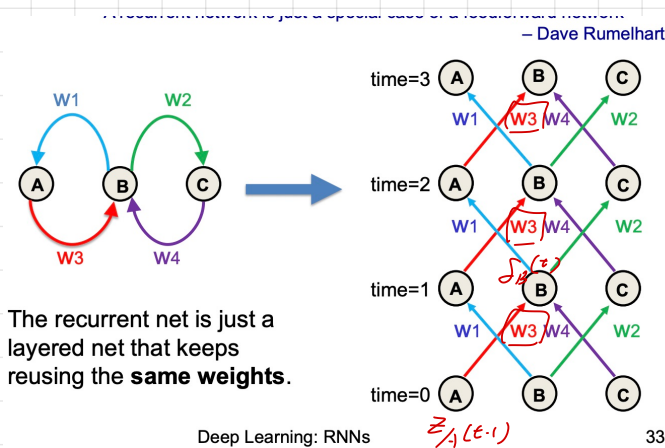
equivolent,

input $\rightarrow$ hidden $\rightarrow$ output



one-step backprop through time.

Unrolling networks in time.



• calculate total gradient over time as an average and then modify the weights.



• propagate activation forward in time ($z_A(t), z_B(t), \dots$)

• propagate delta backward.

• update weights.

$$ex) \quad W_{3+} = d \frac{1}{3} \sum_{t=1}^3 \delta_B(t) z_A(t-1)$$

Hard to train b/c backward pass is linear.

problem of exploding/vanishing gradient:

• weights small, gradient shrinks exponentially.

• weights big, gradient grow exponentially.

→ • switch logists to ReLU.

• In KNN, we can:

• exploding gradient.

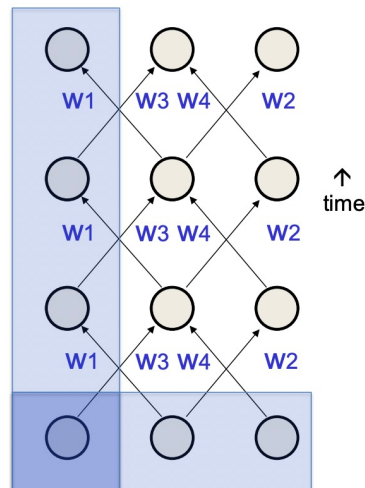
limit length of gradient vector (gradient clipping)

• vanishing gradient:

skip connection

We can specify inputs in several ways:

- Specify the initial states of all the units.
- Specify the initial states of a subset of the units (the rest are hidden)
- Specify the states of the same subset of the units at every time step. I.e., we have inputs at every time step.



-
- A diagram showing a 3x3 grid of nodes (circles) connected by arrows. The nodes are arranged in three rows and three columns. The connections are as follows:
 - Row 1: Node (1,1) connects to (2,1) and (2,2); Node (1,2) connects to (2,1), (2,2), and (2,3); Node (1,3) connects to (2,2) and (2,3).
 - Row 2: Node (2,1) connects to (3,1) and (3,2); Node (2,2) connects to (3,1), (3,2), and (3,3); Node (2,3) connects to (3,2) and (3,3).
 - Row 3: Node (3,1) connects to (4,1) and (4,2); Node (3,2) connects to (4,1), (4,2), and (4,3); Node (3,3) connects to (4,2) and (4,3).
 The nodes are labeled with weights: W1 (bottom-left), W3 (bottom-middle), W4 (bottom-right), and W2 (top-right). The top row of nodes is highlighted in light blue. The middle row of nodes is highlighted in light red. The rightmost column of nodes is highlighted in light green.

A Recurrent Network for Language Generation

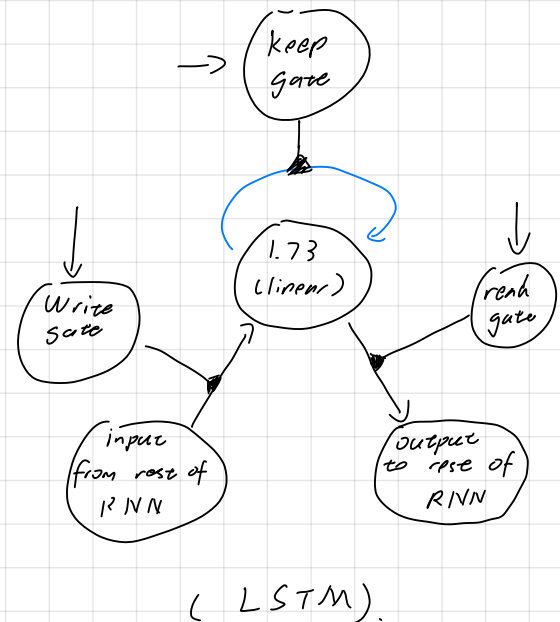
-
- Diagram illustrating a neural network architecture for character classification. The network consists of an input layer, a hidden layer, and an output layer.
- Input Layer:** Takes four inputs: 'h' (red), 'e' (green), 't' (blue), and 't' (blue).
- Hidden Layer:** Consists of three nodes. The first node receives input from 'h' and 'e'. The second node receives input from 'e' and the first hidden node. The third node receives input from 't' and the second hidden node.
- Output Layer:** Consists of four nodes. The first node receives input from the first hidden node and produces output 'e' (red). The second node receives input from the second hidden node and produces output 't' (green). The third node receives input from the third hidden node and produces output 't' (blue). The fourth node receives input from the third hidden node and produces output 'o' (blue).
- Weights:** The weights between the input and hidden layers are labeled W_{oh} . The weights between the hidden and output layers are labeled W_{hh} and W_{ho} .
- Target Characters:** The target characters are 'e', 't', 't', and 'o'.
- Output Layer Weights:** The weights for the output layer are:
- From hidden node 1 to 'e': 1.0, 2.2, -3.0, 4.1
 - From hidden node 2 to 't': 0.5, -0.3, 1.0, 1.2
 - From hidden node 3 to 't': 0.1, 0.5, 1.9, -1.1
 - From hidden node 3 to 'o': 0.2, -1.5, -0.1, 2.2
- Hidden Layer Weights:** The weights for the hidden layer are:
- From input 'h' to hidden node 1: 0.3, -0.1, 0.9
 - From input 'e' to hidden node 1: 1.0, 0.3, 0.1
 - From input 't' to hidden node 2: 0.1, -0.3, -0.3
 - From input 't' to hidden node 3: -0.3, 0.7, 0.7
- Input Layer Weights:** The weights for the input layer are:
- From 'h' to hidden node 1: 1, 0, 0, 0
 - From 'e' to hidden node 1: 0, 1, 0, 0
 - From 't' to hidden node 2: 0, 0, 1, 0
 - From 't' to hidden node 3: 0, 0, 0, 1

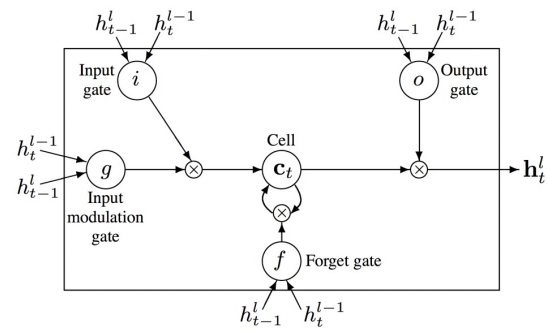
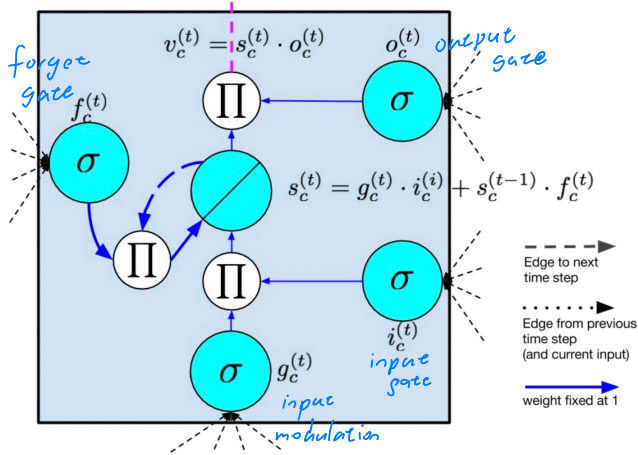
TIME ->

$$h^{(t)} = \sigma(W_{hx} x^{(t)} + W_{hh} h^{(t-1)} + b_h).$$

The diagram illustrates a recurrent neural network structure across two time steps, $t=1$ and $t=2$. At each time step, there are three nodes: an input node (blue circle), a hidden layer node (green circle), and an output node (red circle). Solid black arrows represent connections within the same time step: from input to hidden layer, from hidden layer to output, and from output back to hidden layer. Dashed magenta arrows represent recurrent connections between hidden layer nodes across time steps. A legend at the bottom indicates that a dashed arrow points to the next time step.

- info goes into cell when "write" gate on.
- info stays in cell when "keep" gate on.
- info read from cell when "read" gate on.





$$\begin{pmatrix} i \\ f \\ o \\ g \end{pmatrix} = \begin{pmatrix} \text{sigm} \\ \text{sigm} \\ \text{sigm} \\ \text{tanh} \end{pmatrix} T_{2n,4n} \begin{pmatrix} D(h_{t-1}^{l-1}) \\ h_{t-1}^{l-1} \end{pmatrix}$$

$$c_t^l = f \odot c_{t-1}^l + i \odot g$$

$$h_t^l = o \odot \tanh(c_t^l)$$

LSTM Forward Pass :

$$g^{(t)} = \phi(W_{gx} x^{(t)} + W_{gh} h^{(t-1)} + b_g) \quad \text{input modulation}$$

$$i^{(t)} = \sigma(W_{ix} x^{(t)} + W_{ih} h^{(t-1)} + b_i) \quad \text{input gate}$$

$$f^{(t)} = \sigma(W_{fx} x^{(t)} + W_{fh} h^{(t-1)} + b_f) \quad \text{forget gate}$$

$$o^{(t)} = \sigma(W_{ox} x^{(t)} + W_{oh} h^{(t-1)} + b_o) \quad \text{output gate.}$$

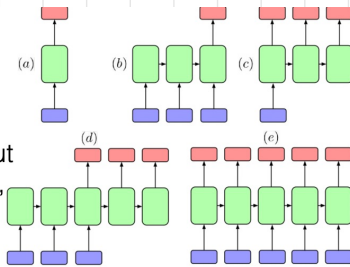
$$s^{(t)} = g^{(t)} \odot i^{(t)} + s^{(t-1)} \odot f^{(t)}$$

$$h^{(t)} = s^{(t)} \odot o^{(t)} \quad \text{control forget}$$

input modulation

RNN architecture

- One to one (not recurrent!)
- Many to one
- One to many
- Many in, then many out
- Simultaneous many in, many out



RNN as generative model

- generative text model ...
- Seq to Seq
- Image to text
- ...

Memory is a matrix of linear neurons.

$$M = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \quad (\text{row as a "work" of memory}).$$

index row as $i = 1, 2, 3$.

$$r_t \leftarrow \sum_i w_t(i) M_t(i)$$

Transformer Network

Map time into space using shared weights

↓
communicate via attention

⇓
when you focus on particular aspect of your env.

- Overt attention: move your eyes towards something
- Covert attention: direct your attention to something without moving your eyes

How to direct attention in Neural Net?

↓
Softmax + multiplicative connections. (Content-based Addressing).

$$w_i^c(i) \leftarrow \frac{\exp(\beta \mathbf{k} \cdot [\mathbf{k}_t, \mathbf{M}_t(i)])}{\sum_j \exp(\beta \mathbf{k} \cdot [\mathbf{k}_t, \mathbf{M}_t(j)])}$$

ex) $\mathbf{M} = \begin{matrix} & \begin{matrix} 3 & 1 & 4 & 1 \end{matrix} \\ \begin{matrix} 1 \\ 2 \\ 2 \end{matrix} & \begin{matrix} 2 & 3 & 4 \\ 7 & 2 & 8 \end{matrix} \end{matrix}$

key $\mathbf{k} = (3, 1, 0, 0)$.

attention weights: $\mathbf{w}^c = (0.998, 0.001, 0.001)$.

read: $\mathbf{r}_t \leftarrow \sum_i w_i^c(i) \mathbf{M}_t(i)$

so, $\mathbf{r}_t = (2.997, 1.007, 3.997, 1.010)$
 $\approx (3, 1, 4, 1)$

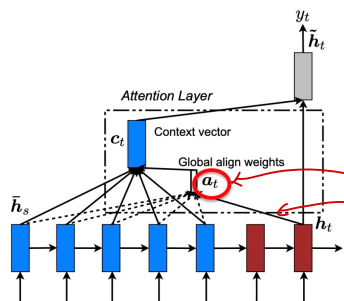
- key vector was computed by the network.
(dynamic, network know what it was looking for).
- Attention weights filtered the memory row via multiplication.
- Match b/t key and memory don't have to be perfect.

Why transformer? $\longrightarrow$ decoder needs different information at different time steps.

- train much faster
- process each position in parallel.

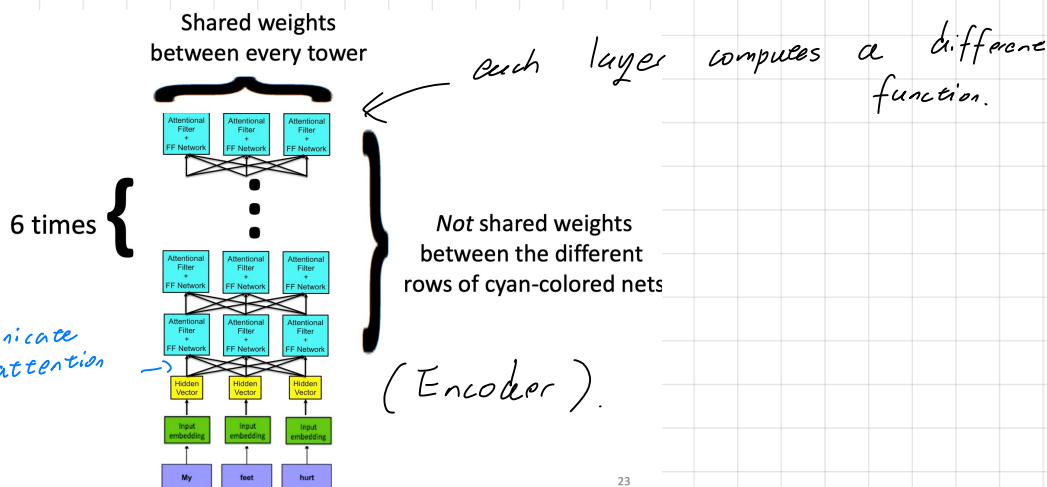
Attention-based methods $\longrightarrow$ problem with recurrence.

No parallelization since hidden states all depend on previous hidden states.



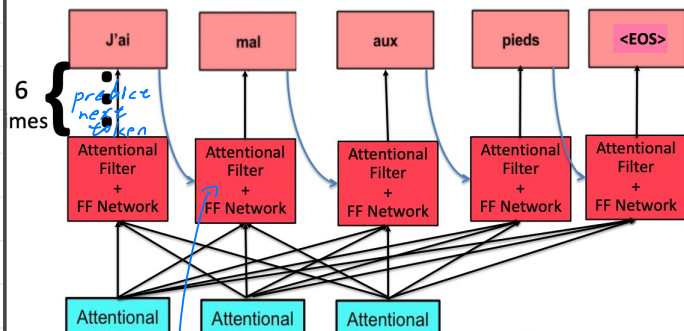
attention vector computed dynamically
used to gate the previous hidden state.

Transformer Network:

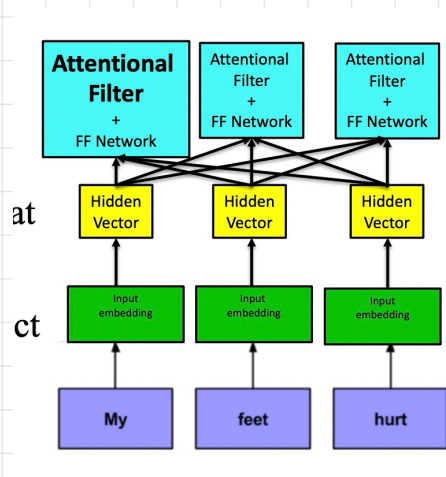


a copy of network for every input position.

(decoder)



retrieves info from all encoder states.



inner produce attention.

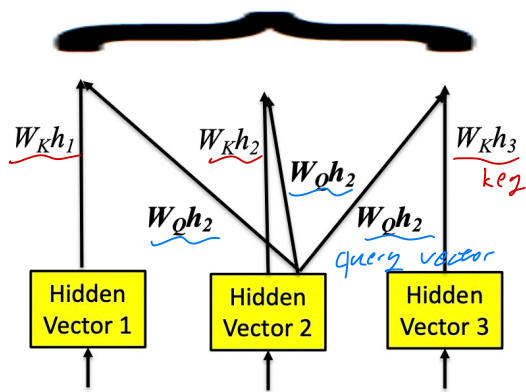
- each tower generates a vector that represents what it is looking for. (Query)
- each tower generates a vector that represents what information it has. (Key).

Query + Key

=> allow one tower to select information from other towers.

(dynamically re-wire itself and communicate between word meaning to construct sentence meaning).

Tower 2 Softmax



$$(query\ vector) \cdot (W_k h_i),$$

↖ the i -th hidden vector

$$\Rightarrow softmax((W_k h_1)^T (W_q h_2), (W_k h_2)^T (W_q h_2), \dots (W_k h_N)^T (W_q h_2)) / \sqrt{d}.$$

softmax "interest" vector over N hidden vectors. $\therefore$ dot product all tower's key.

idea: the most similar key and query vector will be a peak in softmax, selecting where to obtain value.

In matrix notation:

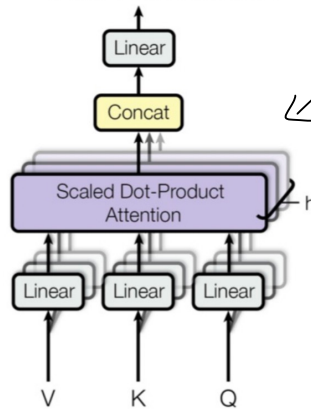
$$W_{key} H = K \quad (key)$$

$$W_{query} H = Q \quad (query)$$

$$W_{value} H = V \quad (value)$$

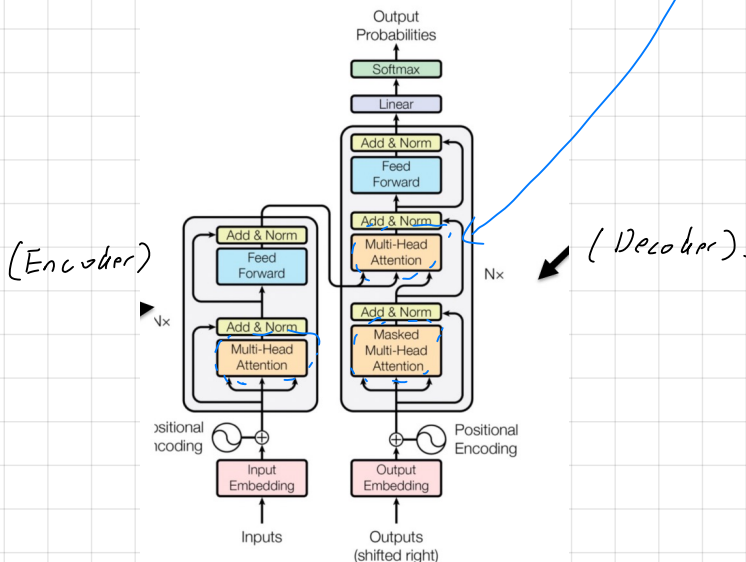
$$result: softmax(K Q) \times V$$

Multi-Head Attention



multi-head attention allowing the network to pay attention to multiple locations in the input at once.

Full Transformer Architecture

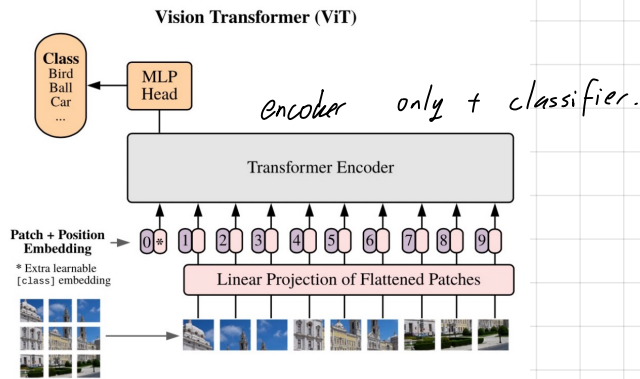


Generative Pre-Training 2.0 (GPT 2.0)

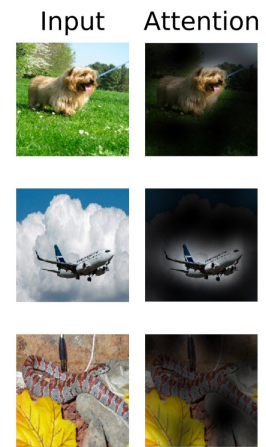
↳ trained to predict next word in a massive dataset.

Image Transformer

• Takes in patches of image (no convolution)



what is it learning?



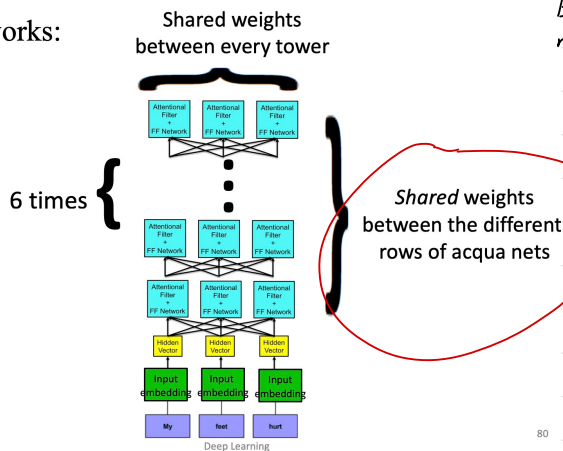
Universal Transformer Network:

↳ shared weights between different layers

(unrolled recurrent network - same computation is performed over and over).

Universal Transformer

Networks:

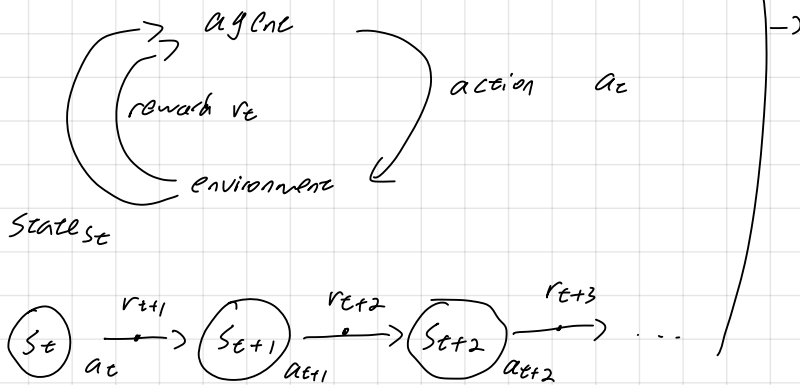


Each tower can decide adaptively how many iterations to run.

between layers.

RL, Deep Nee and Atari

Learning to act in a world.



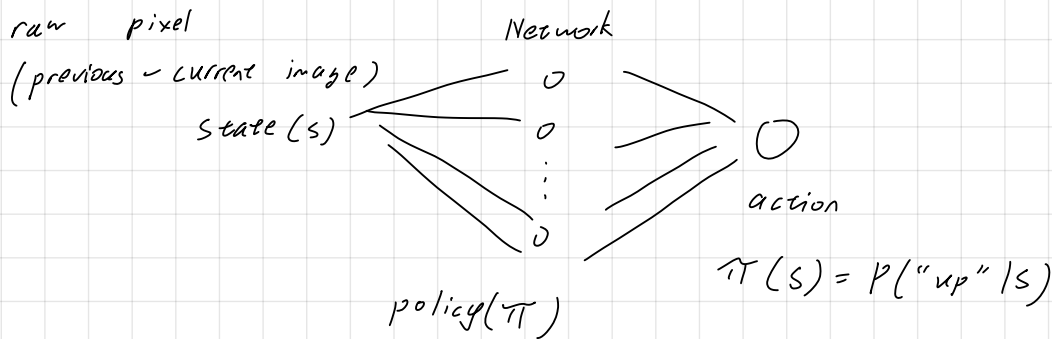
Agent learns a policy π

$$\pi(s) = P(a_1, a_2, \dots, a_n | s)$$

function from states to action probabilities.

↓
"raise prob. of action a_i when it was a good move, lower it if it was a bad move"

ex) PONG policy network



At the end, we use sign of the reward on the gradient

and apply it to every gradient we produced over the whole game.

$\Rightarrow$ win $\rightarrow$ encourage actions that led to win

lose $\rightarrow$ discourage actions that led to loss.

Policy gradient:

① At each step of play, sample from the softmax distribution at output:

$$\pi(s) = (0.1, 0.02, 0.6, \dots, 0.01)$$

↑
network, mapping from states to probabilities of action

② Treat the sample as "teacher"

$$e = (0, 0, 1, 0, \dots, 0)$$

for state/action pair.

③ Compute weight change, add them to a running average of weight change. (some form of "gradient")

④ Multiply weight change by the sign of reward.

⑤ Change the weights after one game.

TD - Gammon and Alpha Go

goal of learning:

maximizing long-term expected reward.

Let r_t be reward at time t .

$$\hookrightarrow \text{maximize } R_t = r_{t+1} + \gamma r_{t+2} + \gamma^2 r_{t+3} + \dots = \sum_{k=0}^{\infty} \gamma^k r_{t+k+1}$$

$0 \leq \gamma \leq 1$ is the discount rate.



ensures that expected reward converges.

Policy: $\pi_t(s, a) = \text{prob. that } a_t = a \text{ when } s_t = s.$

State should satisfy Markov Property:

Model-based and Model-free.

Value Function:

$V^\pi(s)$ = expected long-term return of being in this state, following policy π .

TD - Gammon.

- first application of learning value function using NN function approximator.
- representation of board as input
- learned value of board position as output.
- weighted update using temporal difference rule on every move.

↓

current estimate updated to be closer to next board's value.

TD rule

$$w_{t+1} - w_t = d(Y_{t+1} - \hat{Y}_t) \sum_{k=1}^t \lambda^{t-k} \nabla_w \hat{Y}_k$$

↑
expected value of board position.
(will be minimized when $Y_{t+1} - \hat{Y}_t = 0$).

← control temporal credit assignment of how much of an error detection at a given time step feeds back to previous estimate.

Alpha - Go

- Train two policy networks by supervised training on expert positions.

one is shallow and fast

↓
used for rollouts

one is deep and accurate

↓
train another network by policy gradient method.

↓
After supervised training, network is improved by playing itself.

A 3rd network: value network trained to predict winner from every state.

↓
from every state seen during play, train to predict win rate.

Masked Autoencoder

masking then

reconstruct input in pre-training.

↓
After pre-training, use a simple decoder.

↓
achieve good
results with
minimum decoder

BERT

Step ①: break image with non-overlapping patches.

②: mask 75% of patches.

③: positional encoding of the patches.

④: learns to reconstruct the whole image.

- self-supervised learning without any examples.
- self-supervised learning with few data augmentation.

AGI (advanced ML/AI)

- ↓
- agents learn as much as they can about the world without interaction (by observation).
 - differentiable architectures for planning and reasoning
 - learn multi-level abstraction through long-term prediction and long-term planning.