

Regularization & Dropout

① Regularization

$$\hat{L}(\theta) = L(\theta) + \lambda \|\theta\|_2^2$$

② Bagging: Bootstrap Aggregating (Dropout)

Reduce Variance

In FNN: **Dropout!**

At each training step

each hidden unit or input

has probability p of being dropped out

Implementation: $\hat{h}_i = \mu_{h_i} \cdot h_i$

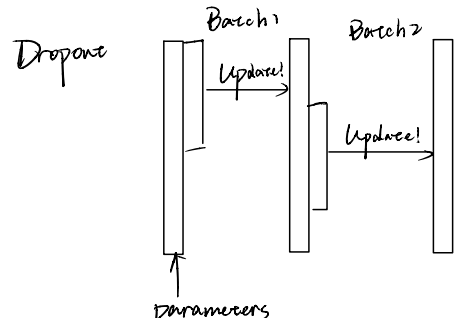
add a binary mask variable: $\mu_{h_i} = \begin{cases} 0, & \text{probability } p \\ 1, & 1-p \end{cases}$

For each minibatch:

$\mu = 1 \Rightarrow$ update weight

$\mu = 0 \Rightarrow$ no update

$\Rightarrow$ training is conducted in a part of network
Meanwhile: still only one network!



Why "Dropout" is like "Bagging"

• 每次随机丢弃神经元，就是在构建一个弱的模型结构。
对于每一个有 N 个神经元的网络，随机丢弃 Dropout ，理论上可以产生 2^N 种不同的网络结构（因为每个神经元可以被保留或丢弃），这相当于 Bagging 中的“多个不同的模型”。
• 由于模型太多，训练模型会在不同模型间产生平均。
• 虽然 Bagging 要求不同的数据子集来训练不同的模型，但 Dropout 是在同一个 batch 的数据上，通过不同的网络结构进行训练，虽然数据相同，但因为网络结构变了，数据流经的路径不同，相当于模型看到了数据的不同“特征组合”。
• 训练时模型会随机，模型在多个子集间进行平均。
在 Bagging 中，测试时我们会对所有模型的输出取平均。在 Dropout 中，测试时我们不再随机丢弃，而是使用完整的网络。为了补偿训练时的随机丢弃，我们将权重乘以保留概率 p ，这可以理解为训练时每个神经元有概率 p 个神经元进行“几倍平均”的效应。
核心思想：Dropout 并不是构建一个巨大的网络，而是在训练过程中，随机丢弃部分子网络共享权重，训练出来的模型在测试时才能取平均。

Optimized Algorithms

Task: $\hat{\theta} = \arg \min_{\theta} L(\theta)$ (always)

Steps: 1. init θ 2. dataset θ to output 3. update θ
Then till meet criterion

Different Optimizer:

① Momentum:

$\bar{\theta}$ has speed $\bar{v}$ at each iteration:

1. current speed reduced, $\bar{v} \leftarrow \mu \bar{v}$

2. acceleration: $-e \bar{g}$

3. $\bar{v} \leftarrow \mu \bar{v} - e \bar{g}$

4. $\theta \leftarrow \theta + \bar{v}$

“记住前次信息，又注意当前参数变化”
(动量)

$$V_t = -\eta (\nabla L_t + \mu \nabla L_{t-1} + \mu^2 \nabla L_{t-2} \dots)$$

learning rate

Momentum ① reduces the variation in overall gradient direction.

② speed up learning.

Algorithm with adaptive learning rate

Parameter that have changed a lot should be allowed to change less in future

often changed $\rightarrow$ maybe already close to optimized value
 $\rightarrow$ fixed rate $\rightarrow$ can not converge
 $\rightarrow$ adaptive rate $\rightarrow$ converge!

seldom changed $\rightarrow$ if learning rate = rate (often changed) (i.e. fixed learning rate) then too slow! $\rightarrow$ needs to speed up

Hadamard product / division

$\Rightarrow$ elementwise

①: product ②: division

① AdaGrad:

① Initialize gradient accumulation variable $\bar{g} = 0$

② While Stop criterion not met do

1. Sample a minibatch of m samples from the training set $\{x^{(1)}, x^{(2)}, \dots, x^{(m)}\}$ with corresponding targets $y^{(1)}, y^{(2)}, \dots, y^{(m)}$

2. Compute Gradient: $g \leftarrow \frac{1}{m} \nabla \sum_{i=1}^m L(f(x^{(i)}; \theta), y^{(i)})$ Gradient

3. "re-re-goes" (Accumulate squared gradients)

4. $\Delta \theta \leftarrow -\frac{e}{\sqrt{\bar{g} + \delta}} g$ "learning rate"

$\bar{g} + \delta$ "Gradient" accumulation
small constant to make sure $\sqrt{\bar{g} + \delta} \neq 0$

5. $\theta \leftarrow \theta + \Delta \theta$

After t iterations:
 $\bar{g} = g_1^2 + g_2^2 + \dots + g_t^2$

② RMSProp

$$\bar{r} \leftarrow \rho \bar{r} + (1-\rho) \bar{g}^2$$

After t iterations:

$$\bar{r} = (1-\rho) (p_1^2 g_1^2 + p_2^2 g_2^2 + \dots + p_t^2 g_t^2 + g_t^2)$$
$$\Delta \bar{\theta} = -\frac{e \bar{g}}{\sqrt{\bar{r} + \delta}}$$

"Exponentially Decaying weight"

③ Adam: Adaptive Moment

RMSProp:

Momentum "Combination"

Bias Correction

Key Procedure:

$s \leftarrow \rho s + (1-\rho) g$ Momentum 1st Moment

$r \leftarrow \rho r + (1-\rho) g \otimes g$ RMSProp 2nd Moment

Correct bias in first moment: $\hat{s} \leftarrow \frac{s}{1-\rho^t}$

Correct bias in second moment: $\hat{r} \leftarrow \frac{r}{1-\rho^t}$

Where t is the count of iterations

$$\Delta \theta = -e \frac{\hat{s}}{\sqrt{\hat{r} + \delta}}$$

$$\theta \leftarrow \theta + \Delta \theta$$

Why doing This?

take s as example:

$$s = (1-\rho) (p_1^t g_1 + p_1^{t-1} g_2 + p_1^{t-2} g_3 + \dots + p_1 g_{t-1} + g_t)$$

$$\star 1-\rho^t = (1-\rho) (p_1^t + p_1^{t-1} + p_1^{t-2} + \dots + 1)$$

$$\hookrightarrow \text{"Geometric Series!" } \sum_{i=1}^t (p_1^i) = \frac{1-p_1^t}{1-p_1}$$

$\therefore$ Bias Correction:

$$\hat{s} = \frac{s}{1-\rho^t} = \frac{(1-\rho) (p_1^t g_1 + p_1^{t-1} g_2 + p_1^{t-2} g_3 + \dots + p_1 g_{t-1} + g_t)}{(1-\rho) (p_1^t + p_1^{t-1} + p_1^{t-2} + \dots + 1)}$$

$$= \frac{(p_1^t g_1 + p_1^{t-1} g_2 + p_1^{t-2} g_3 + \dots + p_1 g_{t-1} + g_t)}{(p_1^t + p_1^{t-1} + \dots + 1)}$$

$$= \frac{\frac{p_1^t}{2} g_1 + \frac{p_1^{t-1}}{2} g_2 + \dots + \frac{1}{2} g_t}{\frac{p_1^t}{2} + \frac{p_1^{t-1}}{2} + \dots + \frac{1}{2}}$$

Sum is 1! $\rightarrow$ "1/2 - 1/2"

A Probability Distribution Over time!

"Momentum"

$$E(x) = \sum_{i=1}^k x_i \cdot P(x_i)$$
$$s = \sum_{i=1}^k g_i \cdot P(g_i)$$

Same!

Through Momentum & RMSProp

We have $E(x)$ and $E(x^2)$

Then $\text{Var}(x) = E(x^2) - [E(x)]^2$

Moment/RMSProp: Create the distribution (with bias)

"记住前次信息，又注意当前参数变化"

Bias correction: $\frac{1}{1-\rho^t}$. Avoid: 1st moment too small $\rightarrow$ slow update at beginning
2nd moment too small $\rightarrow$ Too large learning rate!

$$\Delta \theta \leftarrow -\frac{e \hat{s}}{\sqrt{\hat{r} + \delta}}$$

AdaGrad:

After t iterations:

$$\bar{r} = g_1^2 + g_2^2 + \dots + g_t^2$$

"No Decaying!"

Compare to

Exponentially Decaying weight

Weight Decay & L2 regularization

L2 regularization: $\text{Loss} = \text{Loss}_{\text{original}} + \frac{\lambda}{2} \|\theta\|_2^2$ Weight penalty

Weight penalty $\rightarrow$ smaller model capacity $\rightarrow$ avoid overfitting

Weight update:

$$\nabla_{\theta} L = \nabla_{\theta} L + \lambda \theta$$

$$\theta_{t+1} = \theta_t - \eta (\nabla_{\theta} L + \lambda \theta_t)$$

$$= (1-\eta \lambda) \theta_t - \eta \nabla_{\theta} L$$

Weight decay:

$$\theta_{t+1} = (1-\eta \lambda) \theta_t - \eta \nabla_{\theta} L$$

L2 regularization and Weight Decay are equivalent for SGD

But They are not equivalent for Adam!

AdamW: Adam with decoupled Weight decay

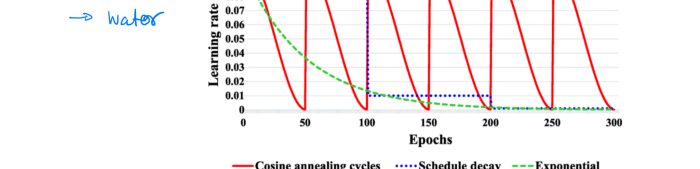
$$g_t = \nabla L + \lambda \theta_{t-1}$$

Learning Rate rescheduling

Learning Rate Scheduling: adaptive w/ time

Momentum and Adam are about the use of gradients.

Here are some ways to schedule the learning rate:



Code: Annealing. The resetting of the learning rate acts like a simulated restart of the learning process and the reuse of good weights as the starting point of the restart is referred to as a "warm restart".

Batch normalization

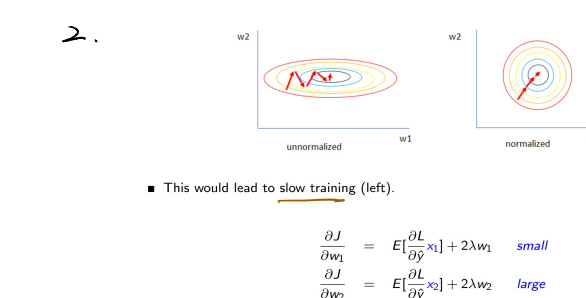
if $x_1 \in [0, 1]$, $x_2 \in [10, 100]$

$$\hat{y} = w_0 + x_1 w_1 + x_2 w_2$$

$$J(w) = E[L(y, \hat{y})] + \lambda (w_1^2 + w_2^2)$$

"Loss Function with L2 regularization"

Problem: 1. regularization affect w_1 more than w_2



Data Normalization

$$\hat{x}_{ij} \leftarrow \frac{x_{ij} - \mu_j}{\sigma_j}$$

In deep model!

if we only normalize x .

x layer m_1 , m_1 's distribution cannot be guaranteed!

different layers becomes dependent / Covariate shift

We're going to normalize each layer

for each mini-batch

Layer n

$x \rightarrow \hat{x} = \frac{x - \mu}{\sigma}$

Normalization

$x \rightarrow \hat{x} \sim N(0, 1)$

is that enough? $\rightarrow$ lack of expression!

Learnable variables

Layer $n+1$

Scaling & shifting

make data more expressive!

Remarks:

1. Batch normalization can reduce "Vanishing Gradient" Problem

2. Batch normalization has different computation results

in training and prediction (i.e. testing) modes. (Since datasets are not the same!)