

Non convex cost functions have multiple optima, only one of which is the global optimum. The cost functions used in deep learning are mostly non-convex and can be very difficult to converge to the global optimum. There are multiple challenges associated with training a neural network

- Learning rate selection
- Avoiding poor local minimas
- Avoiding drastic changes to the loss surface

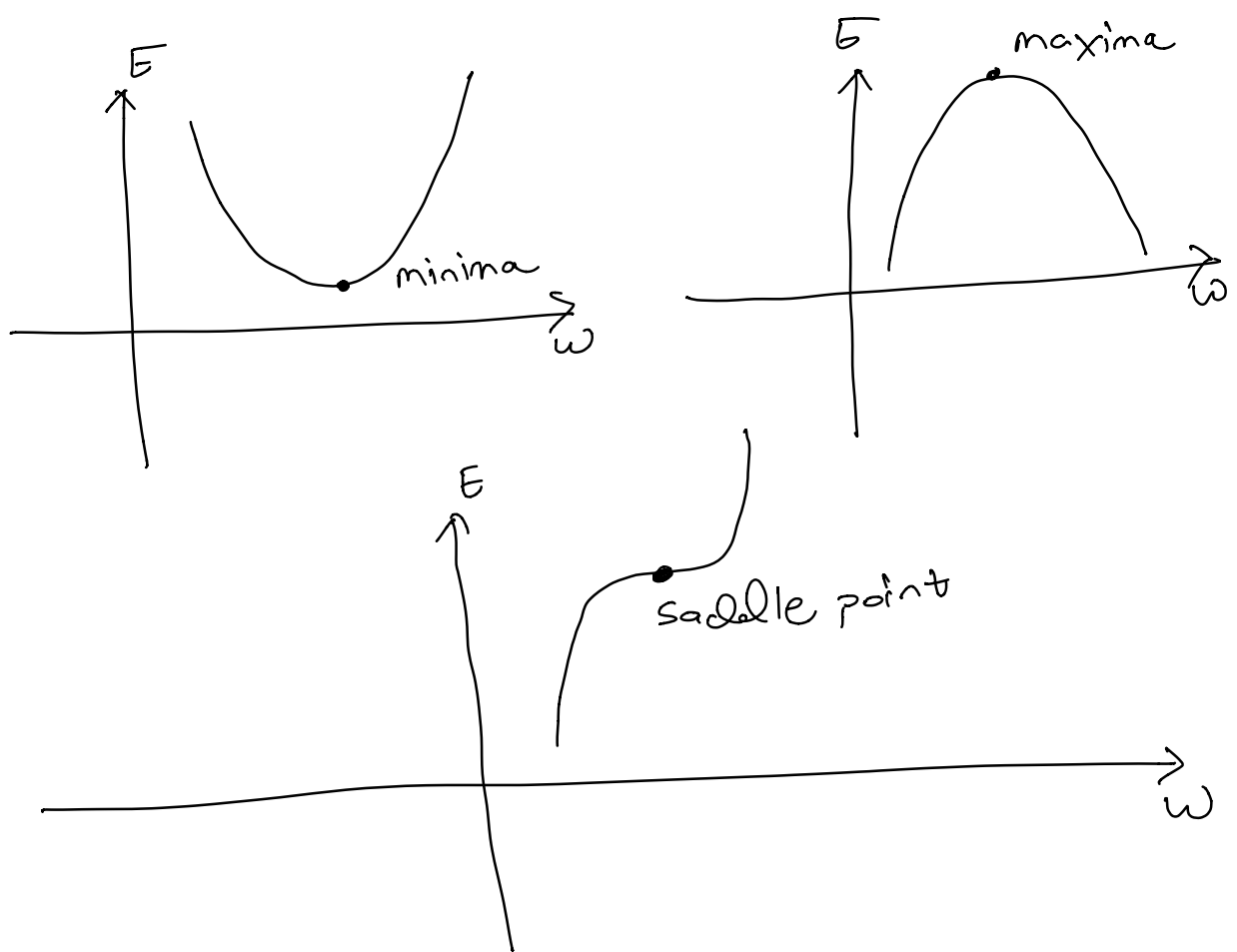
Local minima in deep networks:

Till today it remains an open question whether local minima with a high error rate relative to the global minima are common in practical deep networks.

However, many recent studies seem to indicate that most local minima have error rates and generalization characteristics that are very similar to global minima.

Critical points:

Given an arbitrary function, a point at which the gradient is the zero vector is called a critical point.



For a 1-D cost function, a critical point can take one of the three forms as shown above. Assuming each of these 3 configurations is equally likely then the probability of

a randomly selected critical point being a minima is $1/3$. Hence if we have a total of K critical points then on average we will have a total of $\frac{K}{3}$ minimas.

Now let's consider a cost function in a d -dimensional space. In general, in a d -dimensional space we can slice through a critical point in d -different axes. A critical point can only be a local minima if it appears as a local minima in every single one of the

d one-dimensional subspaces. Hence using the result from earlier we have that the probability that a randomly selected critical point in a d -dimensional space is a local minima is $\frac{1}{3^d}$. $\therefore$ As d increases local minima become exponentially more rare.

Momentum:

In lecture, we observed that loss surfaces with high curvature results in oscillations of the SGD updates. The oscillatory nature leads to a slower convergence of the SGD.

Since SGD will move in the right direction but with oscillations so in order to dampen out the oscillations we will use the average gradient to make the updates.

$$v \leftarrow \alpha v - \epsilon g$$

$$\theta \leftarrow \theta + v$$

v is the velocity and more weight is applied to more recent gradients

creating an exponentially decaying average of gradients.

Adaptive gradients:

A major challenge for training deep networks is appropriately selecting the learning rate. The basic concept behind learning rate adaptation is that the optimal learning rate is appropriately modified over the span of learning to achieve good convergence properties.

Adagrad:

- Adapt global learning rate over time using an accumulation of historical gradients.
- Keep track of a learning rate for each parameter.

$$a \leftarrow a + g \odot g$$

a is the gradient accumulation vector

$$\theta \leftarrow \theta - \frac{\epsilon}{\sqrt{a} + \epsilon} \odot g$$

- Parameters with largest gradients experience a rapid decrease in their learning rates while parameters with smaller gradients only observe a small decrease in their learning rates.

- Since $g \odot g \geq 0$, α keeps growing and in turn causes the learning rate to shrink and eventually become infinitesimally small.

- Flat regions may force AdaGrad to decrease learning rate before it reaches a minima.

RMS Prop:

- Use an exponentially weighted moving average of gradients

$$a \leftarrow \beta a + (1-\beta) g \odot g$$

Where β is the decay factor
Smaller the decay factor the
shorter the effective window

Adam:

- A variant combination of RMSProp and momentum

- Keep track of an exponentially weighted moving average of the gradient

$$v \leftarrow \beta_1 v + (1 - \beta_1) g$$

- Keep track of an exponentially weighted moving average of historical gradients.

$$a \leftarrow \beta_2 a + (1 - \beta_2) g \odot g$$

Practice Problems:

1

$$h^1 = Wx$$

$$= \begin{bmatrix} a & 0 \\ 0 & b \end{bmatrix} x$$

$$h^2 = Wh^1$$

$$= WWx = \begin{bmatrix} a^2 & 0 \\ 0 & b \end{bmatrix} x$$

In general,

$$h^n = W^n x$$

$$= \begin{bmatrix} a^n & 0 \\ 0 & b^n \end{bmatrix} x$$

$$\text{so, } h_1^n = a^n x_1$$

$$h_2^n = b^n x_2$$

Now, $y = h_1^n + h_2^n$

$\Rightarrow y = a^n x_1 + b^n x_2$

now, $\frac{\partial y}{\partial a} = n a^{n-1} x_1$

$\frac{\partial y}{\partial b} = n b^{n-1} x_2$

So, $\nabla y = \begin{bmatrix} n a^{n-1} x_1 \\ n b^{n-1} x_2 \end{bmatrix}$

now,
Assuming $x_1 = x_2 = 1$ and $a = 1, b = 2$, we

have

$y = 1 + 2^n$

$\nabla y = \begin{bmatrix} n \\ n 2^{n-1} \end{bmatrix}$

As $n \rightarrow \infty$ then ∇y explodes

Assuming $x_1 = x_2 = 1$ and $a = 0.5$, $b = 0.9$

we have

$$y = (0.5)^n + (0.9)^n$$

$$\nabla y = \begin{bmatrix} n(0.5)^{n-1} \\ n(0.9)^{n-1} \end{bmatrix}$$

As $n \rightarrow \infty$ then ∇y vanishes.

2

From problem statement, we have the recursive relation

$$V_t = (1 - \beta_1) \sum_{i=1}^t \beta_1^{t-i} g_i^*$$

Let's do a change of variable in the above summation:

$$j = t - i$$

Then,

$$i=1 \rightarrow j=t-1$$

$$i=t \rightarrow j=0$$

$$\beta_1^{t-i} \rightarrow \beta_1^j$$

$$g_i^* \rightarrow g_{t-j}$$

Hence,

$$V_t = (1 - \beta_1) \sum_{j=0}^{t-1} \beta_1^j g_{t-j}$$

Now, taking expectation of both sides we get

$$E[V_t] = E\left[(1 - \beta_1) \sum_{j=0}^{t-1} \beta_1^j g_{t-j}\right]$$

Since $E[\cdot]$ is a linear operator, so

$$E[V_t] = (1 - \beta_1) \sum_{j=0}^{t-1} E[\beta_1^j g_{t-j}]$$

Since β_1^j is a deterministic quantity,

so

$$E[V_t] = (1 - \beta_1) \sum_{j=0}^{t-1} \beta_1^j E[g_{t-j}]$$

From problem statement,

$$E[g_{t-j}] = \mu, \quad j=0, \dots, t-1$$

$$\text{So, } E[V_t] = (1-\beta_1) \sum_{j=0}^{t-1} \beta_1^j M$$

$$= (1-\beta_1) M \sum_{j=0}^{t-1} \beta_1^j$$

Recall,
 $1, \beta_1, \beta_1^2, \beta_1^3, \dots, \beta_1^{t-1}$

is a geometric series with
 $a=1$ and $r=\beta_1$. Then using
 the summation of geometric series

$$\sum_{j=0}^{t-1} \beta_1^j = \frac{1-\beta_1^t}{1-\beta_1}$$

Then,

$$E[v_t] = (1-\beta_1)\mu \frac{(1-\beta_1^t)}{1-\beta_1}$$

$$\Rightarrow E[v_t] = \mu (1-\beta_1^t)$$

Hence,

$$\frac{1}{1-\beta_1^t} E[v_t] = \mu = E[g_t]$$

Similarly,

$$a_t = (1-\beta_2) \sum_{i=1}^t \beta_2^{t-i} g_i^2$$

Doing a change of variables
like before

$$a_t = (1-\beta_2) \sum_{j=0}^{t-1} \beta_2^j g_{t-j}^2$$

Taking $E[\cdot]$ of both sides and simplifying

$$E[a_t] = (1-\beta_2) \sum_{j=0}^{t-1} \beta_2^j E[g_{t-j}^2]$$

since, $E[g_{t-j}^2] = S$

so,

$$E[a_t] = (1-\beta_2) S \sum_{j=0}^{t-1} \beta_2^j$$

Using the summation of geometric series,

$$E[a_t] = (1-\beta_2) S \cdot \frac{(1-\beta_2^t)}{1-\beta_2}$$

$$E[a_t] = (1-\beta_2^t) S$$

Hence,

$$\frac{1}{1-\beta_2^t} E[a_t] = s = E[g_t^2]$$

3

Recall that for $A \in \mathbb{R}^{m \times n}$, its infinity norm is defined as

$$\|A\|_{\infty} = \max_i \sum_{j=1}^n |a_{ij}|$$

Observe that,

$$\sum_{j=1}^n |a_{ij}| \leftarrow \begin{array}{l} \text{Sum of the} \\ \text{absolute values} \\ \text{of elements in} \\ \text{ith row} \end{array}$$

Hence, infinity norm of a matrix is the maximum absolute row sum.

Let

$$A = \begin{pmatrix} 1 & 2 & -4 \\ 3 & 0 & 12 \\ -20 & -1 & 2 \end{pmatrix}$$

Then,

$$\|A\|_{\infty} = \max(\{1+2+4\}, \{3+12\}, \{20+1+2\})$$

$$= 23$$

Now, coming back to our problem

$$W = \begin{bmatrix} w_{11} & w_{12} & \dots & w_{1n} \\ w_{21} & w_{22} & \dots & w_{2n} \\ \vdots & \vdots & & \vdots \\ \vdots & \vdots & & \vdots \\ w_{n1} & w_{n2} & \dots & w_{nn} \end{bmatrix}$$

$$\frac{\partial \|W\|_{\infty}}{\partial w} = \begin{bmatrix} \frac{\partial \|W\|_{\infty}}{\partial w_{11}} & \dots & \frac{\partial \|W\|_{\infty}}{\partial w_{1n}} \\ \vdots & & \vdots \\ \frac{\partial \|W\|_{\infty}}{\partial w_{n1}} & \dots & \frac{\partial \|W\|_{\infty}}{\partial w_{nn}} \end{bmatrix}$$

Suppose, I tell you that the 3rd row of W has the maximum absolute row sum. Then,

$$\begin{aligned} \|W\|_{\infty} &= \sum_{j=1}^n |w_{3j}| \\ &= |w_{31}| + |w_{32}| + \dots + |w_{3n}| \end{aligned}$$

Then,

$$\frac{\partial \|W\|_1}{\partial W} = \begin{bmatrix} 0 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \text{sign}(w_{31}) & \text{sign}(w_{32}) & \dots & \text{sign}(w_{3n}) \\ 0 & 0 & \dots & 0 \\ \vdots & & & \\ 0 & 0 & \dots & 0 \end{bmatrix}$$

Then, from the above pattern we can conclude that

$$\frac{\partial \|W\|_1}{\partial W} = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ \text{sign}(w_k) \\ 0 \\ \vdots \\ 0 \end{bmatrix}$$

where k is the index of the row with maximum absolute sum