

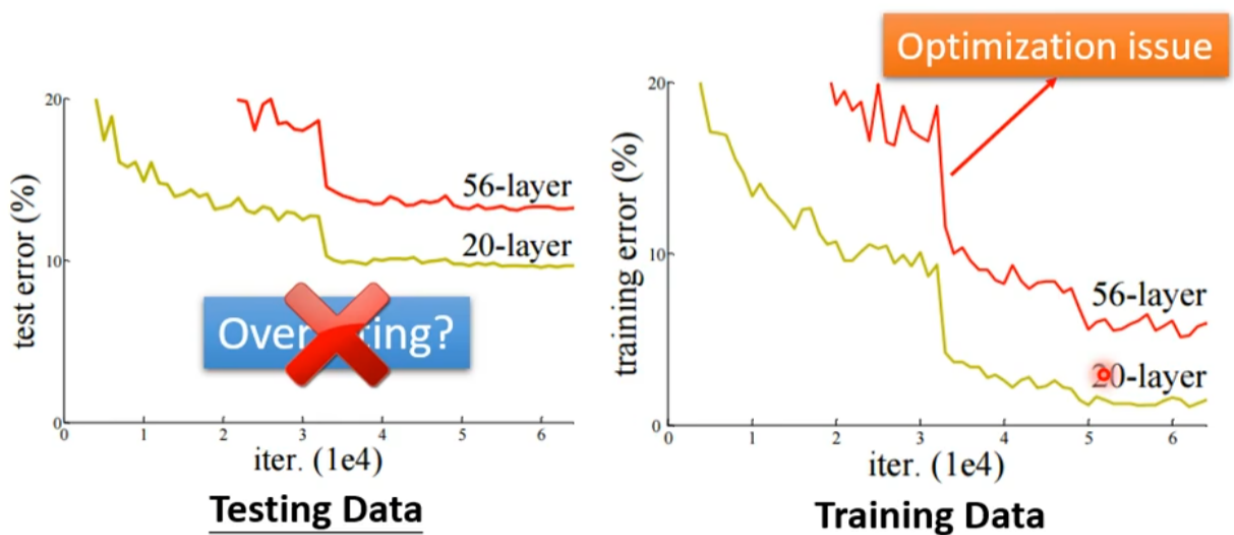
训练模型攻略

回顾：

- step1: 定义含有未知参数函数 $y = f_{\theta}(x)$
- step2: 从训练集中定义损失 $L(\theta)$
- step3: 优化 $\theta^* = \operatorname{argmin}_{\theta} L$

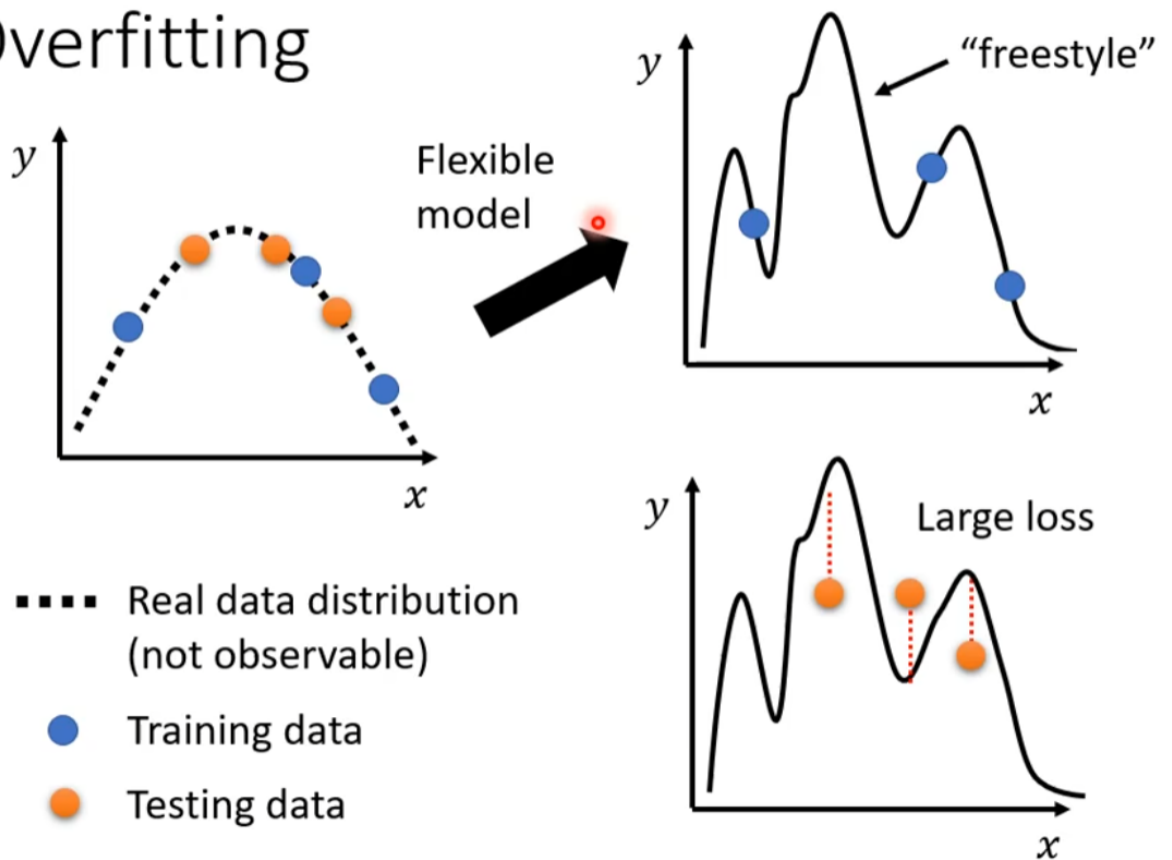
模型训练策略

- 训练集的损失过大
 - 原模型无法描述问题 -> **复杂化模型** (增加更多特征 x , 增加弹性 (激活函数))
 - **优化问题 (optimization issue)** : 没有找到全局最好
 - 更多的层数堆叠引入了更多的未知数, 理论上深的模型可以做到浅模型的功能, 但更多参数optimization也会更难 (**判断: 通过深浅的模型对比**, 如果深的表现更差就存在 optimization issue)



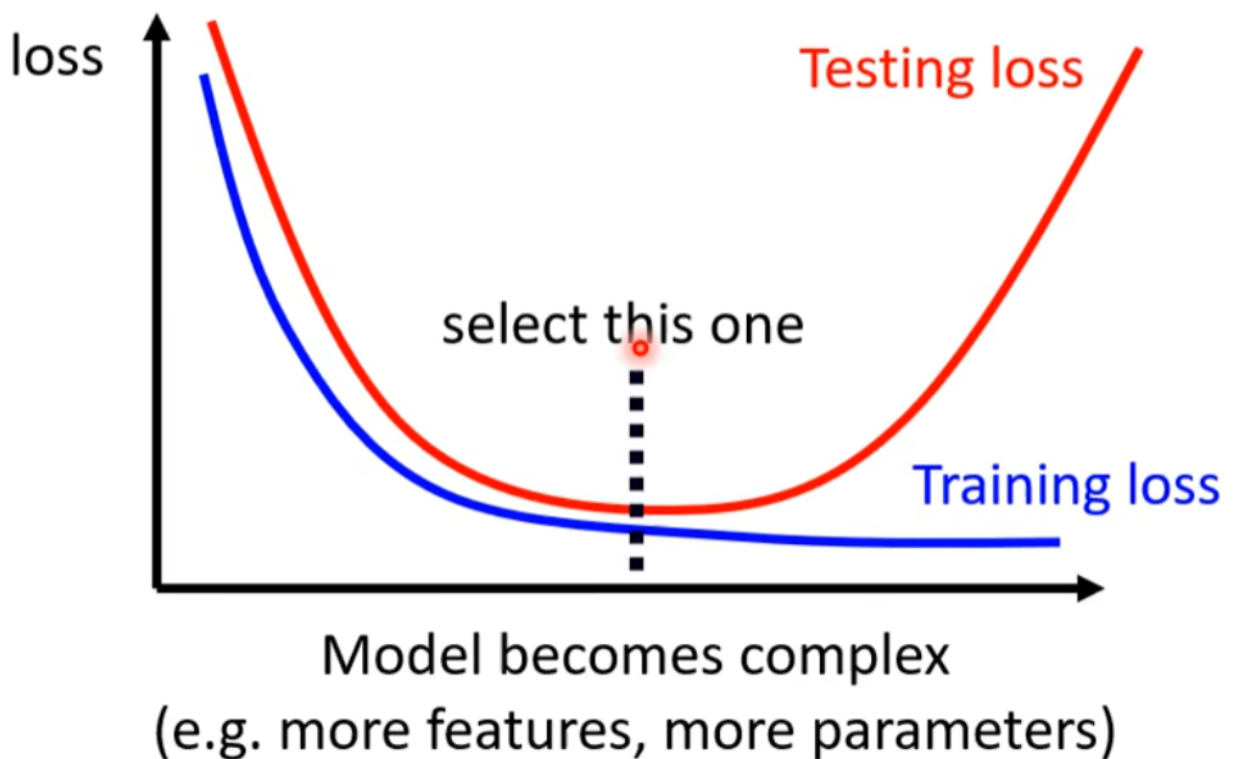
- 训练集的损失小 -> 测试集的损失大
 - **过拟合**

Overfitting



○

- 解决方法：增加数据量，数据增强，减少弹性适当限制模型（更少的参数/共享参数，更少的特征，早停，正则，Dropout）



○

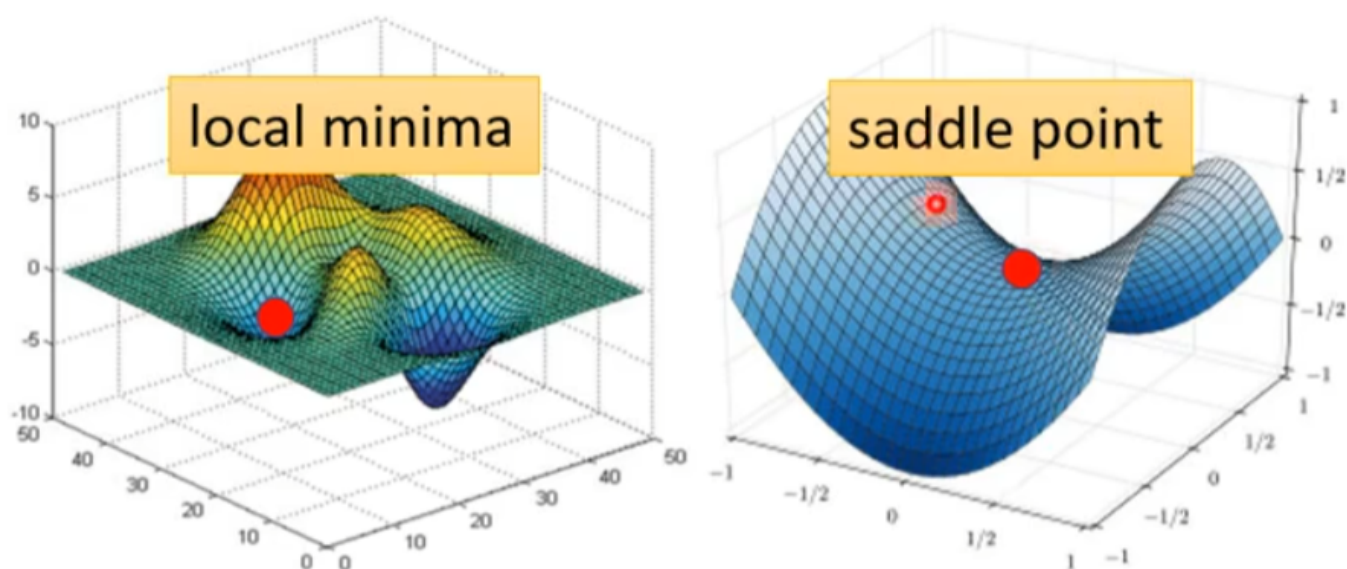
- 如何选择模型 -> 交叉验证：训练集、验证集、测试集

- mismatch: 训练集和测试集是不同的分布

如何解决优化问题

存在问题：当梯度为0，不再更新了，训练停止 -> 局部问题。

梯度为0的情况（零界点）：（1）局部极小值，周边没办法走 （2）鞍部（局部极小值+极大值），可以往两边走



如果是saddle point还可以梯度更新，那如何根据数学判断两种情况：

Taylor Series Approximation

$L(\theta)$ around $\theta = \theta'$ can be approximated below

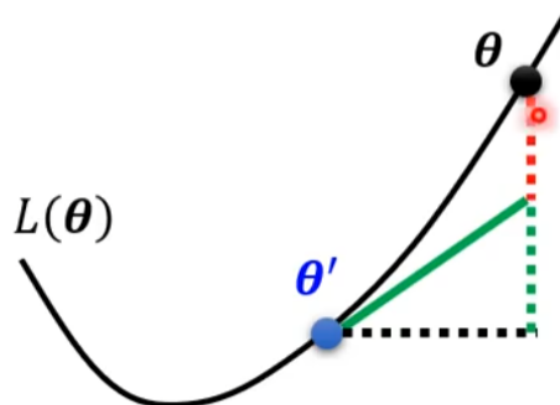
$$L(\theta) \approx L(\theta') + (\theta - \theta')^T \mathbf{g} + \frac{1}{2} (\theta - \theta')^T \mathbf{H} (\theta - \theta')$$

Gradient $\mathbf{g}$ is a vector

$$\mathbf{g} = \nabla L(\theta') \quad g_i = \frac{\partial L(\theta')}{\partial \theta_i}$$

Hessian $\mathbf{H}$ is a matrix

$$H_{ij} = \frac{\partial^2}{\partial \theta_i \partial \theta_j} L(\theta')$$



当走到临界点时 $\Leftrightarrow$ 梯度 $g = 0$, 此时用 H 判断临界点种类:

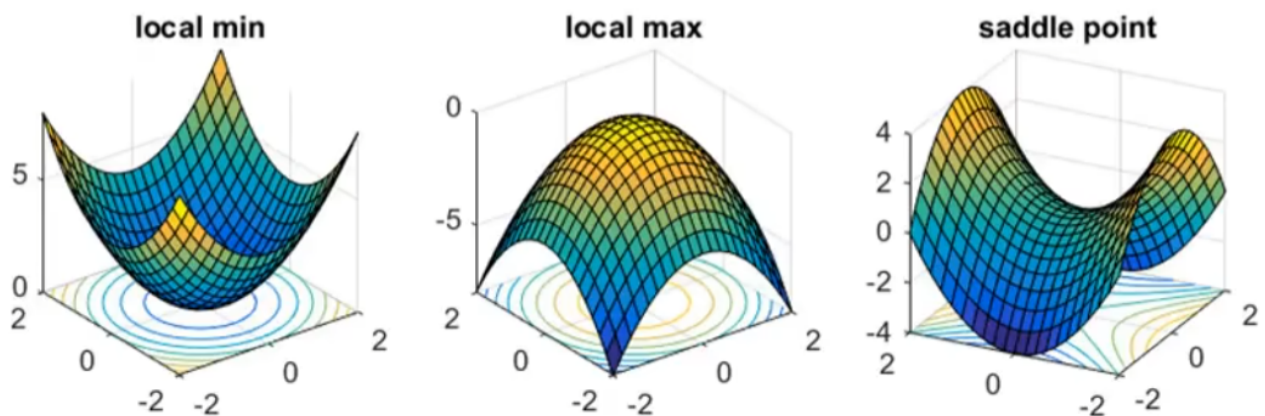
Hessian

$L(\theta)$ around $\theta = \theta'$ can be approximated below

$$L(\theta) \approx L(\theta') + \cancel{(\theta - \theta')^T g} + \frac{1}{2} (\theta - \theta')^T H (\theta - \theta')$$

At critical point

telling the properties of critical points



令 $v = \theta - \theta'$, 我们可以根据推论用 H 的值判断:

For all v

$v^T H v > 0 \implies \text{Around } \theta': L(\theta) > L(\theta') \implies \text{Local minima}$
 $= H \text{ is positive definite} = \text{All eigen values are positive.} \uparrow$

For all v

$v^T H v < 0 \implies \text{Around } \theta': L(\theta) < L(\theta') \implies \text{Local maxima}$
 $= H \text{ is negative definite} = \text{All eigen values are negative.} \uparrow$

Sometimes $v^T H v > 0$, sometimes $v^T H v < 0 \implies \text{Saddle point}$
 Some eigen values are positive, and some are negative. $\uparrow$

可以找一个向量 λ 让 L 变小

Don't afraid of saddle point?

$$\mathbf{v}^T \mathbf{H} \mathbf{v}$$

At critical point: $L(\boldsymbol{\theta}) \approx L(\boldsymbol{\theta}') + \frac{1}{2} (\boldsymbol{\theta} - \boldsymbol{\theta}')^T \mathbf{H} (\boldsymbol{\theta} - \boldsymbol{\theta}')$

Sometimes $\mathbf{v}^T \mathbf{H} \mathbf{v} > 0$, sometimes $\mathbf{v}^T \mathbf{H} \mathbf{v} < 0 \Rightarrow$ Saddle point

$\mathbf{H}$ may tell us parameter update direction!

$\mathbf{u}$ is an eigen vector of $\mathbf{H}$

λ is the eigen value of $\mathbf{u}$

$$\lambda < 0$$



$$\mathbf{u}^T \mathbf{H} \mathbf{u} = \mathbf{u}^T (\lambda \mathbf{u}) = \lambda \|\mathbf{u}\|^2$$

$$< 0$$

$$< 0$$

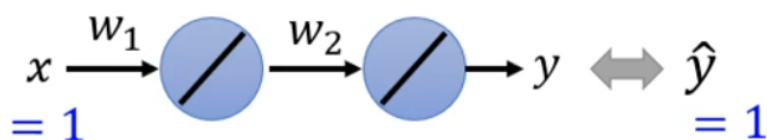
$$L(\boldsymbol{\theta}) \approx L(\boldsymbol{\theta}') + \frac{1}{2} (\boldsymbol{\theta} - \boldsymbol{\theta}')^T \mathbf{H} (\boldsymbol{\theta} - \boldsymbol{\theta}') \Rightarrow L(\boldsymbol{\theta}) < L(\boldsymbol{\theta}')$$

$$\boldsymbol{\theta} - \boldsymbol{\theta}' = \mathbf{u}$$

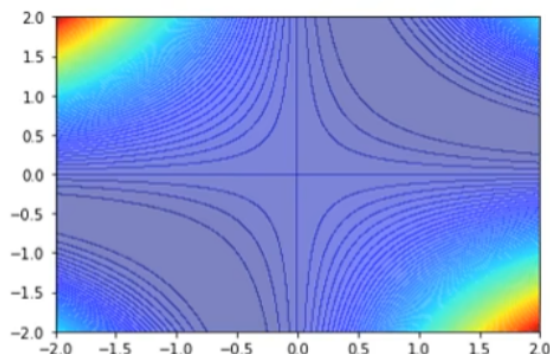
$$\boldsymbol{\theta} = \boldsymbol{\theta}' + \mathbf{u}$$

Decrease L

e.g. (计算H) : (0, 0) 点是saddle point



$$L = (\hat{y} - w_1 w_2 x)^2 = (1 - w_1 w_2)^2$$



$$\frac{\partial L}{\partial w_1} = 2(1 - w_1 w_2)(-w_2) = 0$$

$$\frac{\partial L}{\partial w_2} = 2(1 - w_1 w_2)(-w_1) = 0$$

Critical point: $w_1 = 0, w_2 = 0$

$$H = \begin{bmatrix} 0 & -2 \\ -2 & 0 \end{bmatrix} \quad \lambda_1 = 2, \lambda_2 = -2$$

Saddle point

g

H

$$\frac{\partial^2 L}{\partial w_1^2} = 2(-w_2)(-w_2) = 0$$

$$\frac{\partial^2 L}{\partial w_1 \partial w_2} = -2 + 4w_1 w_2 = -2$$

$$\frac{\partial^2 L}{\partial w_2 \partial w_1} = -2 + 4w_1 w_2 = -2$$

$$\frac{\partial^2 L}{\partial w_2^2} = 2(-w_1)(-w_1) = 0$$

$$\lambda_2 = -2 \quad \text{Has eigenvector } \mathbf{u} = \begin{bmatrix} 1 \\ 1 \end{bmatrix}$$

Update the parameter along the direction of $\mathbf{u}$

You can escape the saddle point and decrease the loss.

但这种二次微分的方法运算复杂，目前实际没有用这个方法逃离。

- 但考虑到维度的问题，路可能不止一条，**局部极小值可能很稀少**。实际中，即使趋向收敛也可能会让梯度继续下降。

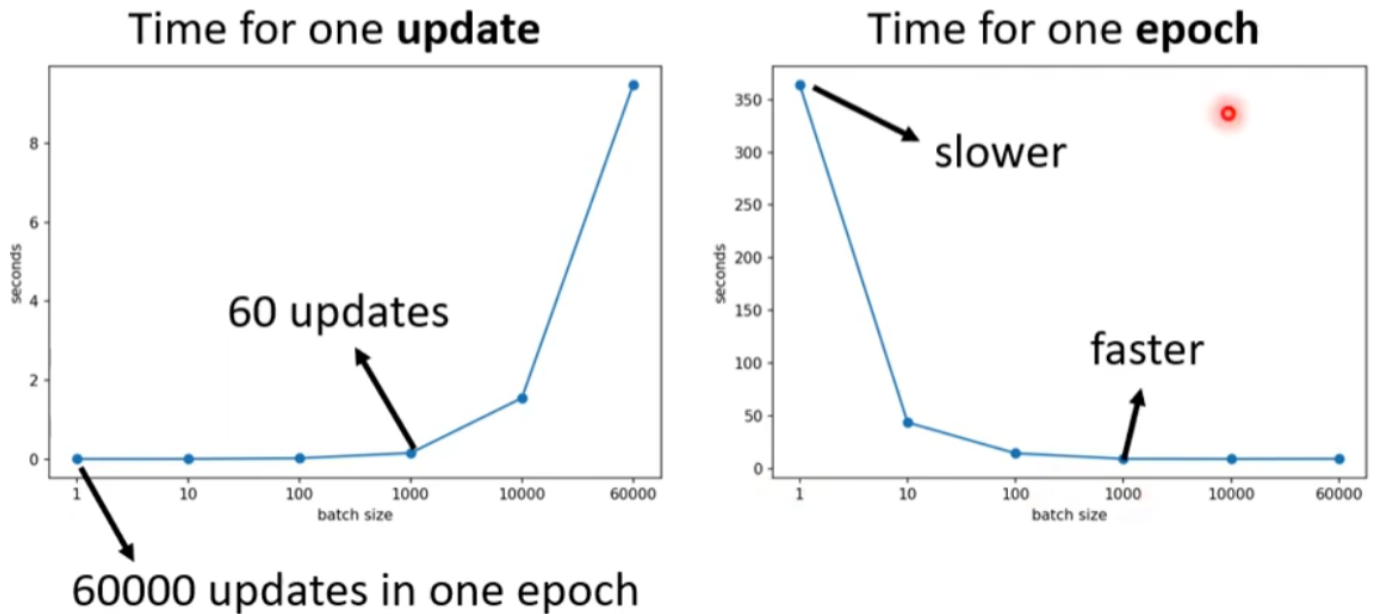
(1) Batch

在没有平行运算的情况下：

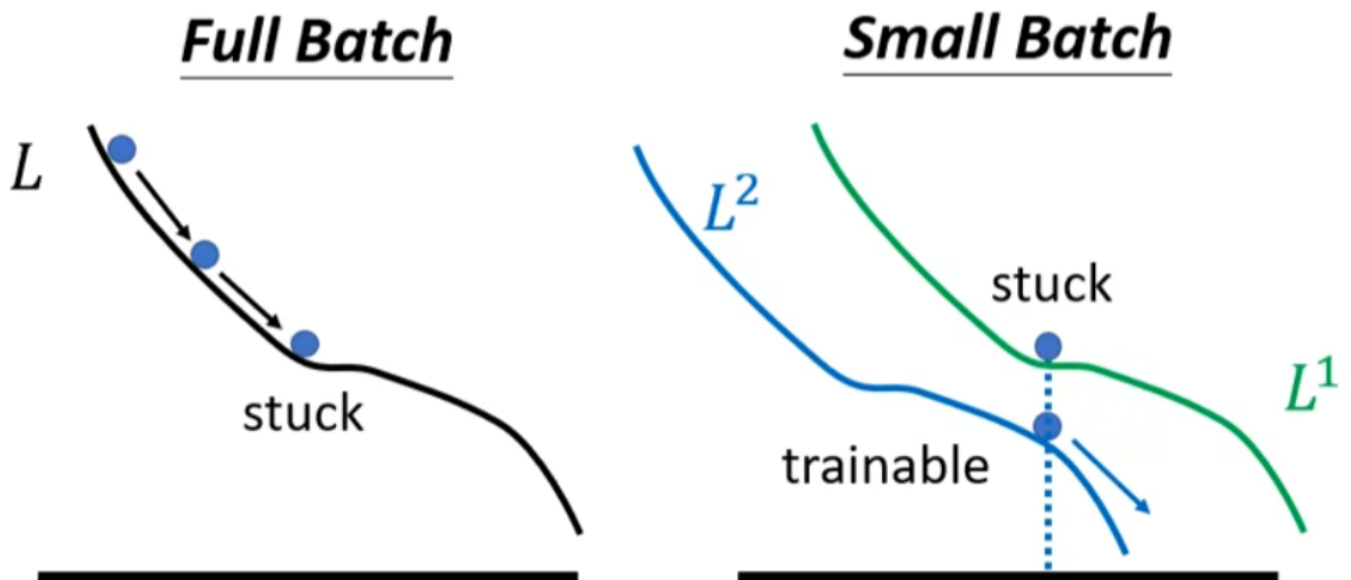
- 大Batch Size：全部数据看完才更新 -> 耗时长但是有用（稳）

- 小Batch size: 分组多次更新 -> 耗时少但是噪声影响大

但又由于GPU并行运算的存在, 在一定范围内大Batch Size似乎表现更好:



但小Batch size有噪声的存在, 可能会解决梯度下降卡住的问题



(2) Momentum

一般的梯度下降: $\theta^{i+1} = \theta^i - \eta g^i$

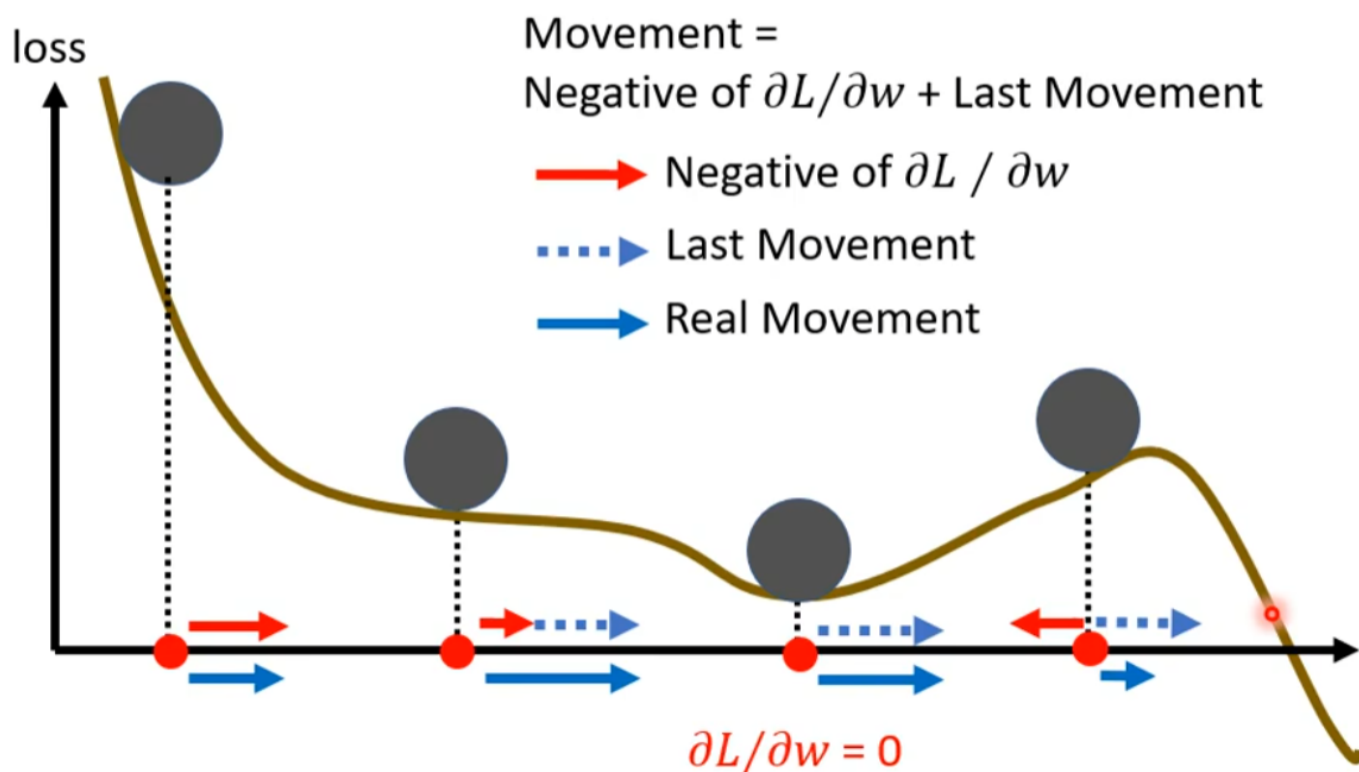
不是只往梯度的反方向移动, 还加入前一步的Momentum:

$$m^{i+1} = \lambda m^i - \eta g^i$$

$$\theta^{i+1} = \theta^i + m^{i+1}$$

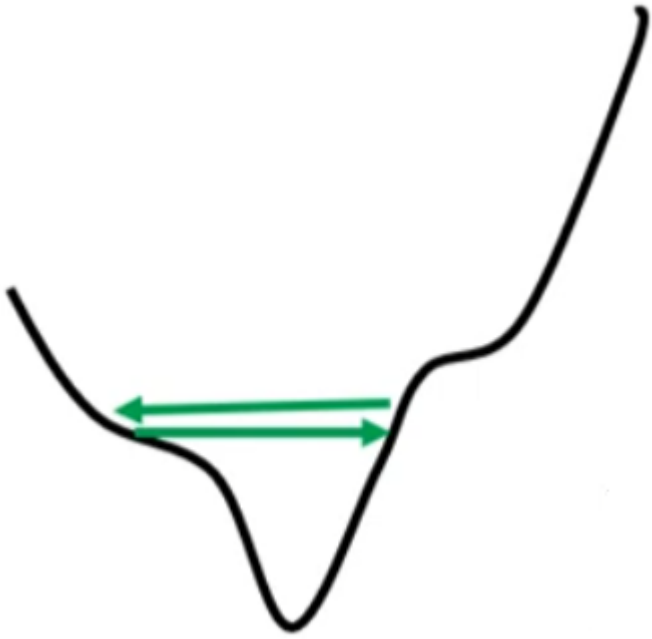
不仅仅和当前梯度有关，还和之前所有梯度有关

Gradient Descent + Momentum



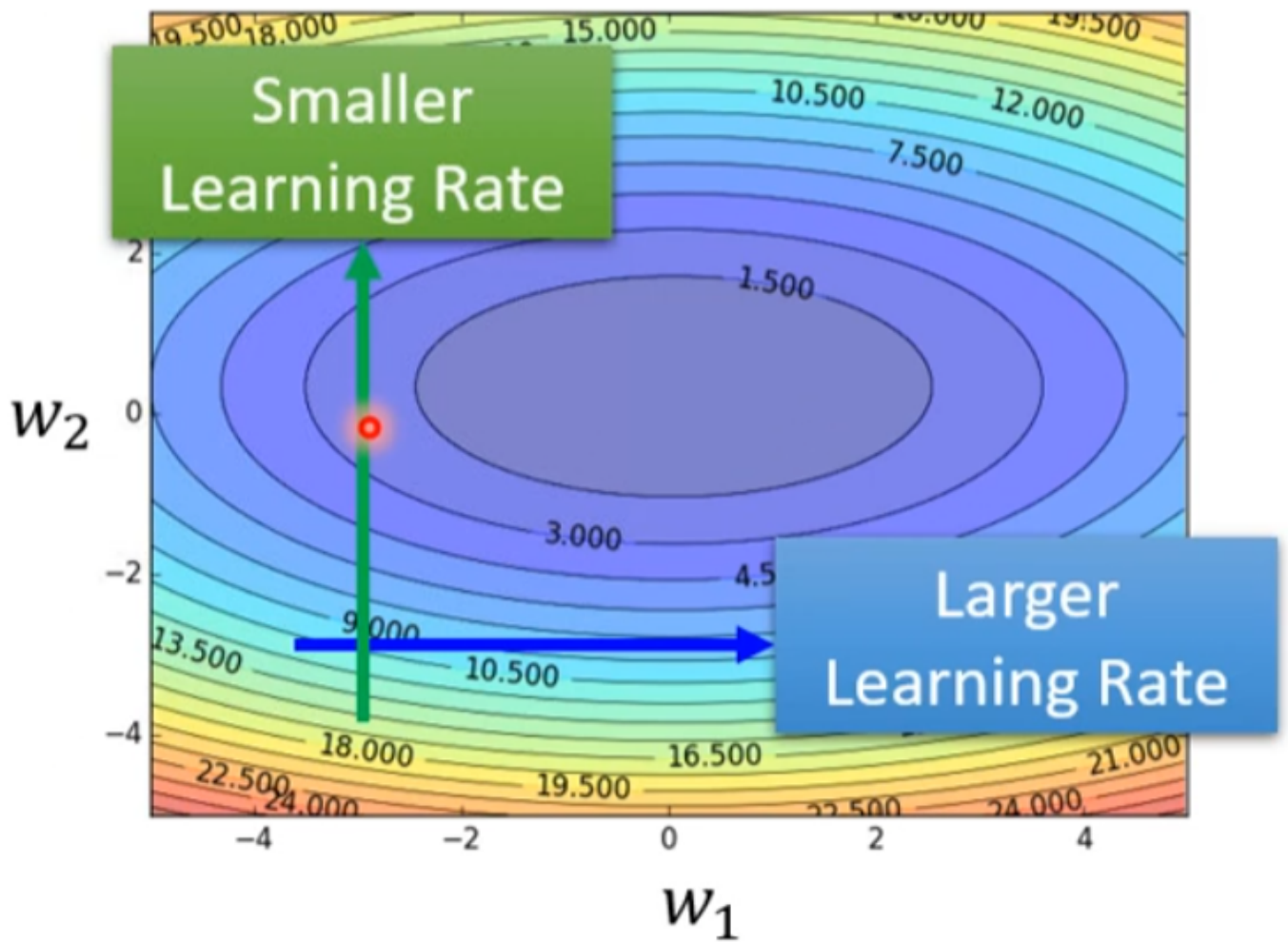
(3) 学习率

loss变小 != gradient变小，下面这个图就是左右来回走，一直没有到最低点

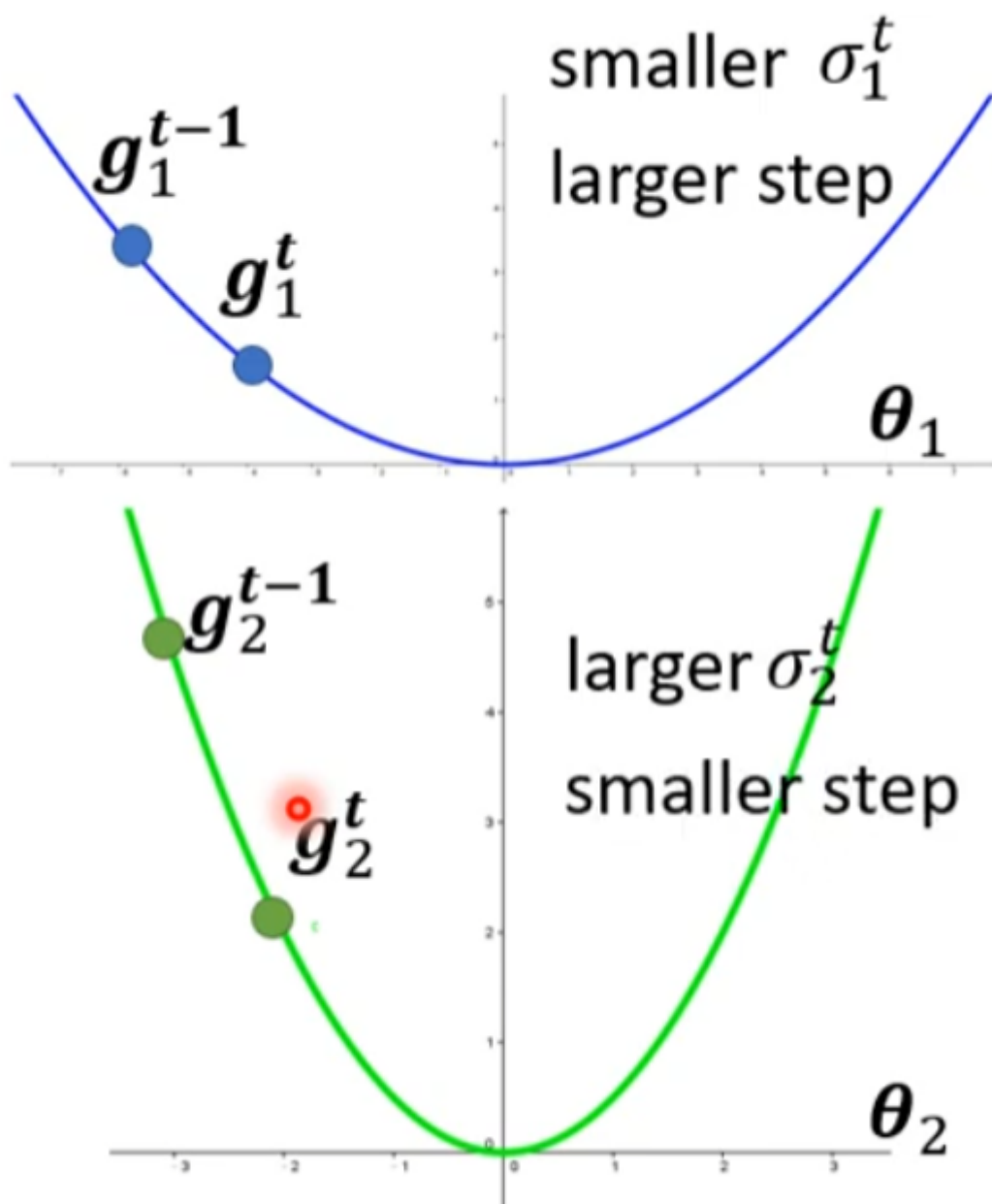


期望的 η 是多次走到谷底，one-size-fits-all的学习率是不能满足的 -> 特质化（期望：陡的时候小一点，缓的时候大一点）

使用 $\theta^{i+1} \leftarrow \theta^i - \eta g^i$ 变成 $\theta^{i+1} \leftarrow \theta^i - \frac{\eta}{\sigma_i^t} g^i$ ，让 $\frac{\eta}{\sigma_i^t}$ 具有参数依赖



- 在Adagrad中 $\sigma_i^t = \sqrt{\frac{1}{t+1} \sum_{i=1}^n (g_i^t)^2}$



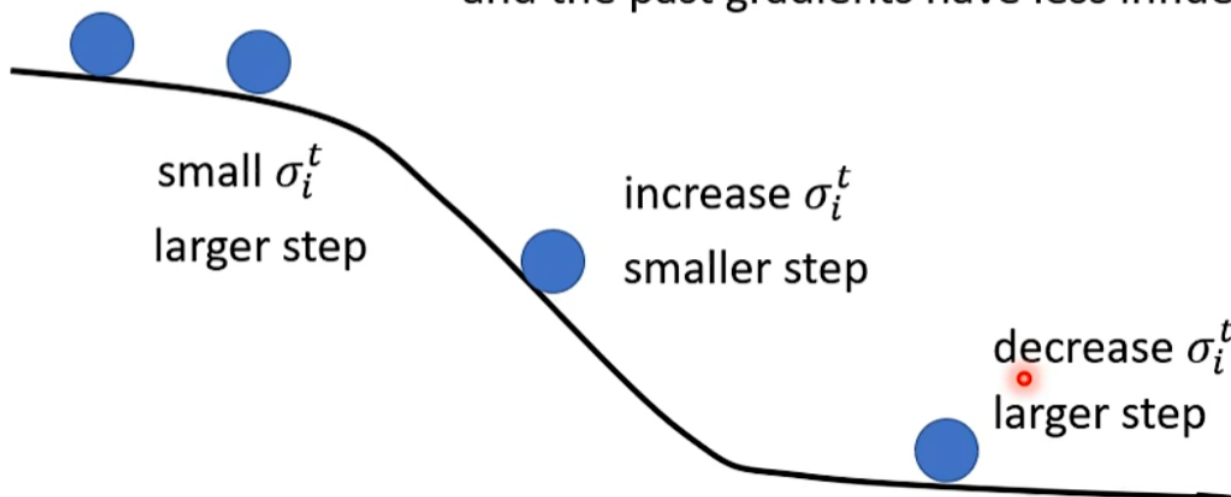
- RMSProp中, $\sigma_i^t = \sqrt{\alpha(\sigma_i^{t-1})^2 + (1 - \alpha)(g_i^t)^2}$

RMSProp

$$\theta_i^{t+1} \leftarrow \theta_i^t - \frac{\eta}{\sigma_i^t} g_i^t \quad \sigma_i^t = \sqrt{\alpha (\sigma_i^{t-1})^2 + (1 - \alpha) (g_i^t)^2} \quad 0 < \alpha < 1$$

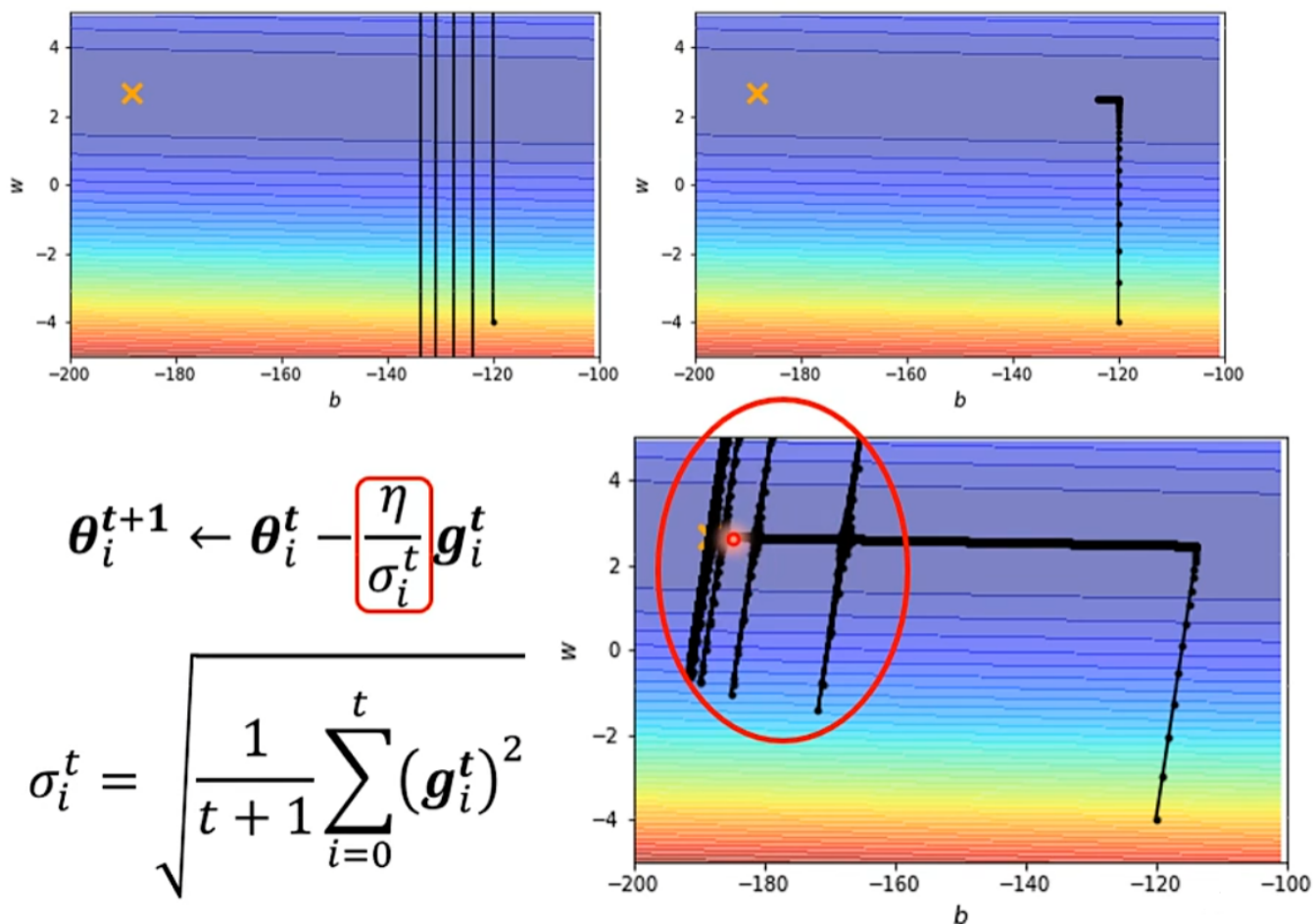
$g_i^1 \quad g_i^2 \quad \dots \quad g_i^{t-1}$


The recent gradient has larger influence, and the past gradients have less influence.



- Adam采用RMSProp+Momentum

Without Adaptive Learning Rate



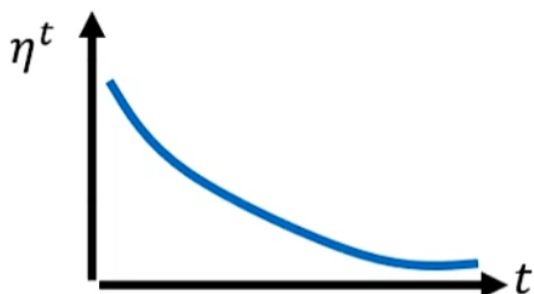
为了解决最后波动的问题，需要减小 η ，采用Learning Rate Scheduling -> 把 η 和时间 t 关联起来

(1) Learning Rate Decay: 逐渐变小

(2) Warm up: 先变大再变小 (目前没有理论原因，解释是先变大可以有一个比较好的探索)

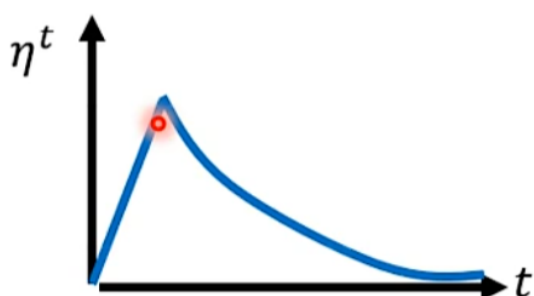
Learning Rate Scheduling

$$\theta_i^{t+1} \leftarrow \theta_i^t - \frac{\eta^t}{\sigma_i^t} g_i^t$$



Learning Rate Decay

After the training goes, we are close to the destination, so we reduce the learning rate.



Warm Up

Increase and then decrease?

At the beginning, the estimate of σ_i^t has large variance.

总结各种梯度下降的提升计算方法:

$$\theta_i^{t+1} \leftarrow \theta_i^t - \frac{\eta^t}{\sigma_i^t} m_i^t$$

Learning rate scheduling

Momentum: weighted sum of the previous gradients

Consider direction

root mean square of the gradients

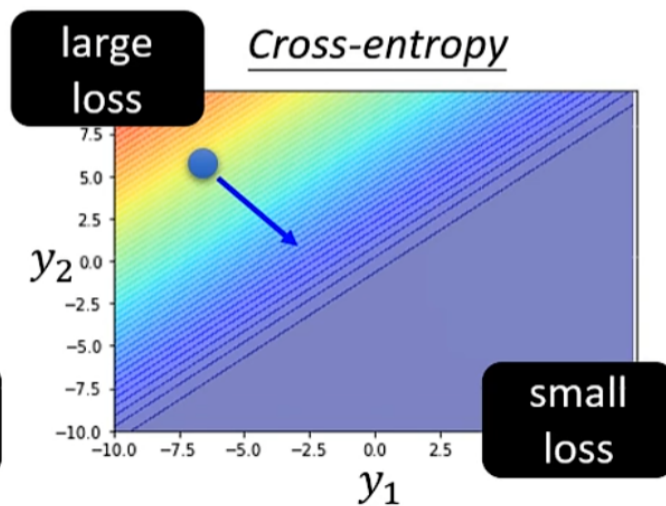
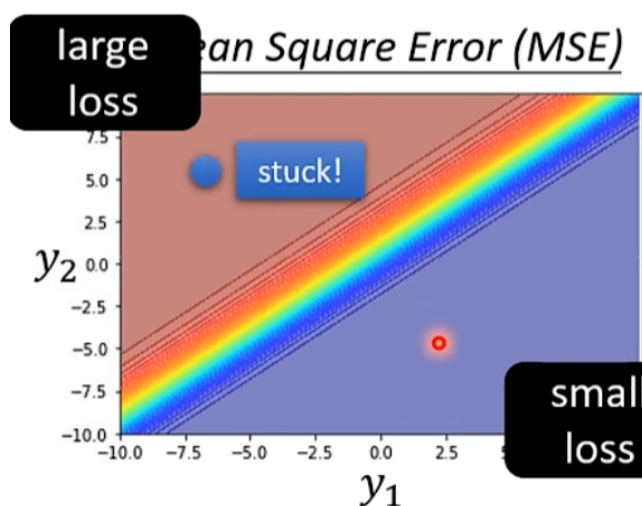
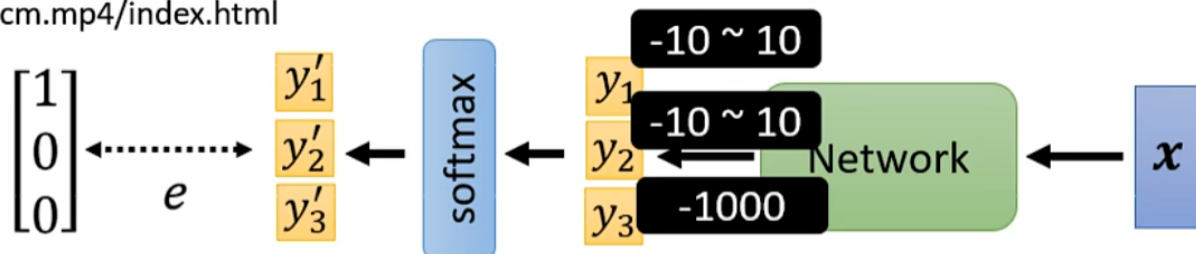
only magnitude

(4) 损失函数

修改损失函数让梯度变得简单好找最小点，要选择更合适的损失函数。

例：在分类任务中，选择MSE还是交叉熵。用MSE在左上角更新的就比较小，不如交叉熵。

[http://speech.ee.ntu.edu.tw/~tlkagk/courses/MLDS_2015_2/Lecture/Deep%20More%20\(v2\).ecm.mp4/index.html](http://speech.ee.ntu.edu.tw/~tlkagk/courses/MLDS_2015_2/Lecture/Deep%20More%20(v2).ecm.mp4/index.html)



个人总结

深度学习的核心是选择学习拟合函数和未知参数的确定。

功能复杂函数的函数有较强的学习能力，但在优化中由于较多的未知参数可能无法较快的确定较优的值（梯度下降面临的困难核心是无法在临界点“走出来”），因此衍生了一系列在保持神经网络结构框架前提下，优化梯度下降的方法。

- (1) 使用Batch会引入噪声，在一定程度上可以跨越一些零界点。
- (2) 使用Momentum向量可以在一定程度上影响梯度下降的方向，而不是固定的往微分的反方向移动。
- (3) 自适应调整学习率，在梯度陡时移动的少点，缓时移动的多点。
- (4) 选择合适的损失函数，可以从根本上减小梯度下降的难度。