

## Batch vs Mini-Batch gradient descent

- Vectorization allows us to compute on  $m$  examples effectively
- If we have a really large number of training examples, we split them into small parts mini-batches

$$X = \left[ \underbrace{x^{(1)} \dots x^{(1000)}}_{x^{[1]} \text{ (n, 1000)}} \mid \underbrace{x^{(1001)} \dots x^{(2000)}}_{x^{[2]}} \mid \dots \mid \underbrace{x^{(5001)} \dots x^{(m)}}_{x^{[5000]} \text{ (n, 1000)}} \right]$$

mini-batch

same for  $y$

## Mini-batch gradient descent

repeat {

for  $t=1, \dots, 5000$       1 step of gradient descent for  $x^{(t)}, y^{(t)}$

forward prop on  $x^{(t)}$  {

$$z^{[1]} = w^{[1]} x^{(t)} + b^{[1]}$$

$$a^{[1]} = g^{[1]}(z^{[1]})$$

...

$$a^{[L]} = g^{[L]}(z^{[L]})$$

compute cost  $J = \frac{1}{1000} \sum_{i=1}^L \mathcal{L}(\hat{y}^{(i)}, y^{(i)}) + \text{regularization}$

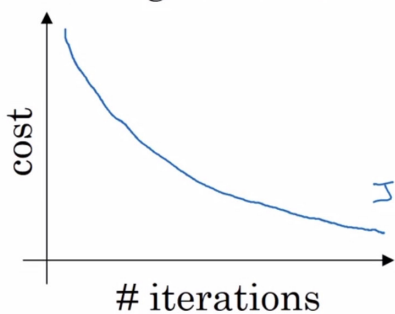
back prop to compute gradients

$$w^{[1]} := w^{[1]} - \alpha \text{d}w^{[1]}, \quad b^{[1]} := b^{[1]} - \alpha \text{d}b^{[1]}$$

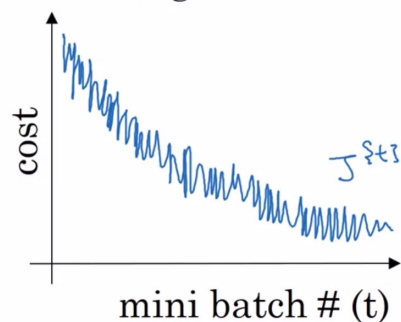
→ "1 epoch"

→ single pass through training set

Batch gradient descent



Mini-batch gradient descent



- Minibatch where every training example is its own mini-batch is called Stochastic Gradient Descent it never hits minimum.

### Stochastic

→ Lose speed from losing vectorization

### In-between

→ Fastest learning  
→ Vectorization  
→ Make progress without processing entire training set

### Batch

→ Too long per iteration  
→ vectorized  
→ all at once

## Choosing Mini-batch size

If small training set: use batch gradient descent  $\propto m \leq 2000$

Typical mini-batch sizes: 64, 128, 256, 512 (power of 2)

\* make sure mini-batch fits in cpu/gpu memory → mini-batch size is a hyperparameter

## Exponentially Weighted Averages

$$v_t = \beta v_{t-1} + (1 - \beta) \theta_t$$

Implementing:

$$v_0 = 0$$

repeat {

get next  $\theta_t$

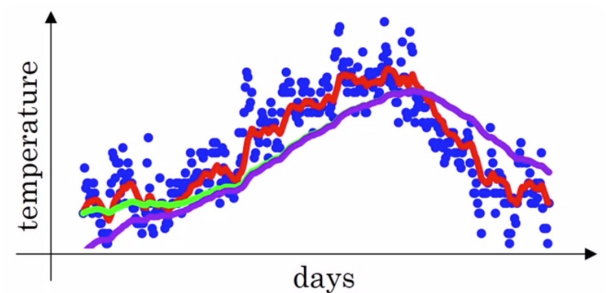
$$v_t := \beta v_0 + (1 - \beta) \theta_t$$

}

## Bias Correction

- The curve starts off really low and that is a problem
- We use  $\beta = \frac{v_t}{1 - \beta^t}$  because it helps get rid

of the bias as over examples  $\beta$  approaches zero and the purple line becomes the same as green line (before fixing bias) (after fixing bias)



## Gradient Descent with Momentum

Momentum: leads to much smaller steps and gradient descent is faster

On iteration  $t$ :

$$V_{dw}, V_{db} = 0$$

Compute  $dw, db$  on current mini-batch / batch

$$V_{dw} = \beta V_{dw} + (1-\beta) dw$$

$$\text{friction} \leftarrow V_{db} = \beta V_{db} + (1-\beta) db \rightarrow \text{acceleration} \rightarrow \text{velocity}$$

$$W = W - \alpha V_{dw}$$

$$b = b - \alpha V_{db}$$

$$"V_t = \beta V_t + (1-\beta) \theta_t"$$

the  $\beta$  helps the gradient descent move roughly in the speed it was moving in already

0.9 is the most common value for  $\beta$

$$\text{Bias correction} \rightarrow \frac{V_{dw}}{1-\beta_t}$$

## RMSprop

On iteration  $t$ :

compute  $dw, db$  on current mini-batch

$$S_{dw} = \beta S_{dw} + (1-\beta) dw^2 \rightarrow \text{element wise}$$

$$S_{db} = \beta S_{db} + (1-\beta) db^2$$

$$W = W - \alpha \frac{dw}{\sqrt{S_{dw} + \epsilon}}$$

$$b = b - \alpha \frac{db}{\sqrt{S_{db} + \epsilon}} \quad \epsilon = 10^{-8} \text{ for numerical stability}$$

## Adam Optimization Algorithm

$$V_{dw} = 0, S_{dw} = 0, V_{db} = 0, S_{db} = 0$$

On iteration  $t$ :

Compute  $dw, db$ , using current mini-batch

$$V_{dw} = \beta_1 V_{dw} + (1-\beta_1) dw$$

$$V_{db} = \beta_1 V_{db} + (1-\beta_1) db$$

$$S_{dw} = \beta_2 S_{dw} + (1-\beta_2) dw^2$$

$$S_{db} = \beta_2 S_{db} + (1-\beta_2) db^2$$

"momentum"  $\beta_1$

"RMSprop"  $\beta_2$

$$\left. \begin{aligned} V_{dw}^{\text{corrected}} &= \frac{V_{dw}^t}{(1 - \beta_1^t)} \\ V_{db}^{\text{corrected}} &= \frac{V_{db}^t}{(1 - \beta_2^t)} \\ S_{dw}^{\text{corrected}} &= \frac{S_{dw}}{(1 - \beta_1^t)} \\ S_{db}^{\text{corrected}} &= \frac{S_{db}}{(1 - \beta_2^t)} \end{aligned} \right\}$$

"bias correction"

$$\left. \begin{aligned} W &= W - \alpha \frac{V_{dw}^{\text{corrected}}}{\sqrt{S_{dw}^{\text{corrected}} + \epsilon}} \\ b &= b - \alpha \frac{V_{db}^{\text{corrected}}}{\sqrt{S_{db}^{\text{corrected}} + \epsilon}} \end{aligned} \right\}$$

update

## Hyperparameter Choice

$\alpha$  : needs to be tuned

$\beta_1$  : 0.9 (momentum)

$\beta_2$  : 0.999 (RMSprop)

$\epsilon$  :  $10^{-9}$  recommended

ADAM: Adaptive Moment Estimation

# Learning Rate Decay

1 epoch = 1 pass through data

$$\alpha = \frac{1}{1 + \text{decay-rate} * \text{epoch-num}} \alpha_0$$

hyperparameter

- Learning rate gradually decreases which helps us take smaller and smaller steps to converge to the minima.

Other ways -

$$\rightarrow \alpha = 0.95^{\text{epochnum}} \cdot \alpha_0 \quad \text{Exponential Decay}$$

some number less than 1

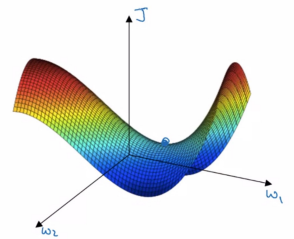
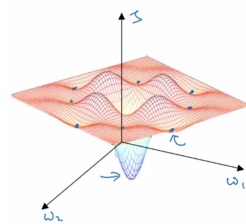
$$\rightarrow \alpha = \frac{k^{\text{constant}}}{\sqrt{\text{epochnum}}} \cdot \alpha_0 \quad \text{or} \quad \frac{k}{\sqrt{t}} \cdot \alpha_0$$

→ decreasing learning rate by a factor "discrete staircase"

→ Manual decay

## Local Optima Problem

- Most local optima are saddle points
- The chances of all  $n$  parameters to be 0 is really small and there is often a curve



## Problem of Plateau

- unlikely to get stuck in a bad local optima
- plateaus make learning slow (methods like ADAM help)

