

Adaptive (Sub)gradient methods

← "gradient-based"

Recall the advantage of first-order methods:

- scalable, easily made stochastic
- don't require a lot of memory or solving large linear systems

Disadvantages:

- sensitive to hyperparameter choices (stepsize)
- doesn't model the curvature of the objective function (slower convergence than 2nd order methods)

Question:

- can we have first order methods that are insensitive to hyperparam choices & better model curvature?

Today:

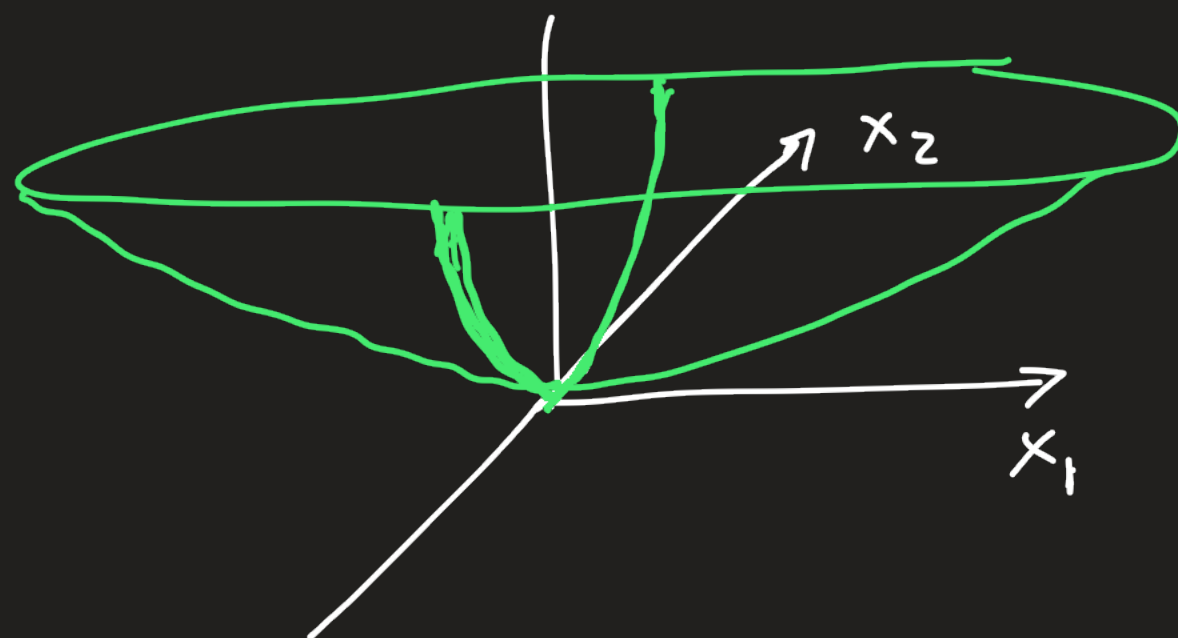
- Adagrad (2011)
- RMSprop (2012)
- AdaDelta (2012)
- * - ADAM (2015)
- SGD w/ momentum (heavy-ball method)
- Nesterov's Accelerated Gradient (NAG)

The ideas:

- use historical information to better model the function
- estimate curvature using first-order information
- handle different update rates for diff features

Ex: $f(x) = \frac{1}{2} \|Dx\|_2^2$ $D = \begin{bmatrix} 1 & 0 \\ 0 & 10 \end{bmatrix}$

$x^* = \begin{bmatrix} 0 \\ 0 \end{bmatrix}$



$f(x) = \frac{1}{2} (x_1^2 + 10x_2^2)$

with gradient descent:

$\nabla f(x) = D^T x = \begin{bmatrix} x_1 \\ 100x_2 \end{bmatrix}$

$x_{t+1} = x_t - \alpha D^T x_t = \begin{bmatrix} x_{t,1} - \alpha x_{t,1} \\ x_{t,2} - \underline{100\alpha} x_{t,2} \end{bmatrix}$

The optimal learning rate for $x_{t,2}$ is $\alpha = \frac{1}{100}$ because regardless of what $x_{0,2}$ is, we have

$x_{1,2} = x_{0,2} - x_{0,2} = 0 = x_2^*$

OTOH, the optimal learning rate for $x_{t,1}$ is $\alpha=1$ because then

$$x_{1,1} = x_{0,1} - x_{0,1} = 0 = x^*,$$

These two features have different optimal learning rates, so no shared learning rate will be superior to assigning them different learning rates.

Fact Newton's exactly accounts for the curvature so it adapts the stepsizes for each feature individually:

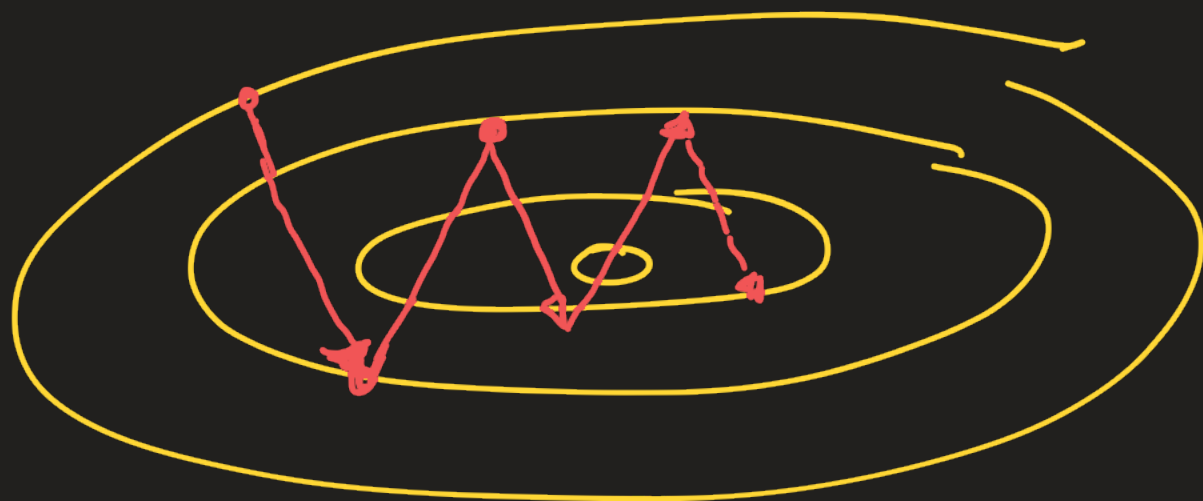
$$x_{t+1} = x_t - \alpha_t \left[\nabla^2 f(x_t) \right]^{-1} \nabla f(x_t)$$

$$\nabla^2 f(x_t) = D^2 = \begin{bmatrix} 1 & \\ & 100 \end{bmatrix} \Rightarrow x_{t+1} = x_t - \begin{bmatrix} \alpha_t & \\ & \alpha_t/100 \end{bmatrix} \nabla f(x_t)$$

$$= x_t - \begin{bmatrix} \alpha_t & \\ & \alpha_t/100 \end{bmatrix} \begin{bmatrix} 1 & \\ & 100 \end{bmatrix} x_t = 0$$

so Newton's method converges in one step when $\alpha_t = 1$

SGD w/ momentum

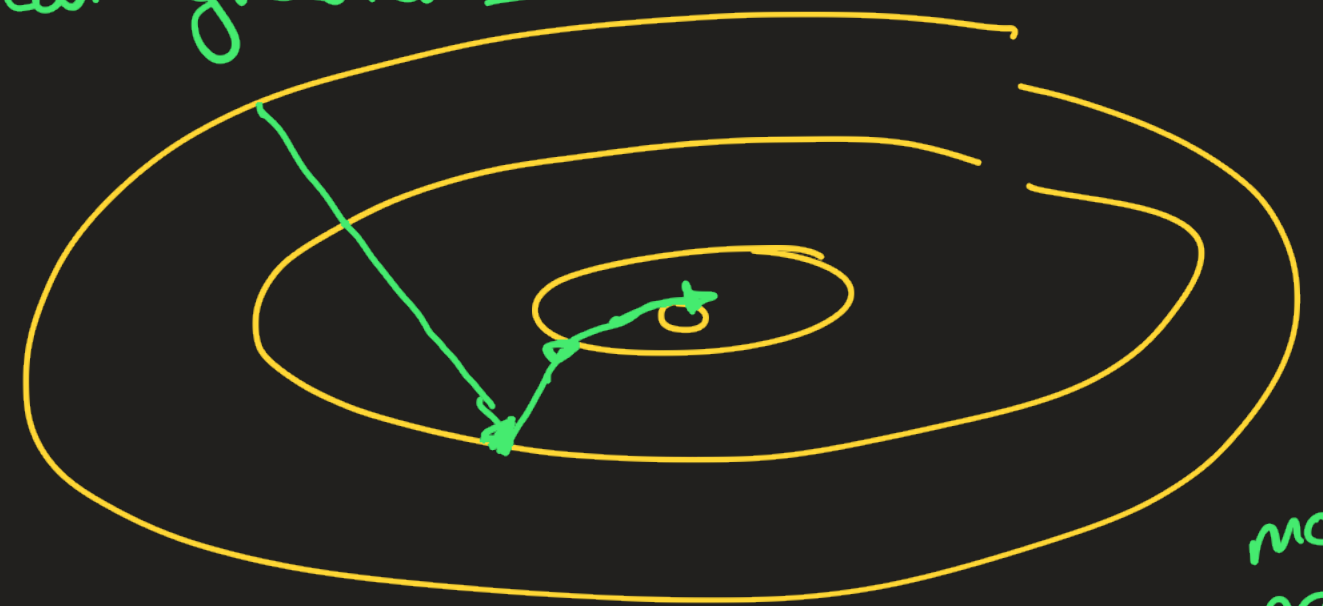


gradient descent oscillates in this case because it always moves perpendicular to the level set at the current point.

$$\text{level set } \mathcal{L}_{f, \alpha} = \{x : f(x) = \alpha\}$$



momentum reduces oscillatory behavior in recurring directions by averaging out historical gradients



SGD w/ momentum

$$\mu_t = \alpha_t \nabla f_t(\omega_t) + \gamma_t \mu_{t-1}$$

update direction $\rightarrow$

momentum param $\downarrow$

$$\omega_{t+1} = \omega_t - \mu_t$$

typically take $\gamma_t = \gamma = 0.9$

here ∇f_t is an estimate of ∇f at iteration t .

called Polyak heavy-ball method and can write these updates as

$$\underline{w_{t+1} = w_t - \mu_t = w_t - [\alpha_t \nabla f_t(w_t) + \gamma_t \mu_{t-1}]}$$

and note

$$w_t = w_{t-1} - \mu_{t-1} \Rightarrow \mu_{t-1} = w_{t-1} - w_t$$

plug in to get

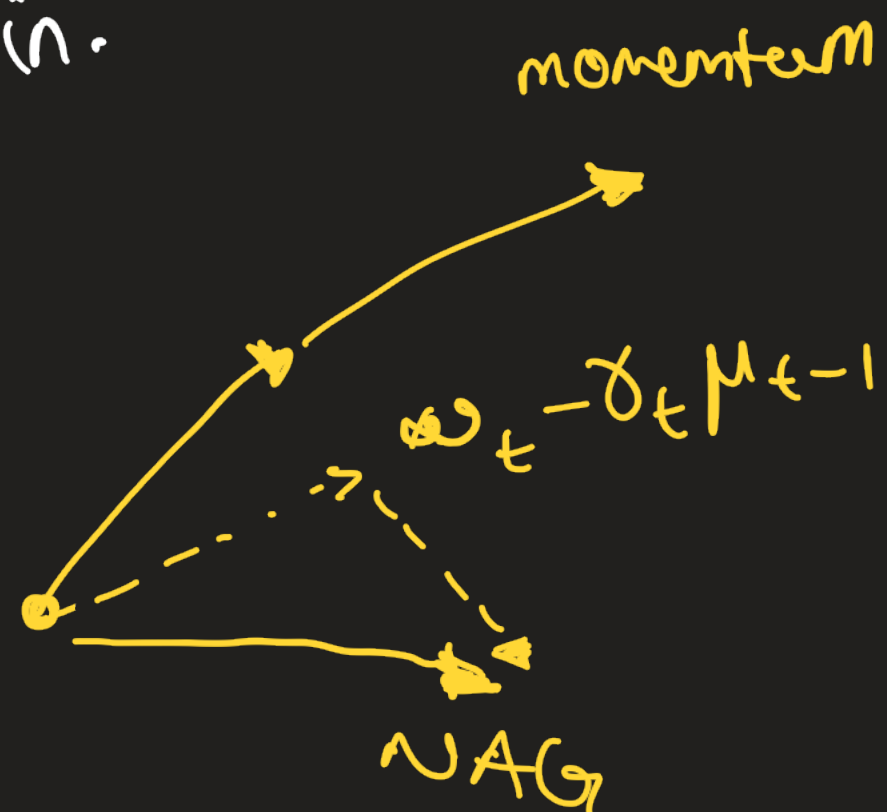
$$\underline{w_{t+1} = w_t - \alpha_t \nabla f_t(w_t) - \gamma_t (w_{t-1} - w_t)}$$

Nesterov's Accelerated Gradient Descent (NAG)

$$\mu_t = \gamma_t \mu_{t-1} + \alpha_t \nabla f_t(\omega_t - \gamma_t \mu_{t-1})$$

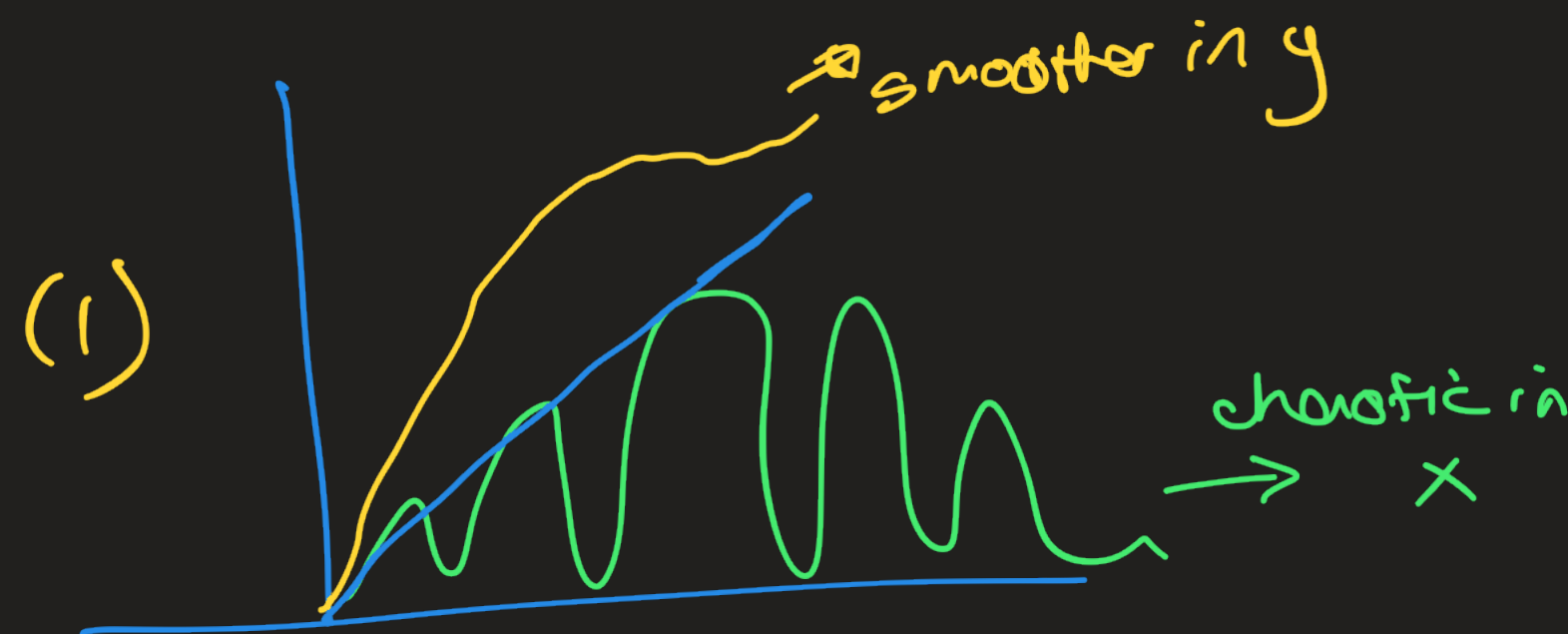
$$\omega_{t+1} = \omega_t - \mu_t$$

instead of trying to use just past information,
look ahead along the current trend (moving in the μ_{t-1}
direction) to choose an even better direction to
move in.



Ada Grad ("Adaptive Gradient Descent") diagonal version

- learn different stepsizes for each feature
- (1) - useful when loss surface has different scales with respect to each feature
- (2) - or, when dealing with sparse data where some features appear infrequently.



(2) Ex: consider a NLP application: fit classifier on bag of word features (BOW) to determine whether a review is positive

$$x = \begin{bmatrix} \overset{\text{'good'}}{\downarrow} 0 & \overset{\text{'bad'}}{\downarrow} 2 & 0 & \dots & \overset{\text{'enjoy'}}{\downarrow} 0 & \dots & \overset{\text{'sublime'}}{\downarrow} 0 & 1 \end{bmatrix}$$

$|V|$ dimension vector where V is our vocab.

then, e.g. 'sublime' is very discriminative but seen infrequently $\rightarrow$ so want to update w_{10K} more aggressively when it is seen than say w_1 (corresp to 'good')

Idea: use per-feature learning rates that are inversely proportional to how much that feature has changed before.

$$w_{t+1} = w_t - \alpha \left(D_t + \epsilon I \right)^{-1/2} \nabla f_t(w_t) \quad \epsilon \approx 10^{-8}$$

where

$$D_t = \text{diag} \left(\sum_{i=r}^t \nabla f_i(w_i) \nabla f_i(w_i)^T \right)$$

$$(\omega_{t+1})_i = (\omega_t)_i - \frac{\alpha \cdot 1}{\left(\sum_{j=1}^t [\nabla f(\omega_j)]_i^2 \right)^{1/2}} [\nabla f(\omega_t)]_i$$

where $i \in [d]$ is a feature index

Advantages:

- relatively insensitive to α compared to SGD
- works impressively for sparse data

Drawbacks:

- the learning rate goes to zero with time because

$$\sum_{j=1}^t \nabla f(\omega_j)_i^2 \rightarrow \infty \quad \text{as } t \rightarrow \infty$$

regardless of whether you're close to convergence or not.

- you need to store a learning rate for every parameter.