

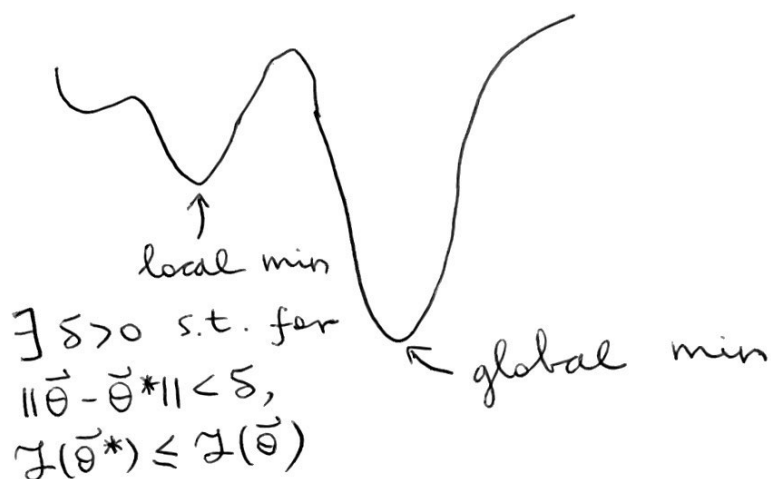
Lecture 12

Optimization

Problem: minimize a scalar-valued loss or cost function $\mathcal{J}(\vec{\theta})$:

$$\vec{\theta}^* = \underset{\vec{\theta}}{\operatorname{argmin}} \mathcal{J}(\vec{\theta})$$

Assume that $\vec{\theta} \in \mathbb{R}^D \xleftarrow{\# \text{ prms}} \Rightarrow \underline{\text{continuous optimization}}$

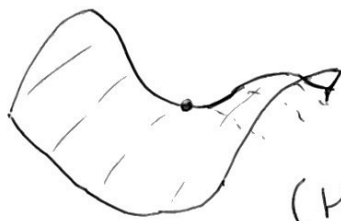


Define $\begin{cases} g_i(\vec{\theta}) = \frac{\partial}{\partial \theta_i} \mathcal{J}(\vec{\theta}), & \text{gradient} \\ H_{ij}(\vec{\theta}) = \frac{\partial^2}{\partial \theta_i \partial \theta_j} \mathcal{J}(\vec{\theta}). & \text{Hessian} \end{cases}$

$\nabla \mathcal{J}(\vec{\theta}^*) = \vec{0}$ and $H(\vec{\theta}^*)$ is pos. def. (that is, all its eigenvalues are > 0),

$\vec{\theta}^*$ is a local min.

Saddle points:



some directions downhill, some uphill
 $(H(\bar{\theta}^*))$ has both pos. & neg. eigenvalues)

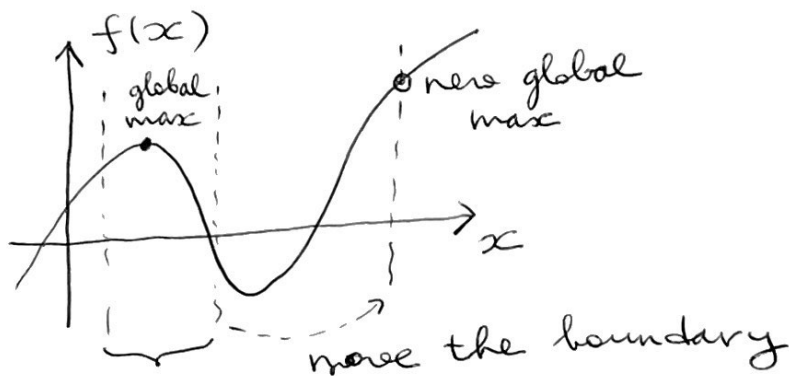
Constraints: $\begin{cases} g_j(\bar{\theta}) \leq 0, & j=1,2,\dots & \text{inequality} \\ h_k(\bar{\theta}) = 0, & k=1,2,\dots & \text{equality} \end{cases}$

Ex. (a) $h(\bar{\theta}) = 1 - \sum_{i=1}^D \theta_i = 0$ prms sum to 1.0

(b) $g_i(\bar{\theta}) = -\theta_i \leq 0$ all prms pos.
 $i=1,\dots,D$

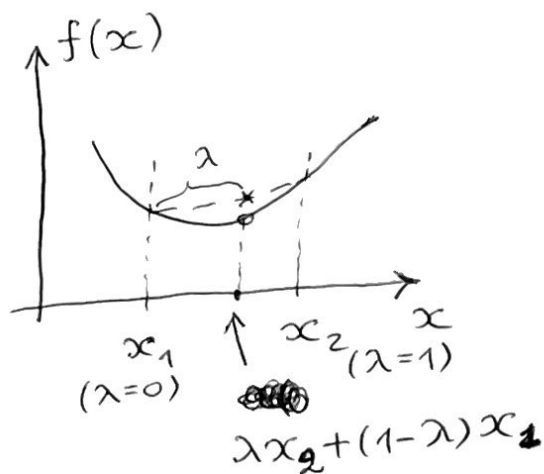
Feasible set = subset of $\bar{\theta}$ that satisfies all the constraints.

Constraints can change the number of min/max of a function:



Convex functions: $f(\lambda \bar{x} + (1-\lambda)\bar{y}) \leq \lambda f(\bar{x}) + (1-\lambda)f(\bar{y})$
 $\forall \bar{x}, \bar{y}; 0 \leq \lambda \leq 1$
 strictly convex if $<$

$\left[\begin{array}{l} f(\bar{x}) \text{ is (strictly) concave} \\ \bar{y} - f(\bar{x}) \text{ is (strictly) convex} \end{array} \right]$



If $f(\vec{x})$ is strictly convex in some domain $\Rightarrow H = \nabla^2 f(\vec{x})$ is pos. def. in that domain.

First-order methods

(no curvature information)

$$\vec{\theta}_{t+1} = \vec{\theta}_t + \underset{\substack{\uparrow \\ \text{learning rate} \\ \text{(step size)}}}{\eta_t} \underset{\substack{\leftarrow \text{descent direction}}}{\vec{d}_t}$$

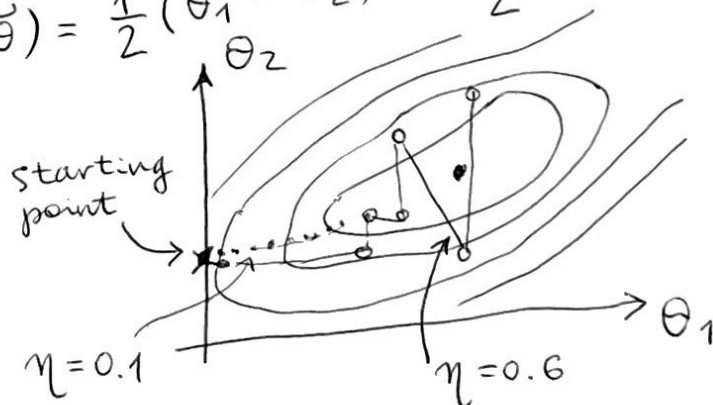
$\vec{\theta}_0$ = starting point

Usually, $\vec{d}_t = - \underbrace{\vec{\nabla}_{\vec{\theta}} \mathcal{I}(\vec{\theta})}_{\vec{g}_t = \text{gradient}} \big|_{\vec{\theta}_t} \Leftarrow \begin{matrix} \text{steepest} \\ \text{descent} \end{matrix}$

Learning rate dependence:

Ex. $\mathcal{I}(\vec{\theta}) = \frac{1}{2}(\theta_1^2 - \theta_2)^2 + \frac{1}{2}(\theta_1 - 1)^2$

$\eta = 0.6$ fails to converge, even on a convex problem



Line search: solve 1D minimization problem along the direction of $\vec{d}_t$:

$$\eta_t = \underset{\eta > 0}{\operatorname{argmin}} \mathcal{L}(\vec{\theta}_t + \eta \vec{d}_t)$$

Ex. $\mathcal{L}(\vec{\theta}) = \frac{1}{2} \theta_i \underbrace{A_{ij} \theta_j}_{\text{sums over repeated indices}} + b_i \theta_i$

$$\begin{aligned} \frac{d}{d\eta} \mathcal{L}(\vec{\theta} + \eta \vec{d}) &= d_i A_{ij} (\theta_j + \eta d_j) + d_i b_i = \\ &= d_i (A_{ij} \theta_j + b_j) + \eta d_i A_{ij} d_j = 0, \\ &\quad \text{s.t.} \end{aligned}$$

$$\eta = - \frac{d_i (A_{ij} \theta_j + b_j)}{d_i A_{ij} d_j}$$

Momentum methods

"heavy ball" method:

$$\begin{cases} \vec{m}_t = \beta \vec{m}_{t-1} + \vec{g}_{t-1}, & 0 < \beta < 1 \\ \vec{\theta}_t = \vec{\theta}_{t-1} - \eta_t \vec{m}_t \end{cases} \quad (\text{typically, } \beta = 0.9)$$

$\beta = 0 \Rightarrow$ steepest descent

note that

$$\vec{m}_t = \beta \vec{m}_{t-1} + \vec{g}_{t-1} = \beta^2 \vec{m}_{t-2} + \beta \vec{g}_{t-2} + \vec{g}_{t-1} = \dots$$

$$= \sum_{\tau=0}^{t-1} \beta^{\tau} \bar{g}_{t-\tau-1}$$

$$\mathcal{H} \quad \bar{g}_t = \bar{g} = \text{const} \quad (\forall t),$$

$$\bar{m}_t = \bar{g} \sum_{\tau=0}^{t-1} \beta^{\tau} \approx \bar{g} \underbrace{\sum_{\tau=0}^{\infty} \beta^{\tau}}_{\frac{1}{1-\beta}} = \frac{\bar{g}}{1-\beta} \quad \text{in this limit}$$

Nesterov accelerated gradient method:

$$\begin{cases} \tilde{\theta}_{t+1} = \bar{\theta}_t + \beta(\bar{\theta}_t - \bar{\theta}_{t-1}), & \Leftarrow \text{extrapolation step} \\ \bar{\theta}_{t+1} = \tilde{\theta}_{t+1} - \eta_t \nabla_{\bar{\theta}} \mathcal{L}(\tilde{\theta}_{t+1}) \end{cases} \quad 0 < \beta < 1$$

One can write

$$\begin{cases} \bar{m}_{t+1} = \beta \bar{m}_t - \eta_t \nabla_{\bar{\theta}} \mathcal{L}(\bar{\theta}_t + \beta \bar{m}_t), \\ \bar{\theta}_{t+1} = \bar{\theta}_t + \underbrace{\bar{m}_{t+1}}_{\substack{\text{new location for the} \\ \text{gradient, updated} \\ \text{by momentum}}} \bar{\theta}_{t+1} - \bar{\theta}_t \end{cases}$$

Second-order methods

Newton's method:

$$\bar{\theta}_{t+1} = \bar{\theta}_t - \eta_t H_t^{-1} \bar{g}_t, \quad \text{where}$$

$$H_{ij,t} = \frac{\partial^2}{\partial \theta_i \partial \theta_j} \mathcal{L}(\bar{\theta}) \Big|_{\bar{\theta}_t}$$

Consider
$$\mathcal{J}(\vec{\theta}) = \mathcal{J}(\vec{\theta}_t) + g_{i,t}(\vec{\theta}_i - \theta_{i,t}) + \frac{1}{2}(\theta_i - \theta_{i,t}) H_{ij,t} (\theta_j - \theta_{j,t}).$$

↑ quadratic expansion

$$\frac{\partial}{\partial \theta_k} \mathcal{J}(\vec{\theta}) = g_{k,t} + H_{kj,t} (\theta_j - \theta_{j,t}) = 0, \text{ or}$$

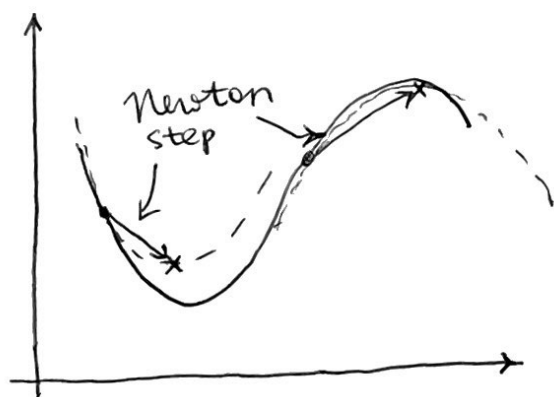
$$\vec{g}_t = -H_t (\vec{\theta} - \vec{\theta}_t),$$

$$\boxed{\vec{\theta} = \vec{\theta}_t - H_t^{-1} \vec{g}_t}$$

↑ "intelligent" step size

In practice, can use line search to find the empirical step size η_t .

Note that Newton's method finds the min of a quadratic function in one step, by construction.



Newton steps can go towards maxima instead of minima in non-convex functions (i.e., functions in which the sign of curvature changes)

Tikhonore regularization

Define $\delta_i = \theta_i - \theta_{i,t}$, then

$$\mathcal{I}(\vec{\theta}) \approx \mathcal{I}(\vec{\theta}_t) + g_{i,t} \delta_i + \frac{1}{2} \delta_i H_{ij,t} \delta_j.$$

So we minimize $\mathcal{I}(\vec{\theta}) + \underbrace{\frac{\lambda}{2} \vec{\delta}^2}_{\text{penalty for large steps}}$ instead of

$$\mathcal{I}(\vec{\theta}), \text{ we get } \vec{\delta} = - \underbrace{(H + \lambda \mathbb{I})^{-1}}_{\text{Tikhonore regularization}} \vec{g}$$

Stochastic gradient descent (SGD)

$$\text{Consider } \mathcal{I}(\vec{\theta}_t) = \frac{1}{N} \sum_{n=1}^N \underbrace{\mathcal{I}_n(\vec{\theta}_t)}_{\text{cost per datapoint}}$$

$$\text{idea: replace } \vec{g}_t = \frac{1}{N} \sum_{n=1}^N \vec{\nabla}_{\vec{\theta}} \mathcal{I}_n(\vec{\theta}_t)$$

$$\text{with } \vec{g}_t = \frac{1}{|B_t|} \sum_{n \in B_t} \vec{\nabla}_{\vec{\theta}} \mathcal{I}_n(\vec{\theta}_t)$$

B_t = subset of datapoints to use at iteration t ;

$|B_t| \ll N$. Typically, $|B_t| \sim 10^1 - 10^2$

Ex. SGD for linear regression

$$\mathcal{L}(\vec{\theta}) = \frac{1}{2N} \sum_{n=1}^N (\vec{x}_n \cdot \vec{\theta} - y_n)^2.$$

$$\text{Then } \vec{g}_t = \frac{1}{N} \sum_{n=1}^N (\vec{x}_n \cdot \vec{\theta} - y_n) \vec{x}_n$$

If we use $|B|=1$, we obtain: diff. between model & data

$$\vec{\theta}_{t+1} = \vec{\theta}_t - \eta_t (\vec{x}_{n_t} \cdot \vec{\theta} - y_{n_t}) \vec{x}_{n_t}, \text{ where}$$

n_t = index of the example @ iteration t

preconditioned SGD:

$$\vec{\theta}_{t+1} = \vec{\theta}_t - \eta_t M_t^{-1} \vec{g}_t, \text{ where}$$

M_t is a preconditioning matrix (typically pos. def.); a 'poor man's Hessian'

AdaGrad ('adaptive gradient')

$$\theta_{i,t+1} = \theta_{i,t} - \eta_t \frac{1}{\sqrt{S_{i,t} + \epsilon}} g_{i,t},$$

where $S_{i,t} = \sum_{\tilde{t}=1}^t \underbrace{g_{i,\tilde{t}}^2}_{\text{sum over past grad}^2}$, $\epsilon > 0$ is a small hyperparam

Typically, $\eta_t = \eta = \text{const}$.

clearly, $M_t = \text{diag}((\vec{S}_t + \vec{\epsilon})^{1/2})$, where

$$\vec{S}_t = \begin{pmatrix} S_{1,t} \\ S_{2,t} \\ \vdots \\ S_{D,t} \end{pmatrix} \quad \text{and} \quad \vec{\epsilon} = \begin{pmatrix} \epsilon \\ \epsilon \\ \vdots \\ \epsilon \end{pmatrix}$$

$\uparrow$
 dim of $\vec{\theta}$

[adaptive learning rate]

$\rightarrow$ components of $\vec{S}_t$
 may grow large
 as $t \uparrow \rightarrow$ learning
 rate $\downarrow$.

RMSProp

Idea: use a weighted average to compute $S_{i,t}$:

$$\begin{aligned} S_{i,t+1} &= \beta S_{i,t} + (1-\beta) g_{i,t}^2 = \\ &= (1-\beta) g_{i,t}^2 + \beta [\beta S_{i,t-1} + (1-\beta) g_{i,t-1}^2] = \\ &= (1-\beta) [g_{i,t}^2 + \beta g_{i,t-1}^2] + \beta^2 S_{i,t-1} = \dots = \\ &= (1-\beta) \underbrace{[g_{i,t}^2 + \beta g_{i,t-1}^2 + \dots + \beta^{t-1} g_{i,1}^2]}_{\text{weighted sum of grad}^2} \end{aligned}$$

Typically, $\beta = 0.9$.

The update is as with AdaGrad:

$$\theta_{i,t+1} = \theta_{i,t} - \eta_t \frac{1}{\sqrt{S_{i,t} + \epsilon}} g_{i,t}$$

AdaDelta

$$\theta_{i,t+1} = \theta_{i,t} - \eta_t \frac{\sqrt{\delta_{i,t-1} + \epsilon}}{\sqrt{\delta_{i,t} + \epsilon}} g_{i,t}, \text{ where}$$

$$\begin{cases} \delta_{i,t+1} = \beta \delta_{i,t} + (1-\beta) g_{i,t}^2, \\ \delta_{i,t} = \beta \delta_{i,t-1} + (1-\beta) (\underbrace{\Delta \theta_{i,t}}_{\theta_{i,t+1} - \theta_{i,t}})^2. \end{cases} \leftarrow \text{same as RMSProp}$$

Adam = RMSProp + momentum
"adaptive moment estimation"

Compute $\begin{cases} m_{i,t} = \beta_1 m_{i,t-1} + (1-\beta_1) g_{i,t}, \\ \delta_{i,t} = \beta_2 \delta_{i,t-1} + (1-\beta_2) g_{i,t}^2. \end{cases}$

Update: $\theta_{i,t+1} = \theta_{i,t} - \eta_t \frac{1}{\sqrt{\delta_{i,t} + \epsilon}} m_{i,t}$

Recommended values:

$$\begin{cases} \beta_1 = 0.9, & \beta_2 = 0.999 \\ \epsilon = 10^{-6}, & \eta_t = \eta = 10^{-3} \end{cases}$$

Finally, one can 'bias-corrected' moments:

$$\begin{cases} \vec{m}_t \rightarrow \hat{\vec{m}}_t = \frac{\vec{m}_t}{1-\beta_1^t}, \\ \vec{s}_t \rightarrow \hat{\vec{s}}_t = \frac{\vec{s}_t}{1-\beta_2^t}. \end{cases}$$

These are only important for small t (in the beginning of optimization)