

HW 1 Released

Simon OH → Thursday 10:30AM
This Week

Clarke Differential



Subdifferential and Subgradient

Definition: Given $f: \mathbb{R}^d \rightarrow \mathbb{R}$, for every x , the subdifferential set is defined as

$\partial_s f(x) \triangleq \{s \in \mathbb{R}^d : \forall x' \in \mathbb{R}^d, f(x') \geq f(x) + s^\top (x' - x)\}$. The elements in the subdifferential set are subgradients.

$$\min f(w)$$

$$h() : x_{t+1} = x_t - \eta \nabla f(x_t)$$

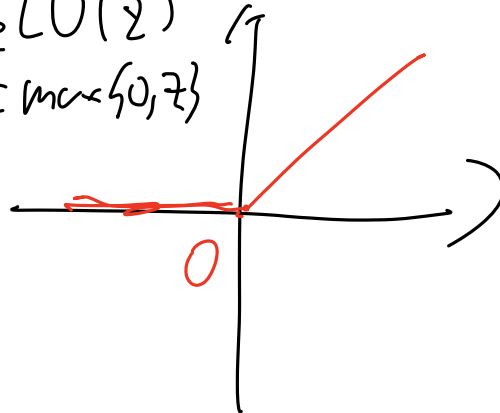
Subgradient descent

$$x_{t+1} = x_t - \eta g_t$$

$$g_t \in \partial_s f(x_t)$$

$$p_{\ell}(\mathcal{U}(z))$$

$$= \max\{0, z\}$$



$$\partial_s f(0) = [\bar{0}, 1]$$

Subdifferential and Subgradient

Definition: Given $f: \mathbb{R}^d \rightarrow \mathbb{R}$, for every x , the subdifferential set is defined as

$\partial_s f(x) \triangleq \{s \in \mathbb{R}^d : \forall x' \in \mathbb{R}^d, f(x') \geq f(x) + s^\top(x' - x)\}$. The elements in the subdifferential set are subgradients.

• If f is convex $\rightarrow \partial_s f$ exists everywhere

• If f is convex & differentiable
 $\partial_s f = \{ \nabla f \}$

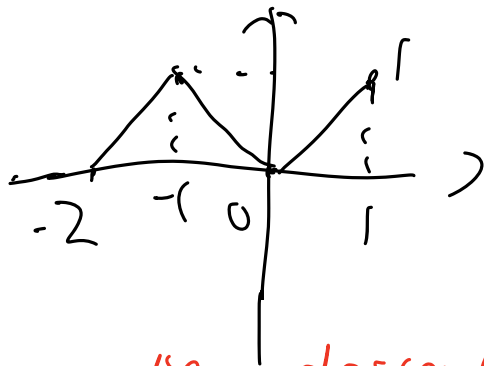
- $\mathcal{O}\left(\frac{1}{\sqrt{T}}\right)$ in convex

Subdifferential is not enough

Definition: Given $f: \mathbb{R}^d \rightarrow \mathbb{R}$, for every x , the subdifferential set is defined as

$\partial_s f(x) \triangleq \{s \in \mathbb{R}^d : \forall x' \in \mathbb{R}^d, f(x') \geq f(x) + s^\top (x' - x)\}$. The elements in the subdifferential set are subgradients.

Problem: $\mathbb{N}$ is not convex



subgradient descent
is not well-defined

$x = -1$
we need $s, \forall x'$
 $f(x') \geq f(-1) + s^\top (x' - (-1))$
let $x' = -2$
 $0 \geq 1 + s(-1) \Rightarrow s \geq 1$
let $x' = 1$
 $1 \geq 1 + 2s \Rightarrow s \leq 0$

Clarke Differential

$$x_{t+1} = x_t - \eta g_t$$

$$g_t \in \partial f(x)$$

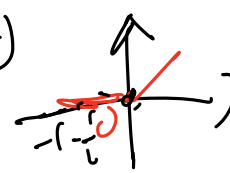
Definition: Given $f: \mathbb{R}^d \rightarrow \mathbb{R}$, for every x , the Clark differential is defined as

$$\partial f(x) \triangleq \text{conv} \left(\{s \in \mathbb{R}^d : \exists \{x_i\}_{i=1}^{\infty} \rightarrow x, \{ \nabla f(x_i) \}_{i=1}^{\infty} \rightarrow s\} \right).$$

The elements in the subdifferential set are subgradients.

$$\text{conv}(S) = \left\{ v : v = \sum_{i=1}^n \lambda_i u_i, u_i \in S, \lambda_i \geq 0, \sum_{i=1}^n \lambda_i = 1 \right\}$$

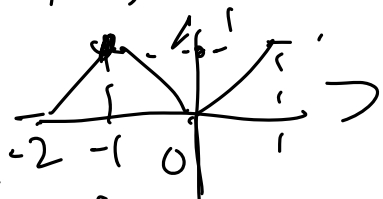
example 1 $\text{ReLU}(z)$

$$= \max\{0, z\}$$


$$\{x_i\}, -1, -\frac{1}{2}, \dots \rightarrow 0$$

$$\nabla f(x_i) = 0$$

example 2:



$$\{x_i\}, 1, \frac{1}{2}, \dots \rightarrow 0$$

$$\nabla f(x_i) = 1$$

$$\{x_i\} : -2, -1.5, \dots \rightarrow -1 \quad \nabla f(x_i) = 1$$

$$\{x_i\} : 0, -\frac{1}{2}, \dots \rightarrow -1 \quad \nabla f(x_i) = -1$$

$$\Rightarrow \partial f(0) = [0, 1]$$

$$\partial f(1) = [-1, 1]$$

When does Clarke differential exists

Definition (Locally Lipschitz): $f : \mathbb{R}^d \rightarrow \mathbb{R}$ is locally Lipschitz if $\forall x \in \mathbb{R}^d$, there exists a neighborhood S of x , such that f is Lipschitz in S .

$$\text{local} \\ \forall x, x' \in S, |f(x) - f(x')| \leq L \|x - x'\|$$

- If f is locally Lipschitz $\Rightarrow \partial f$ exists everywhere
- If f is convex $\Rightarrow \partial f = \partial_s f$
- If f is differentiable $\Rightarrow \partial f = \{ \nabla f \}$
Satisfies chain rule

Positive Homogeneity

Definition: $f: \mathbb{R}^d \rightarrow \mathbb{R}$ is positive homogeneous of degree L if $f(\alpha x) = \alpha^L f(x)$ for any $\alpha \geq 0$. L -homogeneous

(1) $L \in \mathbb{Z} \cup \{0\}$: $G(\alpha \cdot z) = \alpha^L \cdot G(z)$

$$G(z) = \max\{0, z\}$$

(2) monomials of degree L : $\prod_{i=1}^d x_i^{p_i}$, $\sum p_i = L$

$$\prod_{i=1}^d (\alpha x_i)^{p_i} = \alpha^{\sum p_i} \cdot \prod_{i=1}^d x_i^{p_i} = \alpha^L \cdot \prod_{i=1}^d x_i^{p_i}$$

(3) Norm: $\|\alpha \cdot x\| = \alpha \cdot \|x\|$

Positive Homogeneity

$$x \in \mathbb{R}^d, w_1 \in \mathbb{R}^{h_1 \times d}$$

$$w_2, \dots, w_H \in \mathbb{R}^{h_i \times h_{i-1}}$$

$$w_{H+1} \in \mathbb{R}^{m \times h_H}$$

(4) Multi-layer ReLU $\cup$ NN

$$f(x, w_1, \dots, w_{H+1}) = w_{H+1} \sigma(w_H \dots \sigma(w_1 x) \dots)$$

1) for each layer: 1-homogeneous

$$\begin{aligned} f(x, w_1, \dots, \alpha w_h, \dots, w_{H+1}) &= w_{H+1} \sigma(\dots \sigma(\alpha w_h \sigma(\dots \sigma(w_1 x) \dots)) \dots) \\ &= \alpha f(x, w_1, \dots, w_{H+1}) \end{aligned}$$

2) for all layers

$$f(x, \alpha w_1, \dots, \alpha w_{H+1}) = \alpha^{H+1} f(x, w_1, \dots, w_{H+1})$$

$\Rightarrow$ $(H+1)$ -homogeneous

Positive Homogeneity

Fact: , $\forall h$

$$\left\langle w_n, \frac{\partial f(x, w_1, \dots, w_{h+1})}{\partial w_n} \right\rangle = \underbrace{f(x, w_1, \dots, w_{h+1})}_{\text{independent of } h}$$

Pf: $A_n = \text{diag} \left(\underbrace{g'(w_n)}_{g': 0 \text{ or } 1} \underbrace{g(w_1 x) \dots g(w_{h+1} x)}_{\text{pattern whether the activation is on or not}} \right) \in \mathbb{R}^{m \times m}$

$$f(x, w_1, \dots, w_{h+1}) = w_{h+1} A_{h+1} w_h \dots A_1 w_1 x$$

$$\frac{\partial f}{\partial w_n} = \underbrace{(w_{h+1} A_{h+1} \dots w_{n+1} A_{n+1})^T}_{1 \times m} \underbrace{(A_{n-1} w_{n-1} \dots w_1 x)^T}_{m \times 1}$$

$g(z) = g'(z) \cdot z$

$$\langle w_n, \frac{\partial f}{\partial w_n} \rangle = \langle w_n, (w_{H+1} A_H \dots w_{H+1} A_n)^T (A_{n-1} \dots w_1 x)^T \rangle$$

$$A^T B^T$$

$$= (BA)^T$$

$$\text{tr}(AB) = \text{tr}(BA)$$

$$= \text{tr}((w_{H+1} A_H \dots w_n)^T (A_{n-1} \dots w_1 x)^T)$$

$$= \text{tr}((A_{n-1} \dots w_1 x)^T (w_{H+1} \dots A_n w_n)^T)$$

$$= \text{tr}(w_{H+1} A_H \dots A_1 w_1 x)$$

$$= f(x, w_1, \dots, w_{H+1})$$

□

Positive Homogeneity and Clark Differential

Clark Differential exists

Lemma: Suppose $f : \mathbb{R}^d \rightarrow \mathbb{R}$ is Locally Lipschitz and L -positively homogeneous. For any $x \in \mathbb{R}^d$ and $s \in \partial f(x)$, we have $\langle s, x \rangle = Lf(x)$.

View $\chi = \bigcup \eta : \neg \text{hom}$

Norm Preservation

$$f(x, w_1, w_2, w_3) = w_3 \sigma(w_2 \sigma(w_1 x)),$$

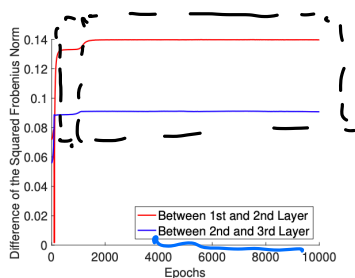
quadratic loss

$$(\|w_1\|_F^2, \|w_2\|_F^2, \|w_3\|_F^2, \|A\|_F^2 = \sum_{i,j} A_{ij}^2)$$

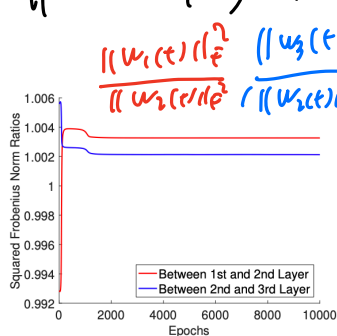
$$w_3 \propto 10$$

$$w_2 / 10$$

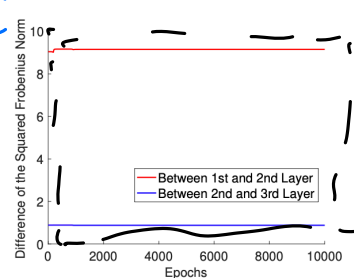
$$w_2(t+1) = w_2(t) + \eta \frac{\partial f}{\partial w_2}$$



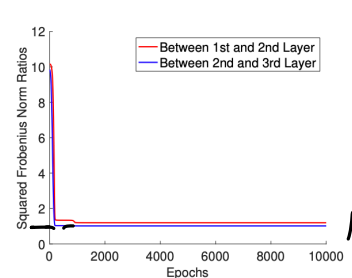
(a) Balanced initialization, squared norm differences.



(b) Balanced initialization, squared norm ratios.



(c) Unbalanced Initialization, squared norm differences.



(d) Unbalanced initialization, squared norm ratios.

$$\bullet \|w_1(0)\|_F^2 \approx \|w_2(0)\|_F^2 \approx \|w_3(0)\|_F^2,$$

$$\bullet \|w_1(t)\|_F^2 - \|w_2(t)\|_F^2$$

$$\bullet \|w_2(t)\|_F^2 - \|w_3(t)\|_F^2$$

$$\|w_1(0)\|_F^2 \neq \|w_2(0)\|_F^2 \text{ small}$$

$$\|w_1(t)\|_F^2 \uparrow \quad \|w_2(t)\|_F^2 \uparrow$$

Gradient flow and gradient inclusion

Discrete-time dynamics can be complex. Let's use continuous-time dynamics to simplify:

$\eta \rightarrow 0$

$$\text{Gradient flow: } x_{t+1} = x_t - \eta \nabla f(x_t) \Rightarrow \frac{dx(t)}{dt} = -\nabla f(x(t))$$

$$\text{Gradient inclusion: } \frac{dx(t)}{dt} \in -\partial f(x(t))$$

$$\frac{x_{t+\eta} - x_t}{\eta} = -\nabla f(x_t)$$

$$\text{let } \eta \rightarrow 0$$

Norm preservation by gradient inclusion

Theorem (Du, Hu, Lee '18) Suppose $\alpha > 0$, $f(x; (W_{H+1}, \dots, \alpha W_i, \dots, W_1)) = \alpha f(x, (W_{H+1}, \dots, W_1))$, i.e., predictions are 1-homogeneous in each layer. Then for every pair of layers $(i, h) \in [H+1] \times [H+1]$, the gradient inclusion maintains: for all $t \geq 0$,

$$\frac{1}{2} \|W_h(t)\|_F^2 - \frac{1}{2} \|W_h(0)\|_F^2 = \frac{1}{2} \|W_{h'}(t)\|_F^2 - \frac{1}{2} \|W_{h'}(0)\|_F^2.$$

$\Leftrightarrow \|W_h(t)\|_F^2 - \|W_{h'}(t)\|_F^2 = \|W_h(0)\|_F^2 - \|W_{h'}(0)\|_F^2$

- If $\|W_{h'}(0)\|_F^2$ small $\rightarrow \|W_h(t)\|_F^2 \approx \|W_{h'}(t)\|_F^2$

Norm preservation by gradient inclusion

Pf: $\text{loss} = \mathcal{L}(f(x, w_1, \dots, w_{H+1}), y)$

$$\frac{dW_n(t)}{dt} = - \mathcal{L}'(f(x, w_1, \dots, w_{H+1}), y) \cdot \frac{\partial f(x, w_1, \dots, w_{H+1})}{\partial w_n}$$

$$\frac{1}{2} \|w_n(t)\|_F^2 = \frac{1}{2} \|w_n(0)\|_F^2$$

$$= \int_0^t \frac{d}{dt} \frac{1}{2} \|w_n(\tau)\|_F^2 d\tau$$

$$= \int_0^t \langle w_n, \frac{dw_n(\tau)}{d\tau} \rangle d\tau$$

(chain rule)

$$= \int_0^t \langle \underbrace{w_n}_{(1)}, \underbrace{- \mathcal{L}'(f(x, w_1, \dots, w_{H+1}), y) \cdot \frac{\partial f(x, w_1, \dots, w_{H+1})}{\partial w_n}}_{(2)} \rangle d\tau$$

independence
of h

$$= \int_0^t \underbrace{\mathcal{L}'(f(x, w_1, \dots, w_{H+1}), y) \cdot f(x, w_1, \dots, w_{H+1})}_{\text{scalar}} d\tau$$

$\mathbb{R}$

Optimization Methods for Deep Learning



Gradient descent for non-convex optimization

$\|A\|_2$: operator norm, largest eigenvalue

Descent Lemma: Let $f: \mathbb{R}^d \rightarrow \mathbb{R}$ be twice differentiable, and $\|\nabla^2 f\|_2 \leq \beta$. Then setting the learning rate $\eta = 1/\beta$, and applying gradient descent, $x_{t+1} = x_t - \eta \nabla f(x_t)$, we have:

$$f(x_t) - f(x_{t+1}) \geq \frac{1}{2\beta} \|\nabla f(x_t)\|_2^2.$$

if $f(x_t) \uparrow$
L2 too large

Pf: by Taylor expansion & mean-value theorem

$$f(x+\Delta) = f(x) + \Delta^T \nabla f(x) + \frac{1}{2} \Delta^T \nabla^2 f(y) \Delta$$

$$\Delta^T \nabla f(y) \Delta \leq \|\nabla^2 f(y)\|_2 \cdot \|\Delta\|_2^2 \leq \beta \|\Delta\|_2^2 \quad \left(y = \lambda x + (1-\lambda)(x+\Delta) \right)$$

$0 \leq \lambda \leq 1$

Let $\Delta = -\eta \cdot \nabla f(x)$

$$\begin{aligned} f(x_{t+1}) &\leq f(x_t) - \eta \|\nabla f(x_t)\|_2^2 + \frac{1}{2} \beta \eta^2 \|\nabla f(x_t)\|_2^2 \\ &= f(x_t) - \frac{1}{2} \cdot \eta \|\nabla f(x_t)\|_2^2 \end{aligned}$$

Converging to stationary points

Theorem: In $T = O(\frac{\beta}{\epsilon^2})$ iterations, we have $\|\nabla f(x)\|_2 \leq \epsilon$.

Gradient Descent for Quadratic Functions

Problem: $\min_x \frac{1}{2} x^\top A x$ with $A \in \mathbb{R}^{d \times d}$ being positive-definite.

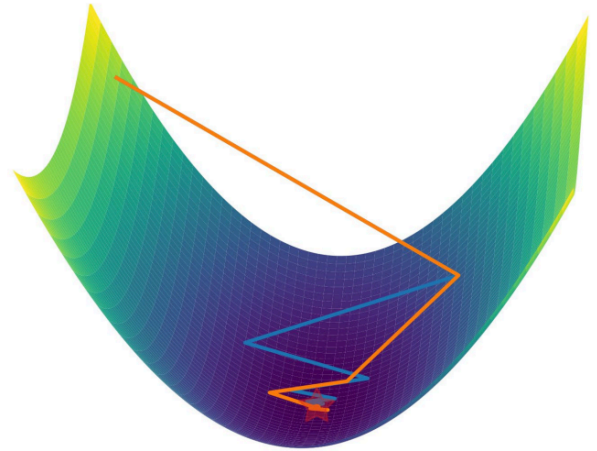
Theorem: Let $\lambda_{\max}$ and $\lambda_{\min}$ be the largest and the smallest eigenvalues of A . If we set $\eta \leq \frac{1}{\lambda_{\max}}$, we have

$$\|x_t\|_2 \leq (1 - \eta \lambda_{\min})^t \|x_0\|_2$$

Momentum: Heavy-Ball Method (Polyak '64)

Problem: $\min_x f(x)$

Method: $v_{t+1} = -\nabla f(x_t) + \beta v_t$
 $x_{t+1} = x_t + \eta v_{t+1}$

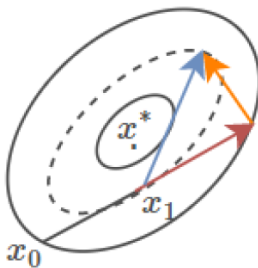


Momentum: Nesterov Acceleration (Nesterov '89)

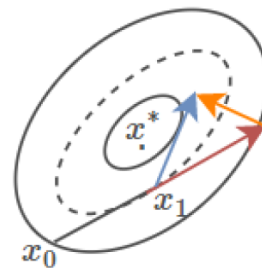
Problem: $\min_x f(x)$

Method: $v_{t+1} = -\nabla f(x_t + \beta v_t) + \beta v_t$
 $x_{t+1} = x_t + \eta v_{t+1}$

Polyak's Momentum

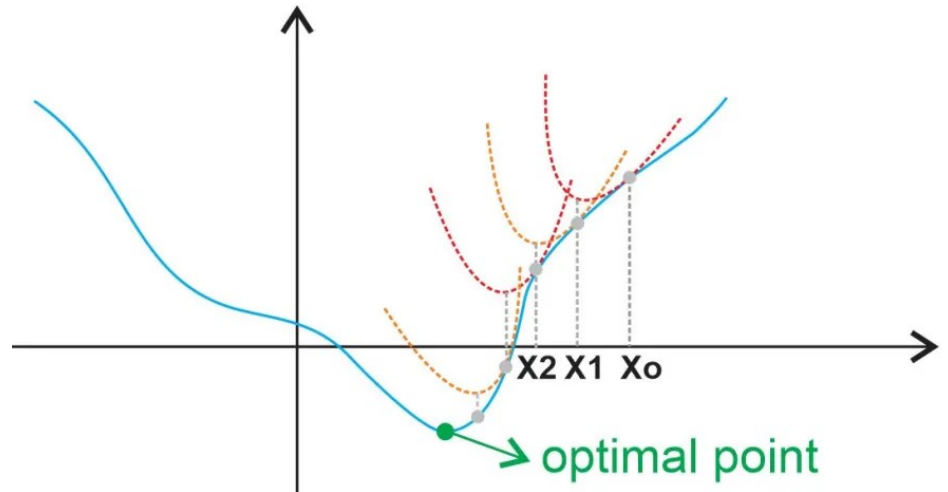


Nesterov Momentum



Newton's Method

Newton's Method: $x_{t+1} = x_t - \eta (\nabla^2 f(x_t))^{-1} \nabla f(x_t)$



AdaGrad (Duchi et al. '11)

Newton Method: $x_{t+1} = x_t - \eta (\nabla^2 f(x_t))^{-1} \nabla f(x_t)$

AdaGrad: separate learning rate for every parameter

$$x_{t+1} = x_t - \eta (G_{t+1} + \epsilon I)^{-1} \nabla f(x_t), (G_t)_{ii} = \sqrt{\sum_{j=1}^{t-1} (\nabla f(x_t)_i)^2}$$

RMSProp (Hinton et al. '12)

AdaGrad: separate learning rate for every parameter

$$x_{t+1} = x_t - \eta(G_{t+1} + \epsilon I)^{-1} \nabla f(x_t), (G_t)_{ii} = \sqrt{\sum_{j=1}^{t-1} (\nabla f(x_t)_i)^2}$$

RMSProp: exponential weighting of gradient norms

$$x_{t+1} = x_t - \eta(G_{t+1} + \epsilon I)^{-1/2} \nabla f(x_t), \\ (G_{t+1})_{ii} = \beta(G_t)_{ii} + (1 - \beta)(\nabla f(x_t)_i)^2$$

AdaDelta (Zeiler '12)

RMSProp:

$$x_{t+1} = x_t - \eta (G_{t+1} + \epsilon I)^{-1/2} \nabla f(x_t),$$
$$(G_{t+1})_{ii} = \beta (G_t)_{ii} + (1 - \beta) (\nabla f(x_t)_i)^2$$

AdaDelta:

$$x_{t+1} = x_t - \eta \Delta x_t,$$
$$\Delta x_t = \sqrt{u_t + \epsilon} \cdot (G_{t+1} + \epsilon I)^{-1/2} \nabla f(x_t)$$
$$(G_{t+1})_{ii} = \rho (G_t)_{ii} + (1 - \rho) (\nabla f(x_t)_i)^2,$$
$$u_{t+1} = \rho u_t + (1 - \rho) \|\Delta x_t\|_2^2$$

Adam (Kingma & Ba '14)

Momentum:

$$v_{t+1} = -\nabla f(x_t) + \beta v_t, \quad x_{t+1} = x_t + \eta v_{t+1}$$

RMSProp: exponential weighting of gradient norms

$$x_{t+1} = x_t - \eta (G_{t+1} + \epsilon I)^{-1} \nabla f(x_t),$$
$$(G_t)_{ii} = \beta (G_t)_{ii} + (1 - \beta) (\nabla f(x_t)_i)^2$$

Adam

$$v_{t+1} = \beta_1 v_t + (1 - \beta_1) \nabla f(x_t)$$
$$(G_{t+1})_{ii} = \beta_2 (G_t)_{ii} + (1 - \beta_2) (\nabla f(x_t)_i)^2$$
$$x_{t+1} = x_t - \eta (G_{t+1} + \epsilon I)^{-1/2} v_{t+1}$$

Default choice nowadays.

Other Optimizers

- AdamW
- NAdam
- RAdam
- GGT
- K-FAC
- ...