

Stochastic Subgradient Method

Lingjie Weng, Yutian Chen

Bren School of Information and Computer Science

UC Irvine

Subgradient

- Recall basic inequality for convex differentiable f :
$$f(y) \geq f(x) + \nabla f(x)^T (y - x)$$
- The gradient $\nabla f(x)$ at x determines a global under-estimator of f .

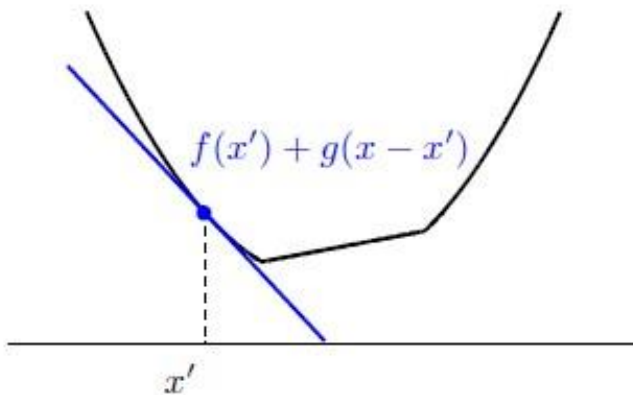
What if f is not differentiable?

- We can still construct a global under-estimator : the vector g is a **subgradient** of f at x if
$$f(y) \geq f(x) + g^T (y - x) \quad \text{for all } y$$
- There can be more than one subgradient at given point: the collection of subgradients of f at x is the subdifferential $\partial f(x)$.

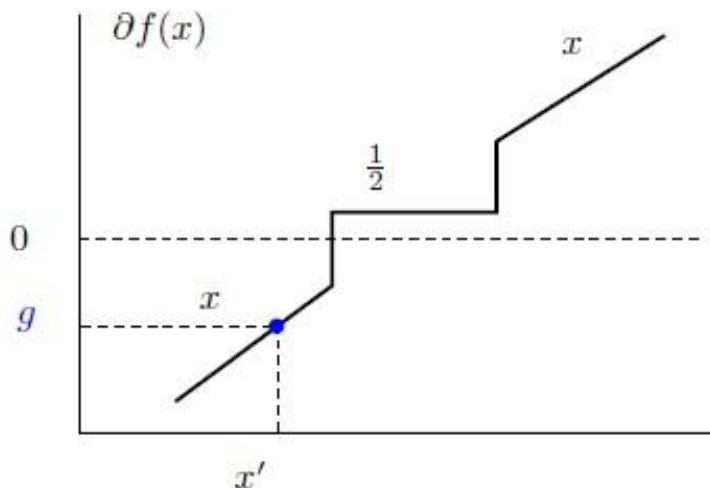
An example: non-differentiable function and its subgradients

$$f = \max\{f_1, f_2\}, \quad \text{with } f_1, f_2 \text{ convex and differentiable}$$

$$f(x) = \max\left\{\frac{1}{2}x^2, \frac{1}{2}x + 2\right\}$$



- $f_1(x') > f_2(x')$: unique subgradient $g = \nabla f_1(x')$
- $f_1(x') < f_2(x')$: unique subgradient $g = \nabla f_2(x')$
- $f_1(x') = f_2(x')$: subgradients form a line segment $[\nabla f_1(x'), \nabla f_2(x')]$



Subgradient methods

- Subgradient method is simple algorithm to minimize nondifferentiable convex function f

$$x^{(k+1)} = x^{(k)} - \alpha_k g^{(k)}$$

$x^{(k)}$ is the k th iterate

$g^{(k)}$ is any subgradient of f at $x^{(k)}$

$\alpha_k > 0$ is the k th step size

- Not a descent method, so we keep track of best point so far

$$f_{best}^{(k)} = \min\{f(x^{(1)}), \dots, f(x^{(k)})\}$$

Why does it work ?--convergence

Assumptions:

- $f^* = \inf_x f(x) > -\infty$, with $f(x^*) = f^*$
- $\|g^{(k)}\|_2 \leq G$ for all k (if f satisfies Lipschitz condition $|f(u) - f(v)| \leq G \|u - v\|_2$)
- $\|x^{(1)} - x^*\|_2 \leq R$

Convergence results for different step size

- Constant step size: $\alpha_k = \alpha$ (positive constant).
 - Converges to a suboptimal solution: $f_{best}^{(k)} - f^* \leq \frac{R^2 + G^2 k \alpha^2}{2k\alpha} \xrightarrow{k \rightarrow \infty} \frac{G^2 \alpha}{2}$
- Constant step length: $\alpha_k = \frac{\gamma}{\|g^{(k)}\|_2}$ (i.e. $\|x^{(k+1)} - x^{(k)}\| = \gamma$).
 - Converges to a suboptimal solution: $f_{best}^{(k)} - f^* \leq \frac{R^2 + \gamma^2 k}{2\gamma k / G} \xrightarrow{k \rightarrow \infty} \frac{G\gamma}{2}$
- Square summable but not summable $\sum_{k=1}^{\infty} \alpha_k^2 < \infty$, $\sum_{k=1}^{\infty} \alpha_k = \infty$ (i.e. $\alpha_k = \frac{a}{b+k}$)
- Nonsummable diminishing step size $\lim_{k \rightarrow \infty} \alpha_k = 0$, and $\sum_{k=1}^{\infty} \alpha_k = \infty$ (i.e. $\alpha_k = \frac{a}{\sqrt{k}}$)
- Nonsummable diminishing step length $\alpha_k = \frac{\gamma k}{\|g^{(k)}\|_2}$
 - Converges to the optimal solution: $\lim_{k \rightarrow \infty} f_{best}^{(k)} - f^* = 0$

Stopping criterion

- Terminating when $\frac{R^2 + G^2 \sum_{i=1}^k \alpha_i^2}{2 \sum_{i=1}^k \alpha_i} \leq \varepsilon$ is really, really slow.
- $\alpha_i = (R/G)/\sqrt{k}$, $i = 1, \dots, k$ minimizes the suboptimality bound.
- The number of steps to achieve a guaranteed accuracy of ε will be at least $(RG/\varepsilon)^2$, no matter what step sizes we use.
- There really isn't a good stopping criterion for the subgradient method
 - are used only for problems in which very high **accuracy** is **not** required, typically around 10%

Subgradient method for unconstrained problem

1. Set $k := 1$ and. Choose an infinite sequence of positive step size value $\{\alpha^{(k)}\}_{k=1}^{\infty}$.
2. Compute a subgradient $g^{(k)} \in \partial f(x^{(k)})$.
3. Update $x^{(k+1)} = x^{(k)} - \alpha_k g^{(k)}$
4. If algorithm has not converged, then set $k := k + 1$ and go to step 2.

Properties:

- ☐ Simple and general. Any subgradient from $\partial f(x^{(k)})$ is sufficient.
- ☐ Convergence theorems exist for properly selected step size
- ☐ Slow convergence. It usually requires a lot of iterations
- ☐ No monotonic convergence. $-g^{(k)}$ need not be decreasing direction $\rightarrow$ line search cannot be used to find optimal α_k .
- ☐ No good stopping condition.

Example: L_2 -norm Support Vector Machines

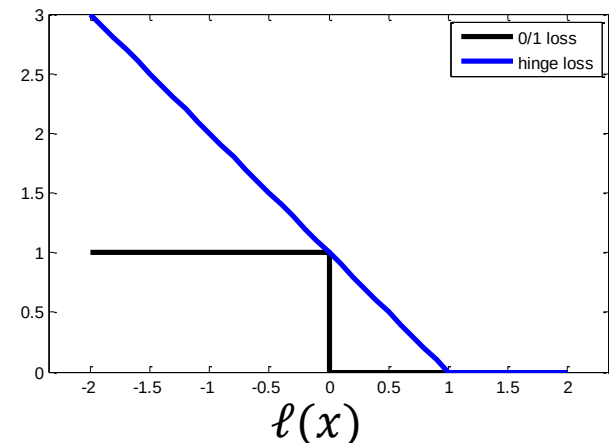
- $(\mathbf{w}^*, b^*, \xi^*) = \underset{\mathbf{w}, b, \xi}{\operatorname{argmin}} \left[\frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^m \xi_i \right],$
s.t. $y_i(\mathbf{w}^T x_i + b) \geq 1 - \xi_i, i = 1, \dots, m$
 $\xi_i \geq 0, \quad i = 1, \dots, m$
- Note that $\xi_i \geq \max\{0, 1 - y_i(\mathbf{w}^T x_i + b)\}$ and the equality is attained in the optimum, i.e., we can optimize
s.t. $\xi_i = \max\{0, 1 - y_i(\mathbf{w}^T x_i + b)\} = \ell(\langle \mathbf{w}, x_i \rangle y_i)$ (hinge loss)
- We can reformulate the problem as an unconstrained convex problem

$$(\mathbf{w}^*, b^*) = \underset{\mathbf{w}, b}{\operatorname{argmin}} \left[\frac{1}{2} \|\mathbf{w}\|^2 + C \ell(\langle \mathbf{w}, x_i \rangle y_i) \right]$$

Example L_2 -norm Support Vector Machines(cont')

- Function $f_0(\mathbf{w}, b) = \frac{1}{2} \|\mathbf{w}\|^2$ is differentiable and thus $\partial_{\mathbf{w}} f_0(\mathbf{w}, b) = \mathbf{w}$, $\partial_b f_0(\mathbf{w}, b) = 0$.
- Function $f_i(\mathbf{w}, b) = \max\{0, 1 - y_i(\mathbf{w}^T x_i + b)\} = \ell(y_i(\mathbf{w}^T x_i + b))$ is a point-wise maximum thus subdifferential is convex combination of gradients of active functions

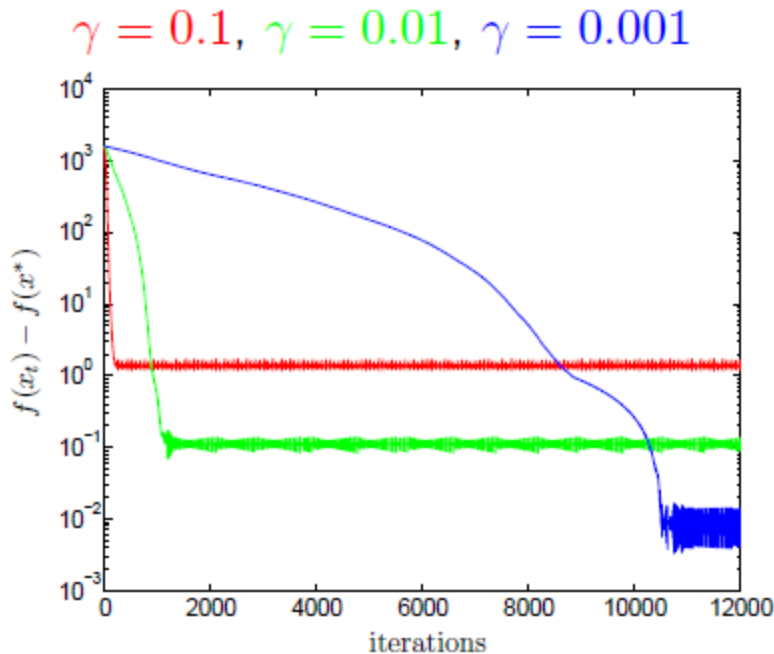
i -th function	$\partial_{\mathbf{w}} f_i(\mathbf{w}, b)$	$\partial_b f_i(\mathbf{w}, b)$
$I_+ = \{i \mid y_i(\mathbf{w}^T x_i + b) > 1\}$	$\mathbf{0}$	0
$I_0 = \{i \mid y_i(\mathbf{w}^T x_i + b) = 1\}$	$\text{Co}\{0, -y_i x_i\}$	$\text{Co}\{0, -y_i\}$
$I_- = \{i \mid y_i(\mathbf{w}^T x_i + b) < 1\}$	$-y_i x_i$	$-y_i$



Constant step length: $\alpha_t = \frac{\gamma}{\|g^t\|}$

Converges to a suboptimal solution:

$$\lim_{t \rightarrow \infty} f_{best}^{(t)} - f^* \leq \frac{G^2 \alpha}{2}$$

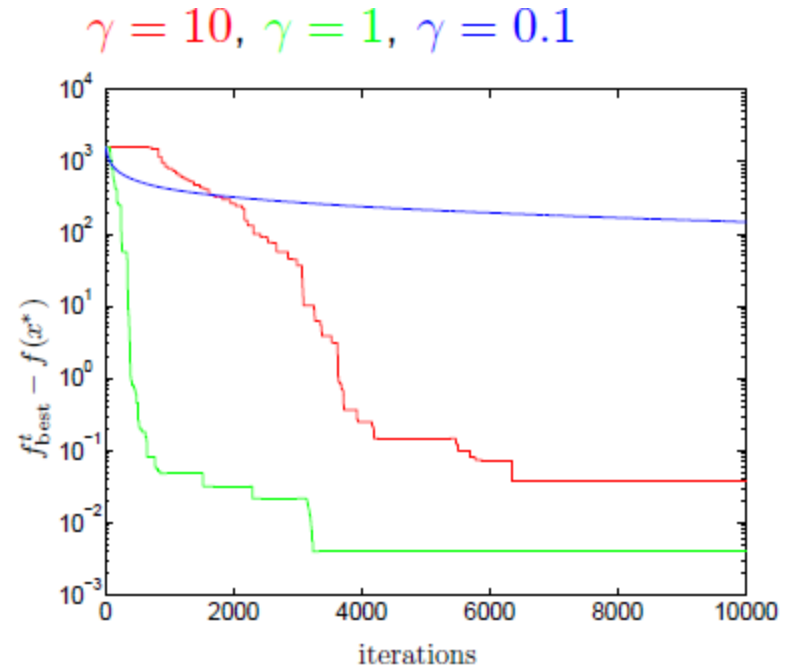


The value of : $f_{best}^{(k)} - f^*$ versus iteration number t

Diminishing step size: $\alpha_t = \frac{\gamma}{t}$

Converges to an optimal solution:

$$\lim_{t \rightarrow \infty} f_{best}^{(t)} - f^* = 0$$



The value of : $f_{best}^{(k)} - f^*$ versus iteration number t

Trade-off: larger γ gives faster convergence, but larger final suboptimality

Subgradient method for constrained problem

- Let us consider the constraint optimization problem
minimize $f(x)$
s.t. $x \in S$ (S is a convex set)
- Projected subgradient method** is given by $x^{(k+1)} = P_S(x^{(k)} - \alpha_k g^{(k)})$
 P_S is (Euclidean) projection on S ,

- Set $k := 1$ and. Choose an infinite sequence of positive step size value $\{\alpha^{(k)}\}_{k=1}^{\infty}$.
- Compute a subgradient $g^{(k)} \in \partial f(x^{(k)})$.
- Update $x^{(k+1)} = P_S(x^{(k)} - \alpha_k g^{(k)})$
- If algorithm has not converged, then set $k := k + 1$ and go to step 2.

Example L_2 -norm Support Vector Machines

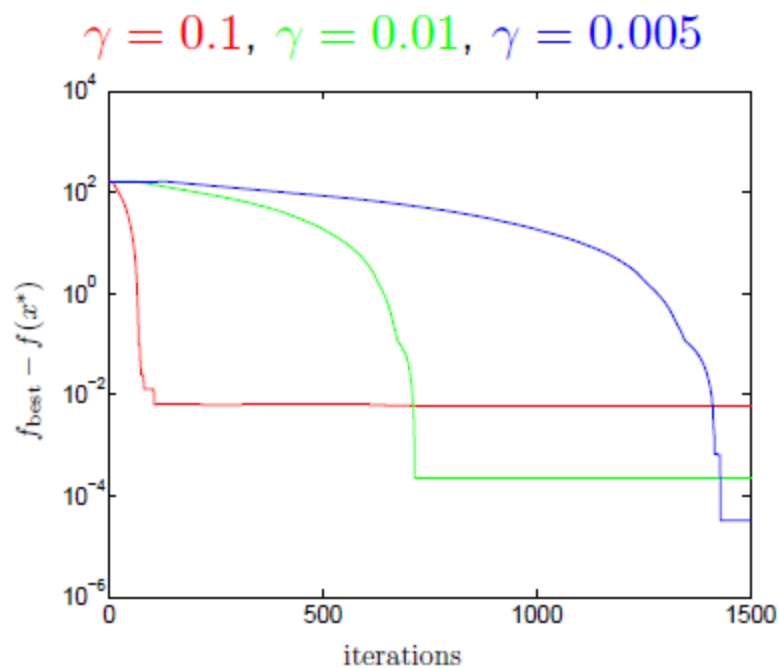
- $(\mathbf{w}^*, b^*) = \operatorname{argmin}_{\mathbf{w}, b} \left[\frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^m \max\{0, 1 - y_i(\mathbf{w}^T x_i + b)\} \right]$
- An alternative formulation of the same problem is
$$(\mathbf{w}^*, b^*) = \operatorname{argmin}_{\mathbf{w}, b} \max\{0, 1 - y_i(\mathbf{w}^T x_i + b)\}$$
$$\text{S.t. } \|\mathbf{w}\| \leq W$$
- Note that the regularization constant $C > 0$ is replaced here by a less abstract constant $W > 0$.

- To implement subgradient method we need only a projection operator

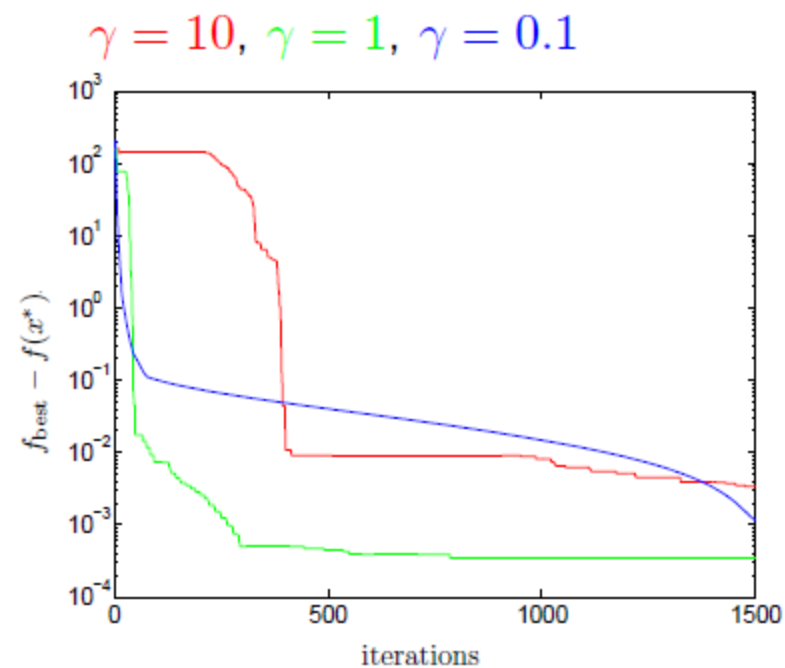
$$P_S(\mathbf{w}, b) = \{\mathbf{w} \rightarrow \frac{\mathbf{w}}{\|\mathbf{w}\|} W, b \rightarrow b\}$$

- So after update $\mathbf{w} := \mathbf{w} - \alpha_k \mathbf{w}'$ and $b := b - \alpha_k b'$,
we project $\mathbf{w} := \frac{\mathbf{w}}{\|\mathbf{w}\|} W$

Constant step length: $\alpha_t = \frac{\gamma}{\|g^t\|}$



Diminishing step size: $\alpha_t = \frac{\gamma}{t}$



Noisy unbiased subgradient

- Random vector $\tilde{g} \in \mathbf{R}^n$ is a **noisy unbiased subgradient** for $f: \mathbf{R}^n \rightarrow \mathbf{R}$ at x if for all z

$$f(z) \geq f(x) + (\mathbf{E} \tilde{g})^T (z - x)$$

i.e., $g = \mathbf{E} \tilde{g} \in \partial f(x)$

- The noise can represent error in computing g , measurement noise, Monte Carlo sampling error, etc
- If x is also random, $\tilde{g}$ is a noisy unbiased subgradient of f at x if

$$\forall z \quad f(z) \geq f(x) + \mathbf{E}(\tilde{g}|x)^T (z - x)$$

holds almost surely

- Same as $\mathbf{E}(\tilde{g}|x) \in \partial f(x)$ (a.s.)

Stochastic subgradient method

- **Stochastic subgradient method** is the subgradient method, using noisy unbiased subgradients

$$x^{(k+1)} = x^{(k)} - \alpha_k \tilde{g}^{(k)}$$

- $x^{(k)}$ is k th iterate
- $\tilde{g}^{(k)}$ is any noisy unbiased subgradient of (convex) f at $x^{(k)}$
- $\alpha_k > 0$ is the k th step size
- Define $f_{best}^{(k)} = \min\{f(x^{(1)}), \dots, f(x^{(k)})\}$

Assumptions

- $f^* = \inf_x f(x) > -\infty$, with $f(x^*) = f^*$
- $\mathbf{E} \|g^{(k)}\|_2^2 \leq G^2$ for all k
- $\mathbf{E} \|x^{(1)} - x^*\|_2^2 \leq R^2$
- Step sizes are square-summable but not summable

$$\alpha_k \geq 0, \sum_{k=1}^{\infty} \alpha_k^2 < \infty, \sum_{k=1}^{\infty} \alpha_k = \infty$$

These assumptions are stronger than needed, just to simplify proofs

Convergence results

- Convergence in expectation:

$$\mathbf{E}f_{best}^{(k)} - f^* \leq \frac{R^2 + G^2 \|\alpha\|_2^2}{2 \sum_{i=1}^k \alpha_i}$$

$$\lim_{k \rightarrow \infty} \mathbf{E}f_{best}^{(k)} = f^*$$

- Convergence in probability (Markov's inequality): for any $\epsilon > 0$,

$$\lim_{k \rightarrow \infty} \mathbf{Prob} \left(f_{best}^{(k)} \geq f^* + \epsilon \right) = 0$$

- Almost sure convergence:

$$\lim_{k \rightarrow \infty} f_{best}^{(k)} = f^* \text{ (a.s.)}$$

Example: Machine learning problem

- Minimizing loss function

$$\min_{\mathbf{w}} \hat{L}(\mathbf{w}) = \frac{1}{m} \sum_{i=1}^m \text{loss}(\mathbf{w}, (x_i, y_i))$$

Subgradient estimate: $\tilde{g}^{(k)} = \nabla_{\mathbf{w}} \text{loss}(\mathbf{w}^{(k)}, (x_i, y_i))$

- Examples: L_2 regularized linear classification with hinge loss, aka. SVM

$$\min_{\|\mathbf{w}\|_2 \leq B} \frac{1}{m} \sum_{i=1}^m \ell(\langle \mathbf{w}, \mathbf{x}_i \rangle y_i) \equiv \min \frac{1}{m} \sum_{i=1}^m \ell(\langle \mathbf{w}, \mathbf{x}_i \rangle y_i) + \frac{\lambda}{2} \|\mathbf{w}\|^2$$

Initialize at $\mathbf{w}^{(0)}$

Iteration k :

$$\mathbf{w}^{(k+1)} \leftarrow \mathbf{w}^{(k)} - \alpha^{(k)} \tilde{g}^{(k)}$$

If $\|\mathbf{w}^{(k+1)}\| > B$

$$\mathbf{w}^{(k+1)} \leftarrow B \frac{\mathbf{w}^{(k+1)}}{\|\mathbf{w}^{(k+1)}\|}$$

Stochastic vs Batch Gradient Descent

$$\min_{\mathbf{w}} \hat{L}(\mathbf{w}) = \frac{1}{m} \sum_{i=1}^m \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_i, y_i))$$

$\mathbf{w} \leftarrow \mathbf{w} - \mathbf{g}_1$	$\mathbf{g}_1 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_1, y_1))$	$\mathbf{x}_1, y_1$	$\mathbf{g}_1 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_1, y_1))$	$\left. \begin{array}{l} \mathbf{g}_1 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_1, y_1)) \\ \mathbf{g}_2 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_2, y_2)) \\ \mathbf{g}_3 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_3, y_3)) \\ \mathbf{g}_4 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_4, y_4)) \\ \mathbf{g}_5 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_5, y_5)) \end{array} \right\} \nabla \hat{L}(\mathbf{w}) = \frac{1}{m} \mathbf{g}_i$
$\mathbf{w} \leftarrow \mathbf{w} - \mathbf{g}_2$	$\mathbf{g}_2 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_2, y_2))$	$\mathbf{x}_2, y_2$	$\mathbf{g}_2 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_2, y_2))$	
$\mathbf{w} \leftarrow \mathbf{w} - \mathbf{g}_3$	$\mathbf{g}_3 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_3, y_3))$	$\mathbf{x}_3, y_3$	$\mathbf{g}_3 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_3, y_3))$	
$\mathbf{w} \leftarrow \mathbf{w} - \mathbf{g}_4$	$\mathbf{g}_4 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_4, y_4))$	$\mathbf{x}_4, y_4$	$\mathbf{g}_4 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_4, y_4))$	
$\mathbf{w} \leftarrow \mathbf{w} - \mathbf{g}_5$	$\mathbf{g}_5 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_5, y_5))$	$\mathbf{x}_5, y_5$	$\mathbf{g}_5 = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_5, y_5))$	
		$\vdots$		
$\mathbf{w} \leftarrow \mathbf{w} - \mathbf{g}_{m-1}$	$\mathbf{g}_m = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_m, y_m))$	$\mathbf{x}_m, y_m$	$\mathbf{g}_m = \nabla \text{loss}(\mathbf{w} \text{ on } (\mathbf{x}_m, y_m))$	$\mathbf{w} \leftarrow \mathbf{w} - \sum \mathbf{g}_i$
$\mathbf{w} \leftarrow \mathbf{w} - \mathbf{g}_m$				

$$\bar{\mathbf{w}} = \frac{1}{m} \sum_{i=1}^m \mathbf{w}_i$$

Stochastic vs Batch Gradient Descent

- Runtime to guarantee an error ϵ ($\alpha^{(k)} = \frac{B/G}{\sqrt{k}}$)
 - Stochastic Gradient Descent: $\mathcal{O}\left(\frac{G^2 B^2}{\epsilon^2} d\right)$
 - Batch Gradient Descent: $\mathcal{O}\left(\frac{G^2 B^2}{\epsilon^2} m d\right)$
- Optimal runtime if only using gradients and only assuming Lipchitz condition.
- Slower than second order method

Stochastic programming

$$\begin{array}{ll}\text{minimize} & \mathbf{E} f_0(x, \omega) \\ \text{subject to} & \mathbf{E} f_i(x, \omega) \leq 0, i = 1, \dots, m\end{array}$$

If $f_i(x, \omega)$ is convex in x for each ω , problem is convex

- “certainty-equivalent” problem

$$\begin{array}{ll}\text{minimize} & f_0(x, \mathbf{E}\omega) \\ \text{subject to} & f_i(x, \mathbf{E}\omega) \leq 0, i = 1, \dots, m\end{array}$$

(if $f_i(x, \omega)$ is convex in ω , gives a lower bound on optimal value of stochastic problem)

Machine learning as stochastic programming

Minimizing generalization error:

$$\min_{\mathbf{w} \in \mathcal{W}} E[f(\mathbf{w}, z)]$$

- Two approaches:
 - Sample Average Approximation (SAA):
 - Collect sample $z_1, \dots, z_m$
 - Minimize $\hat{F}(\mathbf{w}) = \frac{1}{m} \sum_{i=1}^m f(\mathbf{w}, z_i)$, $\hat{\mathbf{w}} = \min_{\mathbf{w}} \hat{F}(\mathbf{w})$
 - Minimizing empirical risk
 - Stochastic/batch gradient descent, Newton method, ...
 - Sample Approximation (SA):
 - Update $\mathbf{w}^{(k)}$ based on weak estimator to $\nabla F(\mathbf{w}^{(k)})$
 $\tilde{\mathbf{g}}^{(k)} = \nabla f(\mathbf{w}^{(k)}, z_k)$
 - Stochastic gradient descent, $\bar{\mathbf{w}}^{(m)} = \frac{1}{m} \sum_{i=1}^m \mathbf{w}_i$

Machine learning as stochastic programming

Minimizing generalization error:

$$\min_{\mathbf{w} \in \mathcal{W}} E[f(\mathbf{w}, z)]$$

$$\min_{\mathbf{w}} \hat{L}(\mathbf{w}) = \frac{1}{m} \sum_{i=1}^m \text{loss}(\mathbf{w}, (x_i, y_i))$$

- Two approaches:
 - Sample Average Approximation (SAA):
 - Collect sample $z_1, \dots, z_m$
 - Minimize $\hat{F}(\mathbf{w}) = \frac{1}{m} \sum_{i=1}^m f(\mathbf{w}, z_i)$, $\hat{\mathbf{w}} = \min_{\mathbf{w}} \hat{F}(\mathbf{w})$
 - Minimizing empirical risk
 - Stochastic/batch gradient descent, Newton method, ...
 - Sample Approximation (SA):
 - Update $\mathbf{w}^{(k)}$ based on weak estimator to $\nabla F(\mathbf{w}^{(k)})$
 $\tilde{\mathbf{g}}^{(k)} = \nabla f(\mathbf{w}^{(k)}, z_k)$
 - Stochastic gradient descent, $\bar{\mathbf{w}}^{(m)} = \frac{1}{m} \sum_{i=1}^m \mathbf{w}_i$

SAA vs SA

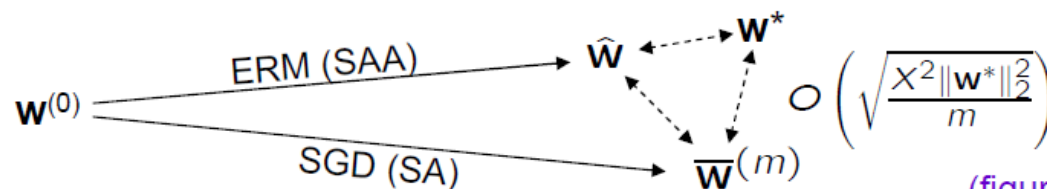
- Runtime to guarantee an error ϵ ($L(\mathbf{w}) \leq L(\mathbf{w}^*) + \epsilon$)
 - SAA (Empirical Risk Minimization)
 - Sample size to guarantee $L(\hat{\mathbf{w}}) \leq L(\mathbf{w}^*) + \epsilon$: $m = \Omega(\frac{G^2 B^2}{\epsilon^2})$
 - Using any method (stochastic/batch subgradient descent, Newton method, ...), runtime:

$$\Omega\left(\frac{G^2 B^2}{\epsilon^2} d\right)$$

- SA (Single-Pass Stochastic Gradient Descent)

- After k iteration, $L(\bar{\mathbf{w}}^{(k)}) \leq L(\mathbf{w}^*) + \mathcal{O}\left(\sqrt{\frac{G^2 B^2}{k}}\right)$
- Runtime:

$$\mathcal{O}\left(\frac{G^2 B^2}{\epsilon^2} d\right)$$



(figure adapted from Leon Bottou)

Reference

- S. Boyd, et. al. Subgradient Methods, Stochastic Subgradient Methods. Lecture slides and notes for EE364b, Stanford University, Spring Quarter 2007–08
- Professor Vojtech Franc's talks. A Short Tutorial on Subgradient Methods. IDA retreat, May 16, 2007
- Nathan Srebro and Ambuj Tewari, Stochastic Optimization: ICML 2010 Tutorial, July, 2010