

today: BCFW
 • stat theory?

Block-coordinate optimization:

"huge scale" optimization $\rightarrow$ [Nesterov 2010-2012]

proposed: block-coordinate projected gradient method

optimization setup: $\min f(x)$
 $\text{s.t. } x \in \bigtimes_{i=1}^n M_i$
 i.e. $x = (x_1, \dots, x_n)$
 blocks

$\rightarrow$ pick i at random

$$\text{then let } \begin{cases} x_i^{(t+1)} = \text{Proj}_{M_i} \left(x_i^{(t)} - \frac{1}{L_i} \nabla_i f(x^{(t)}) \right) \\ x_j^{(t+1)} = x_j^{(t)} \text{ for } j \neq i \end{cases}$$

$\nabla = \begin{pmatrix} \vdots \\ \nabla_i \\ \vdots \end{pmatrix}$
 ∇_i Lipschitz constant for $\nabla_i f(x)$

(only update update block i at iteration t)

$$\text{Nesterov should get } \mathbb{E} f(x^{(t)}) - f(x^*) \leq \frac{2}{t+4} \left[\frac{1}{n} \sum_i L_i \right] \|x_0 - x^*\|^2 \quad (\text{convex } f)$$

Block-coordinate FW: idea: do a FW step on block i

algorithm for $t=0, \dots$

pick i at random

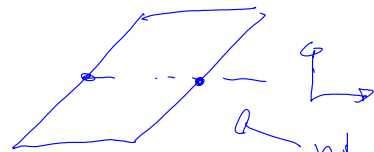
$$\text{let } S_i^t = \underset{S \in M_i}{\text{argmin}} \langle \nabla_i f(x^{(t)}), S \rangle$$

(FW corner for block i) $S_{i,j} = \begin{pmatrix} 0 \\ S_i \\ 0 \end{pmatrix}$

$$x_i^{(t+1)} = x_i^{(t)} (1 - \gamma_t) + \gamma_t S_i^t$$

$$x_j^{(t+1)} = x_j^{(t)} \text{ for } j \neq i$$

$\gamma_t =$ line search:
 $\underset{\gamma \in [0,1]}{\text{argmin}} f(x^{(t)} + \gamma(S_{i,j}^{(t+1)} - x_{i,j}^{(t)}))$



not a product space?

$2n$ # of blocks
 $L+2n$

$\sum_{i=1}^n M_i$ nodes like



an important property: FW gap = $\max_{S \in M} \langle -\nabla f(x), S - x \rangle = \max_{S \in M} \sum_i \langle -\nabla_i f(x), S_i - x_i \rangle$

$$g^{FW}(x) = \sum_i^{BCFW} g_i(x)$$

product structure $\geq \sum_i \max_{S_i \in M_i} \langle -\nabla_i f(x), S_i - x_i \rangle$

$g_i(x)$ "block FW gap"

as before, can show that

$$g_i(x) \geq f(x) - \min_{\substack{y \text{ s.t.} \\ y_j = x_j \forall j \neq i}} f(y)$$

convergence result: define

$$C_f^{(i)} \triangleq \sup_{\substack{\gamma \in (0,1] \\ S_i \in M_i \\ x \in M \\ x_\gamma \triangleq x + \gamma(S_{[i]} - x_{[i]})}} \frac{\gamma}{2} (f(x_\gamma) - f(x) - \gamma \langle \nabla f(x), S_i - x_i \rangle)$$

(only one block update)

$$C_f^{(i)} \leq L_i \text{diam}(M_i)^2$$

"block descent lemma"

during BCFW, if block i is updated:

$$f(x_\gamma) \leq f(x) + \gamma \underbrace{\langle \nabla f(x), S_i - x_i \rangle}_{-g_i(x)} + \frac{\gamma^2}{2} C_f^{(i)}$$

consider γ to be fixed γ_t

$$\mathbb{E}[f(x^{(t+1)}) | x^{(t)}] \leq f(x^{(t)}) - \gamma \underbrace{\mathbb{E}[g_i(x) | x]}_{FW} + \frac{\gamma^2}{2} \underbrace{\mathbb{E}[C_f^{(i)}]}_{FW}$$

$$\Rightarrow \epsilon_{t+1} \leq (1 - \frac{\gamma}{n}) \epsilon_t + \frac{\gamma^2}{2n} C_f^2 \quad \frac{1}{n} g(x) \quad \frac{1}{n} \sum_{i=1}^n C_f^{(i)}$$

→ do induction ...

$$\mathbb{E}[f(x^{(t)})] - f(x^*) \leq \frac{2n[C_f^2 + f(x^{(0)}) - f(x^*)]}{t+2n}$$

$$\text{for } \gamma_t = \frac{2n}{t+2n}$$

if you use line search:

$$\mathbb{E} f(x^{(t)}) - f(x^*) \leq \frac{2n C_f^2}{t - t_0 + 2n}$$

for $t \geq t_0$

$$\text{where } t_0 \leq n \log \frac{2\epsilon_0}{C_f^2}$$

↑ twice to ensure that $\epsilon_t \leq C_f^2$

using from $\epsilon_t \leq \epsilon_0 \exp(-t/n) + \frac{C_f^2}{2}$
when $\epsilon_t \geq C_f^2$

one can show that $C_f^2 \leq C_f$ for quadratic functions

compare batch FW
vs BCFW with
line search

$$\frac{2C_f}{t+2} \quad \frac{2n C_f^2}{t - t_0 + 2n}$$

(say $t_0 = 0$ for simplicity)

BCFW update is n times cheaper than batch FW

⇒ BCFW is "never" slower than batch FW

i.o. # of black oracle call)
to reach ϵ -error

$$\text{FW: } O\left(\frac{n C_f}{\epsilon}\right)$$

$$\text{BCFW: } O\left(\frac{n C_f^2}{\epsilon}\right)$$

extensions : • non-uniform sampling eg using $\{C_f^{(i)}\}$

$$\{g^{(i)}(x)\}$$

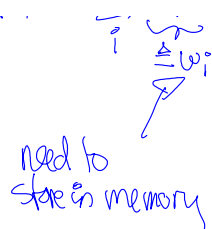
• using away steps etc.. to "get" linear convergence

} ICHL 2016

Applications to SVM struct:

• getting S_i is one loss-augmented decoding call on example i

$$\bullet w = Ax = \sum A_i x_i$$



when you do a BCFW step,

you update $w_i^{(t+1)} = w_i^{(t)} + \gamma (w_S - w_i^{(t)})$

for SVM struct $C_S^{(1)} = \frac{4B^2}{\lambda n^2} \Rightarrow C_S^{(k)} = \frac{4B^2}{\lambda n} = \frac{C_S}{n}$

Theory basics:

decision theory setup:

estimate: $h_w: X \rightarrow Y$

task loss

generalization error = $L_P(w) \triangleq \mathbb{E}_{(x,y) \sim P} [l(y, h_w(x))]$

ultimate goal is find $w^* = \arg \min_{w \in W} L_P(w)$

problem is do not know P (distribution on (x,y))

suppose $\underbrace{(x^{(i)}, y^{(i)})}_{D_n}_{i=1}^n$ iid P training data

$\leadsto$ we can look at $\hat{L}_n(w) = \frac{1}{n} \sum_{i=1}^n l(y^{(i)}, h_w(x^{(i)}))$

learning algorithm $\hat{w}_n = \underset{\text{algorithm}}{A}(D_n)$

from statistics/probs

$\hat{L}_n(w) \xrightarrow{a.s.} L(w) \quad \forall w$

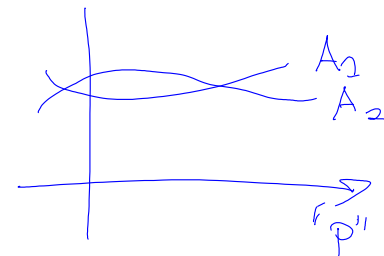
note that this is much weaker
that $\sup_w |\hat{L}_n(w) - L(w)| \xrightarrow{n \rightarrow \infty} 0$

/note that minimizing training error gives not much guarantees in general

regression $\times \times \times \times$

for n points, can get zero training error with polynomial of degree $\geq n-1$

suppose you have no noise
real fct. is quadratic; $\times \times \times \times$



Learning theory, we want to understand properties of learning algorithms
in particular, what can I say about $L_p(A(D_n))$?

"frequentist risk" $\mathbb{E}_{D_n \sim p} [L_p(A(D_n))] = R_p^F(A)$
 D_n is random

PAC framework
"Probably Approximately Correct"

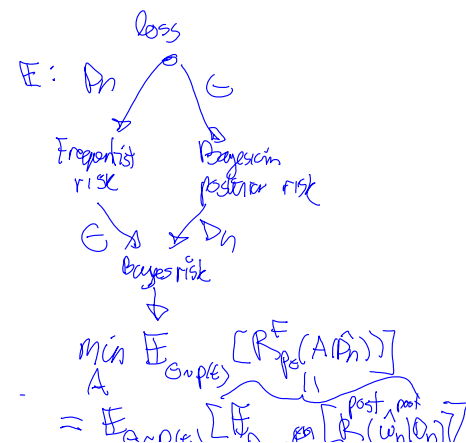
$P\{L_p(A(D_n)) > \text{some bound}\} \leq \delta$

[which means $L_p(A(D_n)) \leq \text{bound}$ with prob bigger than $1-\delta$ on random D_n]
generalization error bound

"Bayesian posterior risk" $R(w|D_n)^{\text{post}} = \mathbb{E}_{\theta \sim p(\theta|D_n)} [L_p(w)]$
prior $p(\theta)$ over distributions
posterior $p(\theta|D_n)$

Bayesian estimate $\hat{w}_n = \arg \min_w R(w|D_n)^{\text{post}}$

could analyze the property of $R_p^F(A(D_n))$ i.e. $R_p^F(\hat{w}_n)$



next time:

Occam's bound:

$\forall w \quad L_p(w) \leq \hat{L}(w) + \text{"complexity}(w)"$ with prob $\geq 1-\delta$

$$\forall w \quad L_p(w) \leq L(w) + \text{"complexity}(w) \quad \text{with prob} \geq 1-\delta$$

$$= \mathbb{E}_{\Theta \sim p(\Theta)} \left[\mathbb{E}_{D_n \sim p(D_n)} \left[\mathbb{E}_{\Theta \sim p(\Theta)} [R(w|D_n)] \right] \right]$$

$p(\Theta)$ ← this is specifying both $p(D_n|\Theta)$ & $p(\Theta)$
(in non-parametric view)

Pointers:

- Nesterov coordinate method:
 - Efficiency of Coordinate Descent Methods on Huge-Scale Optimization Problems, Y. Nesterov, SIAM J. Optim., 22(2), 341–362. 2012
[epubs.siam.org/doi/pdf/10.1137/100802001](https://pubs.siam.org/doi/pdf/10.1137/100802001)
- BCFW in all its gory details:
 - Block-Coordinate Frank-Wolfe Optimization for Structural SVMs, S. Lacoste-Julien, M. Jaggi, M. Schmidt and P. Pletscher, ICML 2013
<http://jmlr.csail.mit.edu/proceedings/papers/v28/lacoste-julien13-suppl.pdf>
code: <https://github.com/ppletscher/BCFWstruct>
- Improvements of BCFW for SVMstruct (non-uniform sampling, away steps, etc.):
 - Minding the Gaps for Block Frank-Wolfe Optimization of Structured SVMs, A. Osokin, J.-B. Alayrac, I. Lukasewitz, P. Dokania and S. Lacoste-Julien, ICML 2016
<https://arxiv.org/abs/1605.09346>
code: <http://www.di.ens.fr/sierra/research/gapBCFW/>