

Lecture 12 - scribbles - FW for SVMstruct

Wednesday, February 15, 2017
10:45

today: FW & BCFW for structured SVM!

recall: batch FW on $\min_{\alpha \in A, \|\alpha\|_1} \underbrace{\frac{\lambda \|\alpha\|_1^2}{2} - b^T \alpha}_{\text{minus dual objective } -d(\alpha)}$

$(\nabla S(\alpha))_{i,y} = -\frac{1}{n} H(y \cdot w^{(k)}) \quad w^{(k)} = A\alpha$
 $= (\lambda A^T A - b)_{i,y}$

FW step on $M = \sum_{i=1}^n \Delta_{\|y_i\|}$ loss-augmented inference $S_i^{(k)} = S_{\hat{y}_i^{(k)}} \quad \text{where } \hat{y}_i^{(k)} = \arg \max_{\tilde{y} \in \mathcal{Y}_i} H_i(\tilde{y}; w^{(k)})$

recall primal $p(w) = \frac{\lambda \|w\|^2}{2} + \sum_{i=1}^n H_i(w)$

$p(w, \hat{y}^*(w)) = \frac{\lambda \|w\|^2}{2} + \sum_{i=1}^n \max_{\tilde{y} \in \mathcal{Y}_i} \ell_i(\tilde{y}) - w^T \ell_i(\tilde{y})$

$\frac{\partial p}{\partial w^{(k)}} = \lambda w^{(k)} - \sum_{i=1}^n \psi_i(\hat{y}_i^{(k)})$
 a subgradient of p at $w^{(k)}$

from $\alpha^{(k+1)} = (1-\gamma) \alpha^{(k)} + \gamma S^{(k)}$
 $w^{(k+1)} = (1-\gamma) w^{(k)} + \gamma \sum_{i=1}^n \psi_i(\hat{y}_i^{(k)})$
 $w^{(k+1)} = (1-\beta\lambda) w^{(k)} + \beta \sum_{i=1}^n \psi_i(\hat{y}_i^{(k)})$

thus batch FW on dual with step-size $\gamma = \beta\lambda$ is equivalent to batch subgradient on primal with step-size $\beta = \frac{\gamma}{\lambda}$

advantage of FW is can use line search (adaptive step-size subgradient method)

note: subgradient convergence proof for μ -strongly convex obj.

needed $\beta \leq \frac{1}{\mu} O(\frac{1}{t})$

$R \triangleq \max_{i,y} \|\psi_i(y)\|_2$

canonical FW step-size $\gamma_t = \frac{2}{t+2} \Rightarrow \beta_t = \frac{1}{\lambda} \frac{2}{t+2}$

step-size in lecture 5?

$p(\hat{w}^{(t)}) - p(w^*) \leq \frac{8R^2}{\lambda(t+2)}$

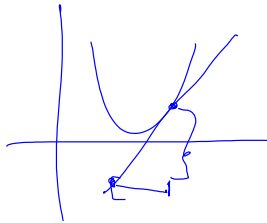
$\hat{w}^{(t)} \triangleq \sum_{s=0}^t \beta_s^t w^{(s)} \quad (\text{weighted average})$

$\min_{S \subseteq [t]} p(w^{(s)}) \leq \sum_{s \in S} \beta_s^t p(w^{(s)})$
 convexity $\leq p(\sum_{s \in S} \beta_s^t w^{(s)})$

for any convex comb.

FW analysis: $d(\alpha^*) - d(\alpha^{(t)}) \leq \frac{2\zeta}{t+2}$; turns out that $\zeta = \frac{4p^2}{\lambda}$ (later today)

let's look at FW gap & $\langle -Df(\alpha^{(t)}), s^{(t)} - \alpha^{(t)} \rangle = \frac{1}{n} \sum_{i=1}^n (H_i(\hat{y}_i^{(t)}; w^{(t)}) - \sum_{y \in \mathcal{Y}_i} \alpha_i(y) H_i(y; w^{(t)}))$



$$= p(w(\alpha^{(t)})) - d(\alpha^{(t)})$$

(Lagrangian gap of lecture 9?)

[note: not always the case;
see e.g. CRF objective]

FW gaps
* one can show, that $\min_{s \leq t} \zeta_s \leq 3 \cdot \frac{2\zeta}{t+2}$

it's also true that $g^{FW}(\hat{\alpha}^{(t)}) \leq 3 \frac{2\zeta}{t+2}$ when $g(\cdot)$ is convex

$\hookrightarrow$ batch FW on dual

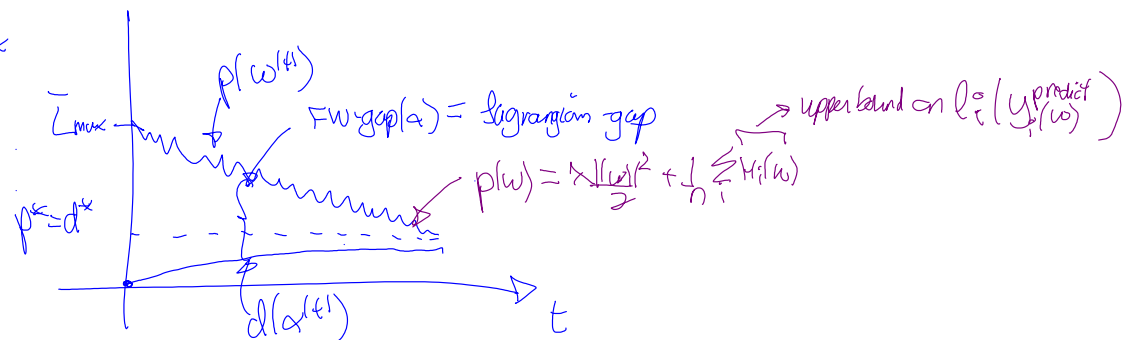
for svm struct $g^{FW}(\hat{\alpha}^{(t)})$

this is case for quadratic functions

also gives $p(\hat{w}^{(t)}) - p(w^*) = O(\frac{1}{t})$

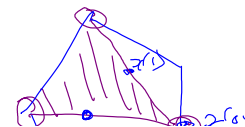
$$p(o) = \frac{1}{n} \sum_{i=1}^n \max_{\tilde{y}} \ell_i(\tilde{y}) = \bar{L}_{\max}$$

$$d(\alpha^{(0)}) = 0$$



another variant of FW: FCFW "fully-corrective FW"

algorithm: re-optimize f over convex hull $(\sum_{s=0}^t S^{(u)})$



turns out that

(batch) FCFW on dual of SVM struct
is equivalent to the
constraint generation / cutting plane approach
for the 1-stack formulation

(see lecture 9)



why?

every $S^{(t)}$ corresponds to $\sum_{i=1}^n \hat{y}_i^{(t)} S_i^{(n)}$

$$\alpha \in \text{convex-hull}(\{S^{(u)}\}_{u \leq t}) \Rightarrow \alpha = \sum_{u \leq t} \tilde{\alpha}_u S^{(u)}$$

$\tilde{\alpha}$ belongs to a subset

$$w = A\alpha = \sum_{u \leq t} \tilde{\alpha}_u A S^{(u)} = \frac{1}{\lambda} \left(\sum_{i=1}^n \eta_i / \hat{y}_i^{(u)} \right) \Delta \left(\sum_{i=1}^n \eta_i S_i \right)$$

to recap:

	primal problem	dual problem	
rate	$O(n)$ iteration cost	batch FW γ	$\otimes$ can do line search γ_t
$O(\frac{R^2}{\lambda \epsilon})$	batch subgradient $\beta = \frac{\gamma}{\lambda}$		
	1-stack cutting plane	FCFW	[see NIPS 2015 get linear convergence rate?]
$O(1)$ iteration cost	stochastic subgradient (no line search)	BCFW	line search

complexity in # of oracle calls

$$O\left(\frac{nR^2}{\lambda \epsilon}\right)$$

at least 2 $\rightarrow n^2 \log(\frac{1}{\epsilon})$

$$O\left(\frac{R^2}{\lambda \epsilon}\right)$$

pointers:

- the usual pointer for FW for structured SVM (as in previous lecture):
 - 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-supp.pdf>
- the linear convergence of FCFW is given in:
 - On the Global Linear Convergence of Frank-Wolfe Optimization Variants, S. Lacoste-Julien and M. Jaggi, NIPS 2015
<http://arxiv.org/abs/1511.05932> | [NIPS link](#)

(already cited in last lectures)

- a good application of using FCFW when the linear oracle is expensive is given in this paper:
Barrier Frank-Wolfe for Marginal Inference, R. Krishnan, S. Lacoste-Julien and D. Sontag, NIPS 2015
<https://arxiv.org/abs/1511.02124>

- the generalization of the observation that Frank-Wolfe optimization sometimes reduced to the subgradient method on the primal is given in:
 - F. Bach. Duality between subgradient and conditional gradient methods.
SIAM Journal of Optimization, 25(1):115-129, 2015
http://www.di.ens.fr/~fbach/sg_cg_fbach_siopt.pdf