

Lecture 20 - variance reduction

Thursday, March 21, 2019 13:27

today: - variance reduction perspective
- application CRT

Variance reduction idea

X, Y be R.V.s

goal: estimate $\mathbb{E}X$ using M.C. samples

suppose: $\mathbb{E}Y$ is cheap to compute and Y is correlated with X

consider estimator $\Theta_\alpha \triangleq \alpha(X - Y) + \mathbb{E}Y$ to approximate $\mathbb{E}X$
 $\alpha \in [0, 1]$

properties: $\mathbb{E}\Theta_\alpha = \alpha \mathbb{E}X + (1-\alpha) \mathbb{E}Y \rightarrow$ unbiased (i.e. $\mathbb{E}\Theta_\alpha = \mathbb{E}X$)
if $\mathbb{E}Y = \mathbb{E}X$ (not interesting)
 $\alpha = 1$

variance:
$$\text{Var}(\Theta_\alpha) = \alpha^2 [\underbrace{\text{Var}(X) + \text{Var}(Y) - 2 \text{cov}(X, Y)}_{\text{variance reduction}}]$$

for $\alpha = 1$
(unbiased setting)

$$\Theta_\alpha = X + \underbrace{(\mathbb{E}X - Y)}_{\text{correction}}$$

SGD setting:

SGD setting :

X is $\nabla f_i(x_t)$; $\mathbb{E}X = \text{batch gradient}$

SAG/SAGA algorithm : Y is g_i [past stored gradient]

$$\mathbb{E}Y = \frac{1}{n} \sum g_i$$

SAG algorithm : $\alpha = \frac{1}{n}$ (biased)

SAGA : $\alpha = 1$ (unbiased)

$$\text{SAG: } x_{t+1} = x_t - \alpha \left[\frac{1}{n} \nabla f_{i_t}(x_t) - g_{i_t}^t + \frac{1}{n} \sum_j g_j^t \right] \quad (\text{biased})$$

$$\text{SAGA: } x_{t+1} = x_t - \alpha \left[\nabla f_{i_t}(x_t) - g_{i_t}^t + \frac{1}{n} \sum_j g_j^t \right] \quad (\text{unbiased})$$

$$\text{SVRG: } x_{t+1} = x_t - \alpha \left[\nabla f_{i_t}(x_t) - \nabla f_{i_t}(x_{\text{old}}) + \frac{1}{n} \sum_{j=1}^n \nabla f_j(x_{\text{old}}) \right] \quad (\text{unbiased})$$

(stochastic variance reduced gradient)

x_{old} is updated from outer loop

SVRG algorithm :

SVRG algorithm:

for $k=0, \dots$ (outer loop)

compute $g_{\text{ref}} \triangleq \frac{1}{n} \sum_j \nabla f_j(x^{(k)})$

for $t=0, \dots, T_{\text{max}}$

sample i_t

$$x_{t+1}^{(k)} = x_t^{(k)} - \gamma [\nabla f_{i_t}(x_t^{(k)}) - \nabla f_{i_t}(x^{(k)}) + g_{\text{ref}}]$$

end

$$x^{(k+1)} = x^{(k)}$$

end

questions:

• what is T_{max} ?

• what is γ ?

original SVRG convergence result: need $\gamma \leq \alpha \frac{1}{L}$

" $T_{\text{max}} \geq \tilde{\kappa} \frac{L}{\mu} = k$ " $\rightarrow$ to run alg., need to know $\tilde{\kappa}$;
 $\Rightarrow$ not adaptive to local strong convexity

break on SVRG:

[Hoffmann et al. NIPS 2015]

$$T_{\text{max}} \sim \text{Geom}(\sim)$$

[at inner loop iterations, do batch gradient comp.]

with prob. $\frac{1}{n}$

then, get same convergence result as SAGA

↳ here, size of inner loop $E[i_{\max}] = n$

overall cost of SVRG $\approx 3 \cdot (\text{SGD for each } n \text{ updates})$

SAG / SAGA / SVRG ^{Karagann}, rate for convex f^* . ($\mu=0$), get $\min_t [E(f_t) - f^*] = O(\frac{1}{t})$

[contrast with $\frac{1}{\sqrt{t}}$ for SGD]

[4n30]

CRF objective:

SVM struct

$$\min_w \frac{\Delta \|w\|^2}{2} + \frac{1}{n} \sum_{i=1}^n H_i(w)$$

primal

$$\max_{\alpha_i \in \Delta_{|\mathcal{Y}_i|}} \frac{-\Delta \|w(\alpha)\|^2}{2} + \frac{1}{n} \sum_{i=1}^n \alpha_i^T x_i$$

dual

CRF:

$$\min_w \frac{\Delta \|w\|^2}{2} + \frac{1}{n} \sum_{i=1}^n -\log p(y^{(i)} | x^{(i)}; w)$$

↓

$$\log \left(\sum_{\tilde{y}} \exp(-w^T \psi(\tilde{y})) \right)$$

$$\max_{\alpha_i \in \Delta_{|\mathcal{Y}_i|}} \frac{-\Delta \|w(\alpha)\|^2}{2} + \frac{1}{n} \sum_{i=1}^n H_i(\alpha_i)$$

entropy

$$\triangleq - \sum_{\tilde{y}} \alpha_i(\tilde{y}) \log \alpha_i(\tilde{y})$$

KKT $\rightarrow w(\alpha) = \frac{1}{\Delta n} \sum_i \sum_{\tilde{y}} \alpha_i(\tilde{y}) \psi_i(\tilde{y})$

$$= \frac{1}{\Delta n} \sum_i \sum_{c \in \mathcal{C}} \mu_{i,c}(\tilde{y}_c) \psi_{i,c}(\tilde{y}_c)$$

(*) at optimality

$$\alpha_i^*(y) = p(y | x^{(i)}; w(\alpha^*))$$

$$= \frac{1}{\Delta n} \sum_i \sum_{c \in \mathcal{C}} \mu_{i,c}(\tilde{y}_c) \psi_{i,c}(\tilde{y}_c)$$

$\downarrow$
 from MRF

$$p(y|x;w) \propto \exp(\langle w, \psi(x,y) \rangle)$$

$$\psi_i(y) - p(y|x;w) \leq 0$$

$\downarrow$
 $\alpha_i^* \in \text{interior of } \Delta_{\psi_i}$
unlike sparse solution in structured SVM

CRF optimization

- primal objective is smooth & strongly convex [vs non-smooth for SVM struct]
- for a while, batch L-BFGS was method of choice [batch $\Rightarrow$ slow for large n]
- [Collins & al. JMLR 2008]: online exponentiated gradient (OEG)

block-coordinate method on dual; exponentiated gradient ^{stop on block}

$$\alpha_i(\tilde{y})^{(t+1)} \propto \alpha_i(\tilde{y})^{(t)} \exp(-\delta_t \nabla_{\alpha_i} f(\alpha^{(t)}))$$

dual fct. $\swarrow$

EG alg $\rightarrow$ proximal gradient step using $KL(\alpha || \alpha^*)$ as Bregman divergence

$\rightarrow$ get linear convergence rate with cheap $O(1)$ updates (like SGD)
 [vs $O(n)$ for batch method]

[can think of it as variance reduced SGD as well]

SAGA for CRF:

[Schmidt & al.]

$$w^{(t+1)} = (1 - \lambda \delta_t) w^{(t)} - \delta_t \left[\nabla f_i(w^{(t)}) - g_i^{(t)} + \frac{1}{n} \sum_j g_j^{(t)} \right]$$

... of CH ... $w = (w_1, w_2, \dots, w_n) \in \mathbb{R}^n$, $y_i \in \{-1, 1\}$

[Schmidt & al.
AISTATS 2015]

SDCA
stochastic dual
coord. ascent.

$$\alpha_{i,t}(\tilde{y})^{(t+1)} = (1 - \gamma_t) \alpha_{i,t}(\tilde{y})^{(t)} + \gamma_t \underbrace{\tilde{S}_i(\tilde{y})^{(t)}}_{p(\tilde{y} | x_i, w^{(t)})} \quad \begin{matrix} \in [0,1] \text{ (stepsize)} \\ \text{[marginal inference]} \end{matrix}$$

[note: BCFW is a
special case of SDCA
on SVM struct obj.]

SOA
for CRF
Leifrid & al.
[UAI 2016]

as a
relaxed fixed point
update for $\alpha_i^* = p(y | x_i, w^*)$

Proximal gradient method

↳ generalization of projected gradient method to other non-smooth functions

composite framework: $F(w) \triangleq f(w) + \Omega(w)$ where f is convex & L -smooth

• constrained opt. $\Omega(w) = \delta_M(w) \triangleq \begin{cases} 0 & \text{if } w \in M \\ +\infty & \text{o.w.} \end{cases}$ $\begin{matrix} \text{Q is convex but not} \\ \text{nec. smooth} \end{matrix}$
"indicator of M "

• ℓ_2 -regularization $\Omega(w) = \|w\|_2$

proximal gradient update:

$$w_{t+1} = \operatorname{argmin}_w f(w_t) + \langle \nabla f(w_t), w - w_t \rangle + \frac{1}{2} \|w - w_t\|^2 + \Omega(w)$$

$$w_{t+1} = \underset{w}{\operatorname{argmin}} \underbrace{f(w_t) + \langle \nabla f(w_t), w - w_t \rangle + \frac{1}{2\delta_t} \|w - w_t\|^2}_{\triangleq B_t(w)} + \Omega(w)$$

if $\delta_t \leq \frac{1}{L}$ then $f(w) \leq B_t(w) \forall w$
 we can rewrite $B_t(w) = \frac{1}{2\delta_t} \|w - [w_t - \delta_t \nabla f(w_t)]\|^2 + \text{const.}$ (by completing the square)

$\rightarrow$ if $\Omega(w) = \delta_M(w)$, we get the projected gradient alg.

$$w_{t+1} = \operatorname{prox}_{\delta_t}^{\Omega}(w_t - \delta_t \nabla f(w_t))$$

$\hookrightarrow$ "proximal operator" $\operatorname{prox}_{\gamma}^{\Omega}(z) \triangleq \underset{w}{\operatorname{argmin}} \{ \Omega(w) + \frac{1}{2\gamma} \|w - z\|^2 \}$

like projection, prox operator is non-expansive (i.e. 1-Lipschitz)

$$\text{i.e. } \|\operatorname{prox}_{\gamma}^{\Omega}(w) - \operatorname{prox}_{\gamma}^{\Omega}(w')\|_2 \leq \|w - w'\|_2$$

$\Rightarrow$ convergence rate of prox. gradient method
 are same as unconstrained gradient descent

$\downarrow$
 could replace
 with Bregman divergence
 to get other generalization
 (e.g. OSG)