

Lecture 20 - variance reduction

Thursday, March 30, 2023 1:29 PM

today: • variance reduction perspective
• application CRF

variance reduction idea:

X & Y are 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:

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

└ convex combo coeff. between $\mathbb{E}X$ & $\mathbb{E}Y$

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 $\mathbb{E}X = \mathbb{E}Y$ [not interesting]
 $\alpha=1$

$$\text{variance: } \text{var}(\Theta_\alpha) = \alpha^2 \left[\text{var}X + \text{var}Y - 2\text{cov}(X,Y) \right]$$

variance reduction?

for $\alpha=1$ (unbiased setting) $\Theta_\alpha = X + \underbrace{(\mathbb{E}Y - Y)}_{\text{correction}}$

SGD setting:

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

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

$$\mathbb{E}X = \frac{1}{n} \sum_{i=1}^n g_i$$

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

SAGA alg.: $\alpha=1 \Rightarrow$ unbiased

$$\alpha(X-Y) + \mathbb{E}Y$$

$\alpha(X-Y) + \mathbb{E}Y$

SAG: $x_{t+1} = x_t - \gamma \left[\frac{1}{n} [\nabla_{i_t} f(x_t) - g_{i_t}^{(t)}] + \frac{1}{n} \sum_{i=1}^n g_i^{(+)} \right]$ (biased)

SAGA: $x_{t+1} = x_t - \gamma \left[\frac{1}{n} [\nabla_{i_t} f(x_t) - g_{i_t}^{(t)}] + \frac{1}{n} \sum_{i=1}^n g_i^{(+)} \right]$ (unbiased)

SVRG: $x_{t+1} = x_t - \gamma \left[\frac{1}{n} [\nabla_{i_t} f(x_t) - \nabla_{i_t} f(x_{\text{old}})] + \frac{1}{n} \sum_{i=1}^n \nabla_{i_t} f(x_{\text{old}}) \right]$ (unbiased)

(stochastic variance reduced gradient)

$\rightarrow x_{\text{old}}$ is updated from outer loop

SVRG algorithm:

```

for k=0, ... (outer loop)
  compute  $g_{\text{ref}} \triangleq \frac{1}{n} \sum_{i=1}^n \nabla f_i(x^{(k)})$ 
  for t=0, ..., Tmax
    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
   $x^{(k+1)} = x_{T_{\text{max}}}^{(k)}$ 
end

```

Questions:

- what is T_{max} ?
- what is γ ?

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

$$T_{\text{max}} \geq \frac{1}{\mu} = k \rightarrow \text{to run alg, need to know } k$$

$\Rightarrow$ not adaptive to local strong convexity

tworack of SVRG (now called "loopless")

[Hoffmann & al. NeurIPS 2015] $T_{\text{max}} \sim \text{Geom}(-)$

[at inner loop, do a batch gradient comp. with prob $\frac{1}{n}$]

then get same convergence as SAGA

$\rightarrow$ comp. cost: size of inner loop $\mathbb{E}[T_{\text{max}}] = n$

overall cost of SURG $\sim 3(\text{SGD cost})$ for n updates

SAG / SAGA / loopless SVRG, convergence for convex fct. ($\mu=0$)

$$\text{get } \min_{t \in T} \{ \mathbb{E} f(x_t) - f^* \} = O\left(\frac{1}{T}\right)$$

[contrast this with $O\left(\frac{1}{\sqrt{T}}\right)$ for SGD]

⊛ note: interpolation regime: $\|\nabla f_i(x^*)\| = 0 \forall i$

$$\text{[vs. just. } \left\| \frac{1}{n} \sum_i \nabla f_i(x^*) \right\| = 0]$$

↳ get similar rate as SVRG / SAGA with SGD

CRF optimization

SVM struct	$\min_w \frac{\lambda \ w\ ^2}{2} + \frac{1}{n} \sum_{i=1}^n H_i(w)$ primal	$\max_{\alpha_i \in \Delta_{ \mathcal{Y} }} \frac{-\lambda \ w(\alpha)\ ^2}{2} + \frac{1}{n} \sum_{i=1}^n \langle \mathbf{1}_i, \alpha_i \rangle$ dual
CRF	$\min_w \frac{\lambda \ w\ ^2}{2} + \frac{1}{n} \sum_{i=1}^n -\log p(\tilde{y}^{(i)} x^{(i)}; w)$ $\log\left(\sum_{\tilde{y}} \exp(-w^T \psi(\tilde{y}))\right)$	$\max_{\alpha_i \in \Delta_{ \mathcal{Y} }} \frac{-\lambda \ w(\alpha)\ ^2}{2} + \frac{1}{n} \sum_{i=1}^n H_i(\alpha_i)$ $\underbrace{H_i(\alpha_i)}_{\text{entropy}} = -\sum_{\tilde{y}} \alpha_i(\tilde{y}) \log \alpha_i(\tilde{y})$

$$\begin{aligned} \text{KKT} \rightarrow w(\alpha) &= \frac{1}{\lambda n} \sum_i \sum_{\tilde{y}} \alpha_i(\tilde{y}) \psi_i(\tilde{y}) \\ &= \frac{1}{\lambda n} \sum_i \sum_{c \in \mathcal{C}} \sum_{\tilde{y}_c} \underbrace{\mu_{i,c}(\tilde{y}_c)}_{\substack{\text{from MRF} \\ \text{2}}} \psi_{i,c}(\tilde{y}_c) \end{aligned}$$

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

⊛ at optimality

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

$\alpha_i^* \in \text{interior of } \Delta_{|\mathcal{Y}|}$

unlike sparse solution in SVM struct

14h33

CRF optimization

- primal is smooth [vs. non-smooth for SVM struct]

- for a while, batch L-BFGS was method of choice [batch $\Rightarrow$ slow large n]
- [Cedric & al. JMLR 2008] = online exponentiated gradient (OEG)

block-coordinate method on dual; exponentiated gradient step on block

$$\alpha_i(\tilde{y})^{(t+1)} \propto \alpha_i(\tilde{y})^{(t)} \exp(-\gamma_t \underbrace{D_{\alpha_i} \tilde{D}(\alpha^{(t)})}_{\text{dual obj.}})$$

EG alg $\rightarrow$ proximal gradient step using $KL(\alpha / \alpha^{(t)})$ as a Bregman divergence for prox term

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

[can think of it as a variance reduced method as well?]

- SAGA for CRF [Schmitt & al. AISTATS 2015]

$$w^{(t+1)} = (1 - \gamma_t) w^{(t)} - \gamma_t \left[D_{f_i}(w^{(t)}) - g_i^{(t)} + \sum_{j \neq i} g_j^{(t)} \right]$$

- SDCA (stochastic dual coordinate ascent)
 $\rightarrow$ SOTA for CRF [Le Prie & al. UAI 2015] thanks line search

$$\alpha_i^{(t+1)}(\tilde{y}) = (1 - \gamma_t) \alpha_i^{(t)}(\tilde{y}) + \gamma_t \underbrace{\tilde{S}_i(\tilde{y})}_{\in [0,1] \rightarrow \text{stabilize}}$$

as a relaxed fixed pt. iteration

$$\alpha_i(\tilde{y}) = p(y|x_i; w(\alpha)) v_i$$

$$\underbrace{p(\tilde{y} | x_i; w(\alpha^{(t)}))}_{\cong \alpha_i(w)}$$

related to subgradient primal

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

proximal gradient method

$\rightarrow$ generalization of projected gradient method to other non-smooth fct.

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

- constrained opt. $\Omega(w) = \Delta_M(w) \triangleq \begin{cases} 0 & \text{if } w \in M \\ +\infty & \text{o.w.} \end{cases}$ Ω " " but not neces. smooth
"indicator fct." on M

- ℓ_1 -regularization $\Omega(w) = \|w\|_1$

proximal gradient update:

$$w_{t+1} = \underset{w}{\operatorname{argmin}} \underbrace{f(w_t) + \langle \nabla f(w_t), w - w_t \rangle + \frac{1}{2\delta_t} \|w - w_t\|_2^2}_{\triangleq B_t(w)} + \Omega(w)$$

• if $\delta \leq \frac{1}{L}$, then $f(w) \leq B_t(w) \quad \forall w$

• we can rewrite $B_t(w) = \frac{1}{2\delta_t} \|w - [w_t - \delta_t \nabla f(w_t)]\|_2^2 + \text{const.}$

$\Rightarrow$ if $\Omega(w) = \delta_M(w)$; we get (by completing square) projected gradient alg.

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

$\rightarrow$ "proximal operator"

$$\operatorname{prox}_{\delta}^{\Omega}(z) \triangleq \underset{w}{\operatorname{argmin}} \left\{ \Omega(w) + \frac{1}{2\delta} \|w - z\|_2^2 \right\}$$

could replace by
proximal divergence to get other
generalization (e.g. OEG)

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

$$\text{i.e. } \|\operatorname{prox}(w) - \operatorname{prox}(w')\|_2 \leq \|w - w'\|_2 \quad (\text{recall lecture 11 landscape of rates})$$

$\Rightarrow$ convergence rate for prox gradient method or $F = f + \Omega$ are same as unconstrained gradient descent on f

* to be useful, need $\operatorname{prox}_{\delta}^{-\Omega}$ to be efficiently computable

$$\operatorname{prox}_{\delta}^{\|\cdot\|_1}(z) = \underset{w}{\operatorname{argmin}} \|w\|_1 + \frac{1}{2\delta} \|w - z\|_2^2$$

$$\text{"soft-thresholding" (component wise)} \triangleq \begin{cases} \operatorname{sgn}(z_d) [|z_d| - \delta] & \text{if } |z_d| \geq \delta \\ 0 & \text{o.w.} \end{cases}$$

used e.g. for lasso: ℓ_1 -reg least-square

FISTA $\rightarrow$ accelerated prox-gradient method

$\rightarrow$ SOTA for batch lasso

$$\min_w \frac{1}{n} \sum_{i=1}^n \ell_i(w) + \Omega(w)$$

* split-learn $\rightarrow$ use (prox) SAGA for loss & ℓ_1 -reg by reg.

$$\text{prox SAGA} \quad w_{t+1} = \operatorname{prox}_{\delta_t}^{-\Omega} \left[w_t - \delta_t \left[\nabla f_t(w_t) - g_t^{(t)} + \frac{1}{n} \sum_{i=1}^n g_i^{(t)} \right] \right]$$

prox SAGA $w_{t+1} = \text{prox}_{\frac{\Omega}{\delta}} \left[w_t - \delta \left[\nabla f_t(w_t) - g_t^{(t)} + \frac{1}{n} \sum_j g_j^{(t)} \right] \right]$

could accelerate prox SAGA using "catalyst" [see next class]