

Lecture 19 - SAG

Tuesday, March 19, 2019 14:40

today: continue BCFW on SVMstruct
 • SAG / SAGA etc.

extensions of BCFW:

- non-uniform sampling e.g. $\left\{ \begin{matrix} C_i^{(i)} \\ g_i(x) \end{matrix} \right\}$
 - using away-step etc. to "get" linear convergence
- } ICML 2016

Application of BCFW on SVMstruct:

- getting s_i is one loss-augmented decoding call for example $\alpha_i^{(t+1)} = \alpha_i^{(t)} + \gamma (s_i^{(t)} - \alpha_i^{(t)})$
 - $w = A\alpha = \sum_i A_i \alpha_i$
 $\quad \quad \quad \triangleq w_i$
 $\quad \quad \quad \nearrow$
 need to store these in memory
 or store α_i (memory/computation tradeoff)
- $\rightarrow$ you update $\left\{ \begin{matrix} w_i^{(t+1)} = w_i^{(t)} + \gamma (s_i^{(t)} - w_i^{(t)}) \rightarrow A_i \\ \quad \quad \quad \hookrightarrow \frac{1}{\lambda n} N_i(\hat{y}_i^{(t)}) \\ w_j^{(t+1)} = w_j^{(t)} \quad \forall j \neq i \end{matrix} \right.$

notes for line search, also need to store $b_i^T \alpha_i \triangleq \ell_{i,\alpha}^{(t)}$

Convergence constants:

for SVMstruct, can show that $C_i^{(i)} \leq 4R^2 \Rightarrow C_i^{(i)} \leq 4R^2 \approx C_f$

for SVMstruct, can show that $C_f^{(i)} \leq \frac{4R^2}{\lambda}$

$$\Rightarrow C_f^{(i)} \leq \frac{4R^2}{\lambda n} \approx \frac{C_f}{n}$$

$$C_f \leq \frac{4R^2}{\lambda}$$

uses affine invariance
critically

ie. BCFW is "n times" faster
than batch FW for SVMstruct?

$$C_f \leq L_{\|\cdot\|} \cdot \text{diam}_{\|\cdot\|}(M)^2$$

if use ℓ_2 -norm, get a bad bound ∇ $\text{diam}(M)^2 = 2n$

Lipschitz constant in ℓ_2 -norm = largest e-value of Hessian

$$\lambda A^T A = \frac{1}{\lambda n^2} \langle \psi_i(y), \psi_j(y) \rangle_{(i,y), (j,y)}$$

say e.g. $\langle \psi_i(y), \psi_j(y) \rangle \approx 1$ for lots of outputs

$$(\mathbf{1}\mathbf{1}^T) \mathbf{1} = \mathbf{1} \cdot \text{dim}$$

$\rightarrow$ get largest e-value can scale with dim. of matrix

(here, dim is exponential
 $\Rightarrow$ useless bound)

• instead, want to use ℓ_2 -norm on $\Delta_{\pi_{S_i}}$
 i.e. ℓ_p -norm for Lipschitz constant...

$$\ell_p(v) \geq \ell_q(v) \text{ when } p \leq q$$

then get $L(\text{diam}(M))^2 \approx \frac{4B^2}{\lambda}$

on M , use ℓ_1 - ℓ_p norm (ℓ_p - ℓ_1 ?)

i.e. $\|x\| = \max_i \|x_i\|_1$

Variance reduced SGD

setup: $\min_x \underbrace{\frac{1}{n} \sum_{i=1}^n f_i(x)}_{\triangleq f(x)}$

where f is μ -strongly convex
 and L -smooth

batch gradient method:
 [Cauchy 18th century]

$$x_{t+1} = x_t - \gamma \underbrace{\frac{1}{n} \sum_{i=1}^n \nabla f_i(x_t)}_{\substack{\nabla f(x_t) \\ O(n) \text{ to compute}}}$$

linear rate
 $\downarrow$
 $\gamma = \frac{1}{L} \quad f(x_t) - f^* \leq (1-\rho)^t (f(x_0) - f^*)$
 where $\rho \approx \frac{\mu}{L} = \frac{1}{\kappa}$
 $\kappa \triangleq \frac{L}{\mu}$ "condition #"

Stochastic gradient method
 [Robbins & Monro 1951]

aka. incremental gradient method

$\ln \mathbb{E} (f(x_t) - f^*) \rightarrow$ batch gradient method

$$x_{t+1} = x_t - \gamma_t \nabla f_{i_t}(x_t)$$

where $i_t \sim \text{unif}(\{1, \dots, n\})$
 $O(1)$ to compute

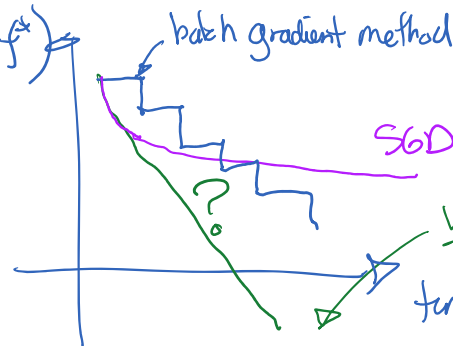
$\gamma_t = \text{const. } \gamma \rightarrow$ linear rate up to a ball of radius γ
 $\gamma_t \approx \frac{1}{\mu t} \rightarrow \tilde{O}\left(\frac{1}{\mu t}\right)$ rate
 ie. sublinear

var. variance of gradient

$\log(f(x_t) - f^*)$

$O(1)$ to compute

ie. sublinear



variance reduced SGD methods?

SAG, SAGA, SVRG, SDCA, OEG, etc...

SAG (stochastic average gradient) [Le Roux, Schmitt & Bach NIPS 2012]
(15h30) (Lagrange opt. prize in 2018)

SAG; • store past gradient for each i (g_i)
• update one at step t

SAG $\left\{ \begin{array}{l} \text{pick } i_t \sim \text{unif}; \text{ update } g_{i_t}^{(t+1)} \triangleq \nabla f_{i_t}(x_t) \\ g_j^{(t+1)} = g_j^{(t)} \quad \forall j \neq i_t \\ x_{t+1} = x_t - \gamma \frac{1}{n} \sum_{i=1}^n g_i^{(t+1)} \end{array} \right.$

$$\frac{1}{n} \left[\sum_{i=1}^n g_i^{(t)} + g_{i_t}^{(t+1)} - g_{i_t}^{(t)} \right]$$

$\rightarrow g_{i_t}^{(t+1)}$ is a stale gradient

$\frac{1}{n} \sum_i g_i \approx$ an approximation of $\nabla f(x_t)$

$O(1)$ cost per iteration
[but $O(n)$ storage cost]

big surprise: converge linearly and fast

$\rightarrow$ vs,

"increment aggregated gradient" (IAG) [Blatt et al. 2007]
 where you cycle deterministically
 through $\Sigma_1, \dots, \Sigma_n$

→ linear rate for quadratic function
 but $\delta_{\max} \approx O(\frac{1}{n})$ (tiny rate)
 "big step size"

SAG convergence rate: thm: with $\delta_t = \frac{1}{16L}$ where $L = \max_i \text{Lipschitz}(Df_i)$

$$\mathbb{E}f(x_t) - f^* \leq \left(1 - \min\left\{\frac{\mu}{16L}, \frac{1}{8n}\right\}\right)^t C_0 \quad \text{constant}$$

$$\rho_{\text{SAG}} = \min\left\{\frac{1}{16K_{\text{SAG}}}, \frac{1}{8n}\right\} \quad \text{vs} \quad \rho_{\text{grad.}} \approx \frac{1}{K_{\text{grad}}}$$

example: ^{b-reg.} log regression on RCV1

$$n = 700k \quad L = 0.25 \quad \mu = \frac{1}{n} \quad d = ? \quad (K = \frac{n}{4})$$

rate comparison:

$$\text{gradient method} \quad \left(\frac{L-\mu}{L+\mu}\right)^2 = 0.99998$$

$$\text{accelerated grad. (Nesterov)} \quad (1 - \sqrt{\frac{\mu}{L}}) = 0.99761$$

$$\text{Nesterov lower bound:} \quad \left(\frac{\sqrt{L} - \sqrt{\mu}}{\sqrt{L} + \sqrt{\mu}}\right) \approx 0.99048$$

$$-VLTVM /$$

$$\text{SAG} \quad (n \text{ iterations}) \quad (1 - \rho_{\text{SAG}})^n = 0.88250$$

practical aspects (see Schmidt Math. prog. & paper) ^{2016 big prize?}

a) storage: if $f_i(w) = h(x_i^T w)$ $\Rightarrow \nabla_w f_i(w) = \underbrace{h'(x_i^T w)}_{\text{scalar}} x_i$ input data
↓
 x_i

instead of $O(n \cdot d)$ storage $\leadsto O(n)$ storage

b) initialization? of α_i 's memory best: run SGD for one pass, then start SAG/SAQA

c) step size?

$$\frac{1}{(4)L}$$

• cheap line search heuristic
(comes from FISTA)

$$\left\{ \begin{array}{l} \text{while } f_i(w_t - \frac{1}{L_i} \nabla f_i(w_t)) \geq f_i(w_t) - \frac{1}{2L_i} \|\nabla f_i(w_t)\|^2 \\ \text{set } \tilde{L}_{i,\text{new}} = 2\tilde{L}_{i,\text{old}} \\ \text{else } \tilde{L}_{i,\text{new}} = \left(\frac{1}{2}\right)^{1/n} \tilde{L}_{i,\text{old}} \end{array} \right.$$

adaptive step size
↓

use $\frac{1}{(\max_i \tilde{L}_i)}$ as step size (see paper)

d) non-uniform sampling?

sample $i \sim \frac{L_i}{\sum_j L_j}$

e) Stopping criterion:

you can use $\frac{1}{n} \sum_j g_j^t$ as approx. of $\nabla f(x_t)$

f) sparse features?

$$x_{t+1} = x_t - \delta \left[\underbrace{\nabla f_{i_t}(x_t)}_{\text{sparse}} - g_{i_t}^t + P_{\lambda} \left[\underbrace{\frac{1}{n} \sum_j g_j^t}_{\text{dense?}} \right] \right] \quad (\text{sparse SAGA})$$

[Leblond et al.]
2017

weighted projection on support of x_t

$$S_t = \{u : (x_t)_u \neq 0\}$$