

today: • finish BCFW
• SAG

application of BCFW to SUMstruct:

• getting s_i^* is an loss-augmented decoding call for example i

$$\alpha_i^{(t+1)} = \alpha_i^{(t)} + \lambda_t (s_i^{(t)} - \alpha_i^{(t)}) \quad \text{AP}$$

• $w = Ax = \sum_i A_i \alpha_i$
 $\triangleq w_i$

(worst case)
 $\rightarrow O(nd)$

need to store
these in memory

$$\begin{cases} w_i^{(t+1)} = w_i^{(t)} + \lambda_t (w_s^{(t)} - w_i^{(t)}) \\ w_j^{(t+1)} = w_j^{(t)} \quad \forall j \neq i \end{cases}$$

$\downarrow \frac{1}{n} \sum_i \psi_i(\hat{y}_i^{(t)})$

or
store α_i 's

(memory/computation tradeoff)

note: for line search, also need to store $b_i^T \alpha_i \triangleq q_i^{(t)}$

convergence constants:

for SUMstruct, can show that $G^{(i)} \leq \frac{4R_i^2}{\lambda n^2}$

$$R_i \triangleq \max_{y \in \mathcal{Y}_i} \|\psi_i(y)\|_2$$

Hessian: $(\lambda A^T A)_{(i,y), (i,y')}$
 $= \lambda \frac{1}{\lambda n^2} \langle \psi_i(y), \psi_i(y') \rangle$

$$d_i^T H_i d_i \leq \frac{1}{\lambda n^2} \max_{y, y'} \langle \psi_i(y), \psi_i(y') \rangle \|d_i\|_2^2$$

$\leq R_i^2 \leq 2$

compare with

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

$$\Rightarrow G \approx \sum_{i=1}^n G^{(i)} \leq \frac{4R^2 n}{\lambda n^2} \approx \frac{G}{n}$$

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

importance of affine invariance:

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

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

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

recall 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. $\langle \psi_i(y), \psi_i(y) \rangle \approx 1$ for lots of output
 $(\mathbb{1}\mathbb{1}^T) \mathbb{1} = (\text{dim}) \cdot \mathbb{1}$

→ get largest e-value can scale with dim of a matrix
 ⇒ really bad here because dim is exponentially

• instead, want to use ℓ_1 -norm of Δ_{RS_1}

i.e. ℓ_1 -norm for Lipschitz constant

then get $L: (\text{diam}(\Delta_{\text{RS}_1}))^2 \approx \frac{4R^2}{\lambda}$

variance reduced SGD

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

where f is μ -strongly convex
 L -smooth

$f(x_t) - f^* \leq (1-\mu) \frac{1}{L} (f(x_t) - f^*)$
 linear rate

batch gradient method
 [Cauchy 18th century]

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

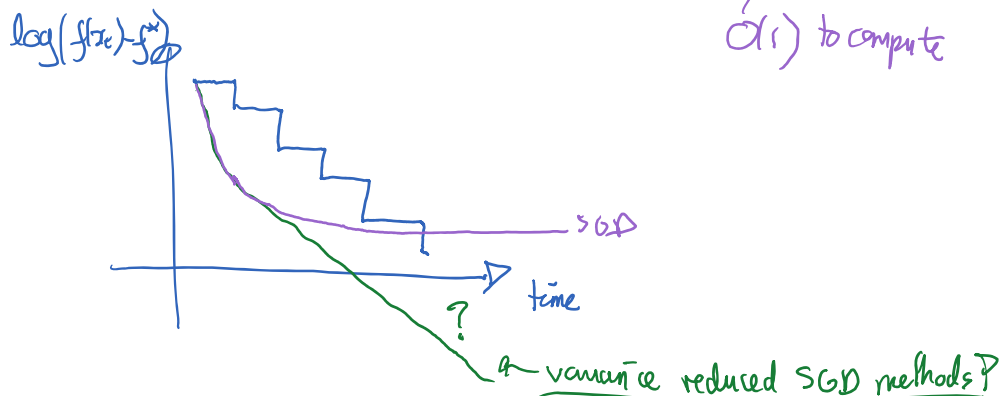
$\gamma = \frac{1}{L}$
 where $\mu \approx \frac{\mu}{L} = \frac{1}{K}$
 $K \triangleq \frac{L}{\mu}$ "condition #"

stochastic gradient method
 aka. incremental gradient method
 [Robbins & Monro 1951]

$x_{t+1} = x_t - \gamma_t \nabla f_{i_t}(x_t)$
 where $i_t \sim \text{unif}(\{1, \dots, n\})$
 $O(1)$ to compute

$\gamma_t = \text{const. } \gamma$
 ⇒ linear rate up to a ball of radius γ

$\gamma_t = \frac{1}{\mu t}$
 ⇒ $\tilde{O}(\frac{1}{t})$ rate
 i.e. sublinear



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

15h33

SAG (stochastic average gradient)

[Le Roux, Schmitt & Bach NeurIPS 2012]

(Sugranger qf. prize 2018)

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

SAG: pick i_t unif., ; update $g_{i_t}^{(t+1)} = \nabla f_{i_t}(x_t)$
 $g_j^{(t+1)} = g_j^{(t)}$ for $j \neq i_t$

$$x_{t+1} = x_t - \gamma \frac{1}{n} \sum_{i=1}^n g_i^{(t+1)} \leftarrow \text{store \underline{stale} gradients}$$

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

$\frac{1}{n} \sum g_i \approx \text{approx of } \nabla f(x_t)$

• $O(1)$ cost per iteration big surprise: converge linearly and fast

[but $O(nd)$ storage cost]
 in worst case

"increment aggregated gradient" (IAG) [Blatt' d. 2007]

where you cycle deterministically through $\{1, \dots, n\}$

$\rightarrow$ linear rate for quadratic f. but $\gamma_{\max} \approx O(1/n)$ (sting rate)
 big step-size vs

SAG convergence rate thm.

with $\gamma_t = \frac{1}{16L}$

where $L = \max_i \text{Lipschitz}(\nabla f_i)$

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

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

example: l_2 -reg. log-regression on RCV1

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

rate comparison

$$\text{gradient method } \left(\frac{L - \mu}{L + \mu} \right)^2 \approx 0.9998$$

accelerated gradient (Nesterov) $(1 - \frac{1}{\sqrt{L}}) = 0.19761$

SAG (iterations) $(1 - \rho_{SAG})^n \approx 0.8825$

practical aspects [see Schmidt & al. Math Prog. 2016 paper]

a) storage: if $f_i(w) = h(x_i^T w) \Rightarrow \nabla_w f_i(w) = h'(x_i^T w) x_i$
 instead of $O(n \cdot d)$ storage $\leadsto$ scalar $O(n)$ storage

b) initialization of g_i 's? best: run SGD for one pass then switch SAG/SA6A

c) step size? • $\frac{1}{(4)L}$ (wrong inequality)

• cheap line search heuristic (comes from FISTA)

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

d) non-uniform sampling? sample $i \sim \frac{\tilde{L}_i}{\sum_j \tilde{L}_j}$

e) stopping criterion?

you can use $\frac{1}{n} \sum_j g_j^{(t)}$ as approx $\nabla f(x_t) = \frac{1}{n} \sum_{i=1}^n \nabla f_i(x_t)$

f) sparse features?

$x_{t+1} = x_t - \gamma \left[\underbrace{\nabla f_{i_t}(x_t)}_{\text{sparse}} - g_{i_t}^{(t)} \right] \underbrace{\sum_{j=1}^n g_j^{(t)}}_{\text{dense}} \quad (\text{sparse SA6A})$

weighted projection on support of x_{i_t}

$S_{i_t} \triangleq \{u: |x_{i_t,u}| > 0\}$