

today: • variance reduced SGD
• prox method, catalyst, etc..

Variance reduction idea:

here $f(x) = \frac{1}{n} \sum_i f_i(x)$

$$\begin{aligned} \nabla f &= \frac{1}{n} \sum_i \nabla f_i \\ &\stackrel{\uparrow}{=} \mathbb{E}[\nabla f_{\xi}] = \mathbb{E}[X] \end{aligned}$$

$I \sim \text{Unif}(\{1, \dots, n\})$

∇f_i

SGD approximate $\mathbb{E}[X]$ with just X (i.e. ∇f_{ξ})

general idea:

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

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

consider estimator $\Theta_{\alpha} \stackrel{\alpha \in [0,1]}{=} \alpha(X - Y) + \mathbb{E}Y$ to approximate $\mathbb{E}X$
(here X & Y are i.i.d.)

$$\mathbb{E}\Theta_{\alpha} = \alpha \mathbb{E}X + (1-\alpha) \mathbb{E}Y \quad \leadsto \text{unbiased if } \mathbb{E}Y = \mathbb{E}X \quad \left[\begin{array}{l} \text{never happens; otherwise} \\ \text{use } \mathbb{E}Y \end{array} \right]$$

$$\text{Variance } \text{Var}(\Theta_{\alpha}) = \alpha^2 [\text{Var}(X) + \text{Var}(Y) - 2\text{Cov}(X, Y)]$$

$\alpha = 1$

→ Variance reduction aspect

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

for SGD software:

~~Stochastic gradient~~

$\mathbb{E}X$ is batch gradient; X is ∇f_i

SAG/SAGA algorithm $\forall$ is g_i [past stored gradient]
 $\mathbb{E}Y = \frac{1}{n} \sum_i g_i$

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

standard SAG $w_{t+1} = w_t - \gamma \frac{1}{n} \sum_i g_i^{t+1}$

$$\sum [g_i^t + \nabla f_i(w_t) - g_i^t]$$

SAG: $w_{t+1} = w_t - \gamma \left[\nabla f_{i_t}(w_t) - \bar{g}_i^t + \frac{1}{n} \sum_j \bar{g}_j^t \right]$ (biased)

SAGA: $w_{t+1} = w_t - \gamma \left[\nabla f_{i_t}(w_t) - \bar{g}_i^t + \frac{1}{n} \sum_j \bar{g}_j^t \right]$ (unbiased
i.e. $\mathbb{E}[w_t] = \nabla f(w_*)$)

SVRG (stochastic variance reduced gradient): $w_{t+1} = w_t - \gamma \left[\nabla f_{i_t}(w_t) - \nabla f(w_0) + \frac{1}{n} \sum_j \nabla f(w_0) \right]$ (for fixed w_0 from outer loop)

SVRG for $K=0, \dots$ (outer loop)

compute $\bar{g}_{\text{ref}} = \frac{1}{n} \sum_i \nabla f_i(w^{(K)})$

$w_0^{(K)} = w^{(K)}$

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

sample i_t

$w_{t+1}^{(K)} = w_t^{(K)} - \gamma \left[\nabla f_{i_t}(w_t^{(K)}) - \nabla f_{i_t}(w_0^{(K)}) + \bar{g}_{\text{ref}} \right]$

end

$w^{(K+1)} = w_{t_{\text{max}}}^{(K)}$

questions: • what is t_{max} ?
• what is γ ?

SAG: • need to store $\bar{g}_i^t$
(O(nd))

• no t_{max} to tune
• no weird 2 loops

SVRG: • fixed t_{max}
• weird 2 loops
• 2 gradients per iteration

• only need to store $w_0^{(K)}$

SVRG convergence result: need $\gamma \leq \frac{1}{L}$
 $t_{\text{max}} \geq \frac{1}{\mu} \frac{1}{\gamma} = K$

$\Rightarrow$ in practice, $t_{\text{max}} = \max\{n, K\}$

important remark: "adapting to local stream connectivity"

important concept: "adaptivity to ^{local} strong convexity"

SAG result $\gamma \approx \frac{1}{L} \Rightarrow$ algorithmic parameters do not depend on μ (strong convexity)

SVRG $\gamma \approx \frac{1}{L}$ but $t_{\max} = \max\{n, \frac{L}{\mu}\}$ depends on μ

for just convex fct. ($\mu=0$), get $\min_t [E f(x_t) - f^*] = O\left(\frac{1}{\sqrt{t}}\right)$

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

SAGA "simple" convergence result:

if you use $\gamma = \frac{\alpha}{L}$ for SAGA for as1 i.e. $E f(x_t) - f^* \leq (1-\rho)^t C_0$
 then rate $\rho \geq \frac{1}{5} \min\left\{\frac{1}{n}, \frac{\alpha}{K}\right\}$

two effects: • bigger step size $\Rightarrow$ bigger variance

• smaller " " $\Rightarrow$ slower gradient descent rate (determined by $\frac{1}{K}$)

if $K < n$; then as long as $\frac{\alpha}{K} \geq \frac{1}{n}$ i.e. $\boxed{1 \geq \alpha \geq \frac{K}{n}}$, we get same rate (roughly)
 $\rightarrow$ "robustness" to step size when $K \leq n$

SAGA complexity to reach ϵ error $O\left((n+K) \log \frac{1}{\epsilon}\right)$ "SGD steps"

Practical aspects of SAG/SA6A:

a) storage: if $f_i(w) = h(x_i^T w)$ $\nabla f_i(w) = \underbrace{h'(x_i^T w)}_{\text{scalar}} x_i$

i.e. instead of $O(nd)$ storage, only need $O(n)$

b) initialization? hack which is to use $\frac{1}{|S_t|} \sum_{i \in S_t} g_i^t$ where $S_t = \{i : i \text{ has been visited before } t\}$

c) step-size?

- $\frac{1}{L}$

- cheap line search heuristic $\left\{ \begin{array}{l} \text{while } f(w_t - \frac{1}{\tilde{L}} \nabla f_i(w_t)) \geq f(w_t) - \frac{1}{2\tilde{L}} \|\nabla f_i(w_t)\|^2 \\ \text{set } \tilde{L}^{new} = 2\tilde{L}^{old} \end{array} \right.$

(comes from FISTA)

$\tilde{L}$ is surrogate for L

use step-size $\frac{1}{\tilde{L}}$

d) non-uniform sampling? $\rightarrow$ sample $i \sim \frac{L_i}{\sum_j L_j}$

then get rate with $\frac{1}{n} \sum_i L_i$ instead of $\max_i L_i$

e) stopping criterion? you can use $\frac{1}{n} \sum_j g_j^t$ as approximate $\nabla f(w_t)$

f) sparse features? two tricks 1) "lagged update" $\rightarrow w_t^d = w_{t-\Delta}^d - (\Delta\gamma) g_{\text{old}}^d$ [Schmidt et al.]

sparse modification $w_{t+1} = w_t - \gamma (\nabla f_i(w_t) - g_i^t + P_{S_i}(\frac{1}{n} \sum_j g_j^t))$

[Lelebrand et al. 2017] weights projection on support of x_i

2) $f_i(w) = \frac{\lambda}{2} \|w\|^2 + h(x_i^T w)$

$$w_{t+1} = (1 - \lambda\gamma) w_t - \gamma \left[\text{sparse} \right]$$

store w as $\left\{ \begin{array}{l} \text{direction} \\ \text{scale} \end{array} \right.$

$S_i = \{d : (x_i)_d \neq 0\}$

Pointers

- variance reduction perspective on SAG / SAGA / SVRG -- in SAGA paper:
 - A. Defazio, F. Bach and S. Lacoste-Julien
SAGA: A Fast Incremental Gradient Method With Support for Non-Strongly Convex Composite Objectives
[NIPS 2014](#)
 - other pointers:
 - SVRG paper:
Rie Johnson and Tong Zhang. Accelerating stochastic gradient descent using predictive variance reduction. [NIPS 2013](#)
 - note that a variant of SVRG that is *adaptive to local strong convexity* is given in the following paper (where the end of the inner loop is decided randomly: at every inner loop iteration, with probability $1/n$, you end the inner loop):
T. Hofmann, A. Lucchi, S. Lacoste-Julien, and Brian McWilliams
Variance Reduced Stochastic Gradient Descent with Neighbors,
[NIPS 2015](#)
- the practical aspects of SAG are described in the massive journal version:
 - M. Schmidt, N. Le Roux, F. Bach.
Minimizing Finite Sums with the Stochastic Average Gradient
[Mathematical Programming](#), 162:83-162, 2017 | [arxiv](#)
 - an alternative to the complicated "lagged updates" when you have sparse features is the Sparse SAGA algorithm; see Section 2 of:
ASAGA: Asynchronous Parallel SAGA,
R. Leblond, F. Pedregosa and S. Lacoste-Julien
[AISTATS 2017](#)