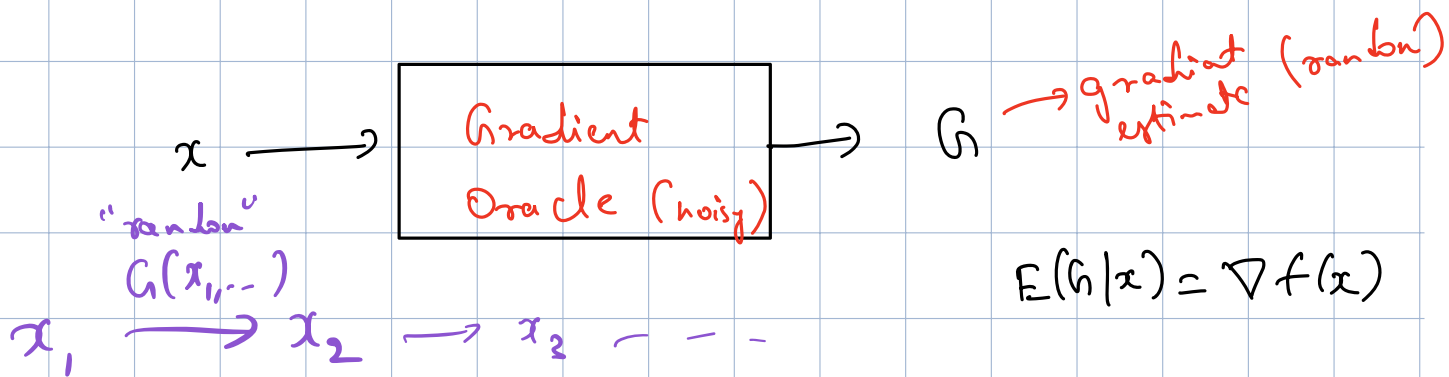


Lecture - 15 (contd)

Non-asymptotic analysis of Stochastic gradient algorithms

Aim: Solve $\min_x f(x)$
Smooth

Setting: Unbiased gradient information



SG algorithm:
$$x_{k+1} = x_k - \alpha(k) G(x_k, \xi_k) \quad (1)$$

Assumption:

step-size or
learning rate

noise

(A1)
$$E_{\xi_k} (G(x_k, \xi_k)) = \nabla f(x_k), \forall k \geq 1$$

$$E_{\xi_k} \|G(x_k, \xi_k) - \nabla f(x_k)\|^2 \leq \sigma^2$$

for some $\sigma > 0$.

(A0) f is L -smooth

$$\|\nabla f(x_1) - \nabla f(x_2)\| \leq L \|x_1 - x_2\|, \quad \forall x_1, x_2 \in \mathbb{R}^d$$

Asymptotic convergence:

$$(A2) \quad \sum_k a(k) = \infty, \quad \sum_k a(k)^2 < \infty$$

e.g. $a(k) = \frac{1}{k} \rightarrow$ standard stochastic approximation conditions.

Under (A0) - (A2), the parameter x_k governed by ① converges almost surely to the set $\{x^* \mid \nabla f(x^*) = 0\}$

Asymptotic guarantee: $x_k \rightarrow \{x^* \mid \nabla f(x^*) = 0\}$ as $k \rightarrow \infty$ w.p.1.

Tool used for proving this claim: Kushner-Clark Lemma

Eq ① ($\Rightarrow$)

$$x_{k+1} = x_k - a(k) (\nabla f(x_k) + M_{k+1})$$

$$M_{k+1} = G(x_k, \xi_k) - \nabla f(x_k)$$

("MDS")

We want to understand the non-asymptotic performance of the SG algorithm above.

Stationary point, say x^* , has $\nabla f(x^*) = 0$

ϵ -stationary point: $\bar{x}$ is an ϵ -stationary point if $E \|\nabla f(\bar{x})\| \leq \epsilon$, where the expectation is over the randomness in the SG algorithm.

[Ghadimi-Lan, 2014, SIAM J. Opt, "Stochastic first/zeroth order"]

"Randomized Stochastic gradient" (RSG)

Suppose we run ① for N iterations.

$\{x_1, x_2, x_3, \dots, x_N\} \xrightarrow{\text{(random)}} \text{Set of iterates visited by SG algorithm.}$

RSG will pick a random iterate as follows!

r is picked uniformly at random from $\{x_1, \dots, x_N\}$

$$\text{i.e., } P(x_R = x_i) = \frac{1}{N} \quad \forall i$$

$$E x_R = \frac{1}{N} \sum_{i=1}^N x_i \quad \left. \vphantom{\sum_{i=1}^N} \right\} \text{average iterate.}$$

The non-asymptotic bound for RSG is of the form

$$E \|\nabla f(x_R)\|^2 \leq \frac{\text{const}}{\sqrt{N}}$$

under suitable choice of step-size.

RL connection: $x \rightarrow$ policy parameter

Objective $f \rightarrow J_\pi(s^0)$ or $J(x)$

"Policy gradient"

↓
Value function
with states s^0
in a discounted/SSP

Proof of the non-asymptotic bound for RSG:

First $\rightarrow$ derive a bound for a general step-size

Second $\rightarrow$ specialize the bound above with a particular stepsize.

First step: Recall $x_{k+1} = x_k - a(k) g(x_k, \xi_k)$

f is L -smooth:

$$\begin{aligned} f(x_{k+1}) &\leq f(x_k) + \nabla f(x_k)^T (x_{k+1} - x_k) + \frac{L}{2} \|x_{k+1} - x_k\|^2 \\ &= f(x_k) - a(k) \nabla f(x_k)^T g(x_k, \xi_k) \\ &\quad + \frac{L}{2} a(k)^2 \|g(x_k, \xi_k)\|^2 \end{aligned} \quad \text{--- ①}$$

Taking expectation wrt $\mathcal{F}_k$ ^{$\rightarrow$ σ -field with info up to " k "}, denote this by E_k

$$\begin{aligned} E_k f(x_{k+1}) &\leq f(x_k) - a(k) \|\nabla f(x_k)\|^2 \\ &\quad + \frac{L}{2} a(k)^2 (\|\nabla f(x_k)\|^2 + \sigma^2) \end{aligned} \quad \text{--- ②}$$

wrt (A)

$$\begin{aligned} &= f(x_k) - \left(a(k) - \frac{L}{2} a(k)^2 \right) \|\nabla f(x_k)\|^2 \\ &\quad + \frac{L}{2} a(k)^2 \sigma^2 \end{aligned}$$

$$\left(a(k) - \frac{L}{2} a(k)^2 \right) \|\nabla f(x_k)\|^2 \leq f(x_k) - E_k(f(x_{k+1})) + \frac{L}{2} a(k)^2 \sigma^2$$

$$\textcircled{\times} a(k) \|\nabla f(x_k)\|^2 \leq \frac{2(f(x_k) - E_k(f(x_{k+1})))}{2 - La(k)} + \frac{La(k)^2 \sigma^2}{(2 - La(k))}$$

↳ assuming $a(k) \leq 1/L \quad \forall k$

Summing $\otimes$ over $k=1$ to N ,

$$\sum_{k=1}^N a(k) \|\nabla f(x_k)\|^2 \leq 2 \sum_{k=1}^N \frac{f(x_k) - \mathbb{E}_k f(x_{k+1})}{2 - L a(k)} + L \sigma^2 \sum_{k=1}^N \frac{a(k)^2}{(2 - L a(k))} \quad \text{--- (3)}$$

Take expectation wrt. $\{x_1, \dots, x_N\}$, i.e., $\mathbb{E}_N(\cdot)$

$$\sum_{k=1}^N a(k) \mathbb{E}_N \|\nabla f(x_k)\|^2 \leq 2 \sum_{k=1}^N \frac{\mathbb{E}_N f(x_k) - \mathbb{E}_N f(x_{k+1})}{2 - L a(k)} + L \sigma^2 \sum_{k=1}^N \frac{a(k)^2}{(2 - L a(k))}$$

"End of lecture 15"

Set $a(k) = a \quad \forall k=1 \dots N$

$$a \sum_{k=1}^N \mathbb{E}_N \|\nabla f(x_k)\|^2$$

$$\leq \frac{2}{2 - La} (f(x_1) - f(x^*)) + \frac{L \sigma^2 N a^2}{2 - La}$$

↗ global optima of f

Choose $a \leq \frac{1}{L}$

Wrt $\mathbb{E}_N f(x_{N+1}) \geq f(x^*)$

$$E_N \|\nabla f(x_R)\|^2 = \frac{1}{N} \sum_{k=1}^N E_N \|\nabla f(x_k)\|^2$$

So,

$$E_N \|\nabla f(x_R)\|^2 \leq \frac{1}{Na} \left[\frac{2}{2-La} (f(x_1) - f(x^*)) + \frac{L\sigma^2 Na^2}{2-La} \right]$$

Initial error
(bias)

$$\leq \frac{2}{Na} (f(x_1) - f(x^*))$$

Sampling error
(variance)

$$+ \frac{L\sigma^2 Na^2}{Na}$$

$$= \frac{2}{Na} (f(x_1) - f(x^*)) + L\sigma^2 a$$

Variance
due to gradient
estimation

We have,

$$E \|\nabla f(x_R)\|^2 \leq \frac{\text{const}}{Na} + \text{const } a$$

Setting $a = \min \left\{ \frac{1}{L}, \frac{1}{\sqrt{N}} \right\}$

$$x_{k+1} = x_k - a_k g(x_k, \xi_k)$$

$$a_k = \frac{1}{k} \text{ diminishing}$$

$$a_k = a \rightarrow \text{constant}$$

$$\text{Nikiforov, } a = \frac{1}{\sqrt{N}}$$

$$E \|\nabla f(x_R)\|^2 \leq \frac{2(f(x_1) - f(x^*))}{N} \max(L, \sqrt{N})$$

$$+ \frac{L\sigma^2}{\sqrt{N}}$$

$$= \frac{2L(f(x_1) - f(x^*))}{N}$$

$$+ \frac{1}{\sqrt{N}} \left(2(f(x_1) - f(x^*)) + L\sigma^2 \right)$$

$$\mathbb{E} \|\nabla f(x_p)\|^2 \leq \frac{\text{const}}{\sqrt{N}}$$

→ Bound for
Smooth f +
unbiased
gradient oracle.

($\Rightarrow$) For RSG to converge to an ϵ -stationary

point, i.e., $\mathbb{E} \|\nabla f(x_p)\|^2 \leq \epsilon$,

an order $\frac{1}{\epsilon^2}$ number of iterations

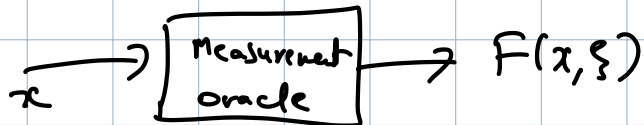
of RSG algorithm are sufficient.

Stochastic optimization with a biased gradient oracle

(regular) stochastic optimization

$$\min_{x \in \mathcal{X}} \{ f(x) = \mathbb{E}_{\xi} (F(x, \xi)) \}$$

Zeroth-order
input



"gradients" not directly available

Stochastic gradient algorithm:

$$x_{k+1} = \Pi \left(x_k - \gamma_k g_k \right)$$

$\uparrow$ projection operator $\rightarrow$ stepsize $\rightarrow$ gradient estimate

With zeroth-order input, can form gradient estimates that satisfy

$$\| \mathbb{E}_{\xi} [g(x, \xi)] - \nabla f(x) \| \leq c_1 \delta^2$$

$\nwarrow$ Bias

$$\mathbb{E}_{\xi} \| g(x, \xi) - \mathbb{E}_{\xi} (g(x, \xi)) \|^2 \leq \frac{c_2}{\delta^2}$$

$\nwarrow$ Variance

" δ " $\rightarrow$ bias-variance tradeoff.

How to form such an estimate! "Simultaneous perturbation" trick.

Here is a sample for $f \in \mathcal{L}^3$



Take two measurements: ① $f(x + \delta \Delta) + \xi^+$
② $f(x - \delta \Delta) + \xi^-$

$$\Delta = (\Delta_1, \dots, \Delta_d)$$

$$\Delta_i = \text{Rademacher} = \begin{cases} +1 & \text{w.p. } 1/2 \\ -1 & \text{w.p. } 1/2 \end{cases} \quad \text{i.i.d.}$$

$$\text{Gradient estimate } g^i = \frac{\textcircled{1} - \textcircled{2}}{2 \delta \Delta^i}$$

Use Taylor expansion

$$f(x \pm \delta \Delta) = f(x) \pm \delta \Delta^T \nabla f(x) + \frac{\delta^2}{2} \Delta^T \nabla^2 f(x) \Delta + o(\delta^3)$$

$$\frac{f(x + \delta \Delta) - f(x - \delta \Delta)}{2 \delta \Delta^i} = \nabla_i f(x) + \sum_{j \neq i} \frac{\Delta_j}{\Delta_i} \nabla_j f(x) + o(\delta^2)$$

↑
zero-mean

$$E(g^i) = \nabla_i f(x) + o(\delta^2)$$

$$\text{Can show } E(\|g - E g\|^2) \leq \frac{C \delta^2}{2}$$

Δ = Rademacher vector $\Rightarrow$ SPSA "Spall 1992"

Can we use other distributions for random perturbations.

Alternate framework: "Biased function measurements".

$$f(x) = \min_y \mathbb{E}_{\xi} (H_x(y, \xi)) \quad \text{--- (x)}$$

↑
The objective at a certain input is the solution to an optimization problem

Suppose a batch-size parameter "m" decides how well (x) is solved, i.e.,

$$F(x, m) = \min_y \mathbb{E}_{\xi} [H_x(y, \xi)] + \underbrace{\epsilon(m)}$$

error term

that has positive mean

"increasing m lowers $\epsilon(m)$ ".

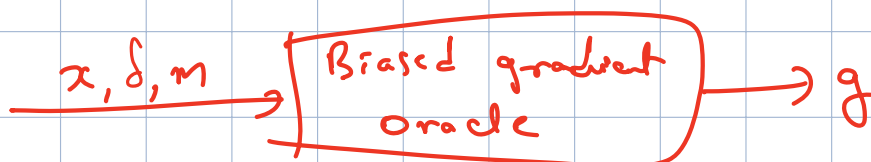
e.g. Risk-sensitive optimization ℓ : risk-measure $\ell_x = \ell(y_x)$

$\hat{\ell}_{m,x} \rightarrow$ estimate of $\ell(y_x)$ using m samples

$\mathbb{E}(\hat{\ell}_{m,x}) \neq \ell(y_x)$, but

$$\mathbb{E} |\hat{\ell}_{m,x} - \ell(y_x)|^2 \leq \frac{\text{const.}}{\sqrt{m}} \quad \forall x$$

With biased function measurements, using the simultaneous perturbation trick leads to



Bound: $\| \mathbb{E}_{\xi} [g(x, \xi, m)] - \nabla f(x) \| \leq c_1 \delta^2 + \frac{c_2}{\delta \sqrt{m}}$

Variance: $\mathbb{E}_{\xi} \| g - \mathbb{E} g \|^2 \leq \frac{c_3}{\delta^2}$

Algorithm: Randomized Stochastic gradient
[Ghadimi-Lan, SIOP 2015]

Perform $x_{k+1} = \Pi(x_k - a_k g(x_k, \xi_k, m_k))$
for N iterations

& return a random iterate R picked uniformly
at random from $\{x_1, \dots, x_N\}$

Assumption (A1) $\| \nabla f(x) - \nabla f(y) \| \leq L \|x - y\| \quad \forall x, y \in \mathcal{X}$.

Assumption (A2) $\| \nabla f(x) \|_1 \leq B < \infty, \forall x \in \mathcal{X}$

Main claim: With $a_k = \min \left\{ \frac{1}{L}, \frac{1}{N^{1/3}} \right\}$, $\delta_k = \frac{1}{N^{1/6}}$,

$$m_k = N$$

$$\forall N \geq 1, \quad \mathbb{E} \| \nabla f(x_R) \|^2 \leq \frac{2L(f(x_1) - f(x^*))}{N} + \frac{\text{const.}}{N^{1/3}}$$

See [N. Bhojan & P. C. A. "Zeroth-order"]

A biased gradient oracle variant:

Bias:

$$\| \mathbb{E}_{\xi} [g(x, \xi, m)] - \nabla f(x) \|_{\infty} \leq c_1 \delta^2 + \frac{c_2}{\delta \sqrt{m}}$$

In this oracle model, we can obtain $\left(\frac{1}{\sqrt{m}}\right)$ rate

Variance:

$$\mathbb{E}_{\xi} \|g - \mathbb{E}g\|^2 \leq c_3 \delta + \tilde{c}_3$$

Motivation! Biased function measurements in a "noise-lux" setting

Recall $g^i = \frac{f(x + \delta \Delta^i) + \xi^+ - f(x - \delta \Delta^i) - \xi^-}{2 \delta \Delta^i}$

$$\mathbb{E} \|g - \mathbb{E}g\|^2 \leq \frac{\text{Const}}{\delta^2} \quad \text{because of noise elements } \xi^+, \xi^-$$

If no noise, then $\mathbb{E} \|g - \mathbb{E}g\|^2 \leq c_3 \delta^2 + \tilde{c}_3$

Now! Use a Gaussian-smoothing estimator.

[Nesterov-Spokoiny, 2017]

$$g = \Delta \left(\frac{f(x + \delta \Delta) - f(x - \delta \Delta)}{2 \delta} \right)$$

$$\Delta = (\Delta^1, \dots, \Delta^d) \quad \Delta^i \sim N(0, 1)$$

$$x_{k+1} = x_k - \gamma_k g(x_k, \xi_k)$$

Main claim:

With $\gamma_k = \min \left\{ \frac{1}{L}, \frac{1}{\sqrt{N}} \right\}$, $\delta_k = \frac{1}{\sqrt{N}}$,

$$m_k = N$$

$$\forall N \geq 1, \quad E \|\nabla f(x_k)\|^2 \leq \frac{2(f(x_1) - f(x^*))}{N} + \frac{\text{const.}}{\sqrt{N}}$$

for an oracle without
bias-variance tradeoff

~~End of Lecture-16~~

Lecture 17 Non-asymptotic bounds for zeroth-order optimization

(0) Input: x Output: $g(x, \xi)$

$$(i) \mathbb{E}_{\xi}(g(x, \xi)) = \nabla f(x) + c_1 \delta^2 \mathbf{1}_{d \times 1}$$

$$(ii) \mathbb{E}_{\xi}[\|g(x, \xi) - \mathbb{E}_{\xi}(g(x, \xi))\|^2] \leq \frac{c_2}{\delta^2}$$

Oracle with
bias variance
tradeoff

Algorithm: Randomized Stochastic gradient

[Ghadimi-Lan, SIOT 2015]

Perform $x_{k+1} = x_k - \gamma_k g(x_k, \xi_k)$

for N iterations

4 return a random iterate R picked uniformly at random from $\{x_1, \dots, x_N\}$

Assumption (A1) $\|\nabla f(x) - \nabla f(y)\| \leq L \|x - y\| \quad \forall x, y \in \mathbb{R}^d$.

Main claim: With $\gamma_k = \min \left\{ \frac{1}{L}, \frac{1}{(d^2 N)^{1/3}} \right\}$, $\delta_k = \frac{1}{(d^5 N)^{1/6}}$,

$$\forall N \geq 1, \quad \mathbb{E} \|\nabla f(x_R)\|^2 \leq \frac{2L(f(x_1) - f(x^*))}{N} + \frac{\text{const.}}{N^{1/3}}$$

Pf: We prove a result for general γ_k & then specialize

With $P_R(k) = \text{Prob}(R=k) = \frac{\gamma_k}{\sum_{i=1}^N \gamma_i}$,

$$\mathbb{E} \|\nabla f(x_R)\|^2 \leq \frac{1}{\sum_{i=1}^N \gamma_i} \left[\frac{2L(f(x_1) - f(x^*))}{(2-L\gamma_1)} + 2B \sum_{k=1}^N c_1 \delta_k^2 \left(\frac{\gamma_k + L\gamma_k^2}{2-L\gamma_k} \right) \right]$$

$$+ L \sum_{k=1}^N \frac{\gamma_k^2}{(2-L\gamma_k)} \left(2C_1^2 \delta_k^4 + \frac{C_2}{\delta_k^2} \right)$$

Pf: Since f is L -smooth,

$$\begin{aligned} f(x_{k+1}) &\leq f(x_k) + \langle \nabla f(x_k), x_{k+1} - x_k \rangle + \frac{L}{2} \|x_{k+1} - x_k\|^2 \\ &= f(x_k) - \gamma_k \langle \nabla f(x_k), g(x_k, \xi_k) \rangle + \frac{L}{2} \gamma_k^2 \|g(x_k, \xi_k)\|^2 \end{aligned}$$

$$\mathbb{E}_k(f(x_{k+1}))$$

Conditional expectation
i.e., w.r.t info up to time k

$$\begin{aligned} &\leq f(x_k) - \gamma_k \langle \nabla f(x_k), \nabla f(x_k) + C_1 \delta_k^2 \nabla^2 f(x_k) \rangle \\ &\quad + \frac{L}{2} \gamma_k^2 \left(\mathbb{E}_k(\|g(x_k, \xi_k)\|^2) + \frac{C_2}{\delta_k^2} \right) \end{aligned}$$

w.r.t $\mathbb{E}_k \|g\|^2 = C_1^2 \delta_k^4 + \frac{C_2}{\delta_k^2}$
w.r.t $\|x_k\| \leq \sum_{i=1}^d x_i$

w.r.t $\mathbb{E}_k g = \nabla f + \delta_k^2 \nabla^2 f$

$$\leq f(x_k) - \gamma_k \|\nabla f(x_k)\|^2 + \gamma_k C_1 \delta_k^2 \|\nabla f(x_k)\|,$$

$$+ \frac{L}{2} \gamma_k^2 \left(\|\nabla f(x_k)\|^2 + 2C_1 \delta_k^2 \|\nabla f(x_k)\| + 2C_1^2 \delta_k^4 + \frac{C_2}{\delta_k^2} \right)$$

Upper bound
w.r.t $\|\nabla f(x_k)\|$

$$\leq f(x_k) - \left(\gamma_k - \frac{L}{2} \gamma_k^2 \right) \|\nabla f(x_k)\|^2 + C_1 \delta_k^2 (\gamma_k + L\gamma_k^2) B$$

$$+ \frac{L}{2} \gamma_k^2 \left(2C_1^2 \delta_k^4 + \frac{C_2}{\delta_k^2} \right)$$

Re-arrange,

$$\begin{aligned} \left(\gamma_k - \frac{L}{2} \gamma_k^2 \right) \|\nabla f(x_k)\|^2 &\leq f(x_k) - \mathbb{E}_k f(x_{k+1}) + C_1 \delta_k^2 (\gamma_k + L\gamma_k^2) B \\ &\quad + \frac{L}{2} \gamma_k^2 \left(2C_1^2 \delta_k^4 + \frac{C_2}{\delta_k^2} \right) \end{aligned}$$

Assuming $\gamma_k \leq \frac{1}{L}$

$$\gamma_k \|\nabla f(x_k)\|^2 \leq \frac{2}{(2-L\gamma_k)} \left[f(x_k) - E_k(f(x_{k+1})) + C_1 \delta_k^2 (\gamma_k + L\gamma_k^2) B \right] + \frac{L\gamma_k^2}{2-L\gamma_k} \left[2C_1^2 \delta_k^4 + \frac{C_2}{\delta_k^2} \right]$$

Summing & taking expectation,

$$\sum_{k=1}^N \gamma_k E_N \|\nabla f(x_k)\|^2 \leq 2 \sum_{k=1}^N \frac{E_N(f(x_k)) - E_N(f(x_{k+1}))}{(2-L\gamma_k)} + 2 \sum_{k=1}^N C_1 \delta_k^2 \left(\frac{\gamma_k + L\gamma_k^2}{2-L\gamma_k} \right) B + L \sum_{k=1}^N \frac{\gamma_k^2}{(2-L\gamma_k)} \left(2C_1^2 \delta_k^4 + \frac{C_2}{\delta_k^2} \right) \leq 2 \left[\frac{f(x_1)}{2-L\gamma_1} - \frac{E_N(f(x_{N+1}))}{2-L\gamma_1} \right] + 2 \sum_{k=1}^N C_1 \delta_k^2 \left(\frac{\gamma_k + L\gamma_k^2}{2-L\gamma_k} \right) B + L \sum_{k=1}^N \frac{\gamma_k^2}{(2-L\gamma_k)} \left(2C_1^2 \delta_k^4 + \frac{C_2}{\delta_k^2} \right)$$

$$\leq 2 \left[\frac{f(x_1) - f(x^*)}{(2-L\gamma_1)} \right] + 2 \sum_{k=1}^N C_1 \delta_k^2 \left(\frac{\gamma_k + L\gamma_k^2}{2-L\gamma_k} \right) B$$

$$+ L \sum_{k=1}^N \frac{\gamma_k^2}{(2-L\gamma_k)} \left(2C_1^2 \delta_k^4 + \frac{C_2}{\delta_k^2} \right)$$

$$E \|\nabla f(x_k)\|^2 \leq \frac{1}{\sum_{k=1}^N \gamma_k} \left(\right)$$

Specialize the result above for

$$\gamma_k = \gamma = \min \left\{ \frac{1}{L}, \frac{1}{(d^2 N)^{1/3}} \right\}, \quad \delta_k = \delta = \frac{1}{(d^5 N)^{1/6}}$$

$$\begin{aligned} E \|\nabla f(x_k)\|^2 &= \frac{1}{N} \sum_{k=1}^N E \|\nabla f(x_k)\|^2 \\ &\leq \frac{1}{N\gamma} \left[\frac{2(f(x_1) - f(x^*))}{2-2\gamma} + \frac{2BN C_1 \delta^2 (\gamma + L\gamma^2)}{(2-L\gamma)} \right. \\ &\quad \left. + \frac{LN\gamma^2}{2-L\gamma} \left(2C_1^2 \delta^4 + \frac{C_2}{\delta^2} \right) \right] \end{aligned}$$

using $\gamma < \frac{1}{L}$

$$\leq \frac{2(f(x_1) - f(x^*))}{N\gamma} + 4B C_1 \delta^2 L N \gamma^2 \left(2C_1^2 \delta^4 + \frac{C_2}{\delta^2} \right)$$

$$\begin{aligned} &\leq \frac{2(f(x_1) - f(x^*))}{N} \max \left(L, (d^2 N)^{2/3} \right) + \frac{4B C_1}{(d^5 N)^{1/3}} \\ &\quad + L \left(\frac{d C_1^2}{(d^5 N)^{2/3}} + \frac{2 d^{5/3}}{N^{-1/3}} \right) \frac{1}{(d^2 N)^{2/3}} \end{aligned}$$

$$= \frac{2L(f(x_1) - f(x^*))}{N} + \frac{2 d^{4/3}}{N^{1/3}} (f(x_1) - f(x^*)) + \frac{4B C_1}{d^{5/3} N^{1/3}}$$

$$+ L \left(\frac{C_1^2}{d^{7/3} N^{2/3}} + \frac{C_2 d^{5/3}}{N^{-1/3}} \right) \frac{1}{(d^2 N)^{2/3}}$$

$$= \underbrace{\frac{2L(f(x_1) - f(x^*))}{N}}_{\text{bias error}} + \underbrace{\frac{1}{N^{1/3}} \left[2d^{4/3} (f(x_1) - f(x^*)) + \frac{4B C_1}{d^{5/3}} + \frac{L C_1^2}{d^{11/3} N} + C_2 d^{1/3} \right]}_{\text{sampling error}}$$

$$E \|\nabla f(x_k)\|^2 \leq \frac{C_1 d^{4/3}}{N^{1/3}} \leq \epsilon$$

$$\text{if } N \geq \frac{d^4}{\epsilon^3}$$

Extend to biased gradient oracle without
bias-variance tradeoff, i.e., a setting like

$$\|Eg - \nabla f\| \leq C_1 \epsilon$$

$$E \|g - Eg\|^2 \leq C_2 \epsilon + C_3$$

End of Lecture-17

Analysis of stochastic gradient algorithms: Strongly convex case

7.1 SG algorithm and a useful lemma

$$x_{k+1} = x_k - \alpha_k g(x_k, \xi_k). \quad (7.1)$$

Stochastic gradient algo

Assumption 7.1. $\mathbb{E}_{\xi_k}[g(x_k, \xi_k)] = \nabla f(x_k)$ and $\mathbb{E}_{\xi_k}[\|g(x_k, \xi_k)\|_2^2] - \|\mathbb{E}_{\xi_k}[g(x_k, \xi_k)]\|_2^2 \leq \sigma^2$.

Assumption 7.2. $f(x) \geq \bar{f}$ for all x

$$\mathbb{E}_{\xi_k} g^2 \leq \sigma^2 + \|\mathbb{E} g\|^2 = \sigma^2 + \|\nabla f\|^2$$

Lemma 7.1 (Improvement in gradient step). Under Assumption 7.1, we have

$$\mathbb{E}_{\xi_k}[f(x_{k+1})] - f(x_k) \leq -\alpha_k(1 - \frac{1}{2}\alpha_k L)\|\nabla f(x_k)\|_2^2 + \frac{1}{2}\alpha_k^2 L\sigma^2. \quad (7.2)$$

Proof. Using L -smoothness of f and the update iteration (7.1), we obtain

∇f is L -Lipschitz

$$f(x_{k+1}) - f(x_k) \leq \nabla f(x_k)^T(x_{k+1} - x_k) + \frac{1}{2}L\|x_{k+1} - x_k\|_2^2 \quad (7.3)$$

$$\leq -\alpha_k \nabla f(x_k)^T g(x_k, \xi_k) + \frac{1}{2}\alpha_k^2 L\|g(x_k, \xi_k)\|_2^2. \quad (7.4)$$

Taking expectation wrt ξ_k , we obtain

$$\mathbb{E}_{\xi_k}[f(x_{k+1})] - f(x_k) \leq -\alpha_k \|\nabla f(x_k)\|_2^2 + \frac{1}{2}\alpha_k^2 L \mathbb{E}_{\xi_k}[\|g(x_k, \xi_k)\|_2^2] \quad (7.5)$$

$$\leq -\alpha_k(1 - \frac{1}{2}\alpha_k L)\|\nabla f(x_k)\|_2^2 + \frac{1}{2}\alpha_k^2 L\sigma^2. \quad (7.6)$$

The claim follows. $\square$

7.2 Strongly convex case

Theorem 7.2. *Let f be a m -strongly convex function. Then, the SG algorithm governed by (7.1) and with $\alpha_k = \alpha$ s.t. $0 < \alpha L < 1$, satisfies*

$$\mathbb{E}[f(x_n) - f(x^*)] \leq \frac{\alpha L \sigma^2}{2m} + (1 - \alpha m)^{n-1} \left(f(x_1) - f(x^*) - \frac{\alpha L \sigma^2}{2m} \right). \quad (7.7)$$

Proof. From Lemma 7.1, we have

$$\mathbb{E}_{\xi_k}[f(x_{k+1})] - f(x_k) \leq -\alpha_k \left(1 - \frac{1}{2}\alpha_k L\right) \|\nabla f(x_k)\|_2^2 + \frac{1}{2}\alpha_k^2 L \sigma^2.$$

Since $\alpha_k = \alpha$ and $0 < \alpha L < 1$, we have

$$\mathbb{E}_{\xi_k}[f(x_{k+1})] - f(x_k) \leq -\frac{1}{2}\alpha \|\nabla f(x_k)\|_2^2 + \frac{1}{2}\alpha^2 L \sigma^2. \quad (7.8)$$

Since f is m -strongly convex, the PL-condition holds, i.e.,

$$f(x) - f(x^*) \leq \frac{1}{2m} \|\nabla f(x)\|_2^2, \quad \forall x. \quad \rightarrow \text{PL-condition}$$

Using PL-condition in (7.8), we obtain

$$\mathbb{E}_{\xi_k}[f(x_{k+1})] - f(x_k) \leq -m\alpha(f(x_k) - f(x^*)) + \frac{1}{2}\alpha^2 L \sigma^2. \quad (7.9)$$

Subtracting $f(x^*)$ on both sides and re-arranging, we obtain

$$\mathbb{E}_{\xi_k}[f(x_{k+1}) - f(x^*)] \leq (1 - \alpha m)[f(x_k) - f(x^*)] + \frac{1}{2}\alpha^2 L \sigma^2. \quad (7.10)$$

Taking expectations followed by straightforward simplifications, we obtain

$$\begin{aligned} \mathbb{E}[f(x_{k+1}) - f(x^*)] - \frac{\alpha L \sigma^2}{2m} &\leq (1 - \alpha m)\mathbb{E}[f(x_k) - f(x^*)] + \frac{\alpha^2 L \sigma^2}{2} - \frac{\alpha L \sigma^2}{2m} \\ &= (1 - \alpha m) \left(\mathbb{E}[f(x_k) - f(x^*)] - \frac{\alpha L \sigma^2}{2m} \right). \end{aligned} \quad (7.11)$$

Since $\alpha m < \frac{m}{L} < 1$, a repeated application of the above inequality leads to the following bound:

$$\mathbb{E}[f(x_n) - f(x^*)] \leq \frac{\alpha L \sigma^2}{2m} + (1 - \alpha m)^{n-1} \left(f(x_1) - f(x^*) - \frac{\alpha L \sigma^2}{2m} \right). \quad (7.12)$$

The claim follows. $\square$

Remark 7.1. *Taking limits as $k \rightarrow \infty$ in (7.7), we obtain*

$$\mathbb{E}[f(x_k) - f(x^*)] \rightarrow \frac{\alpha L \sigma^2}{2m} \text{ as } k \rightarrow \infty. \quad (7.13)$$

The result above implies that a constant stepsize SG algorithm does not converge to the optima, and instead gets to within a ball around the optima.

Diminishing Stepsize:

From improvement lemma,

$$E_k(f(x_{k+1})) - f(x_k) \leq -\alpha_k \left(1 - \frac{\alpha_k L}{2}\right) \|\nabla f(x_k)\|^2 + \frac{\alpha_k^2 L \sigma^2}{2}$$

Suppose $\alpha_k L \leq 1$

$$\leq -\frac{\alpha_k}{2} \|\nabla f(x_k)\|^2 + \frac{\alpha_k^2 L \sigma^2}{2}$$

PL-condition $\rightarrow$ $\leq -\alpha_k m (f(x_k) - f(x^*)) + \frac{\alpha_k^2 L \sigma^2}{2}$

$$E_k(f(x_{k+1})) - f(x^*) \leq (1 - \alpha_k m) (f(x_k) - f(x^*)) + \frac{\alpha_k^2 L \sigma^2}{2}$$

Set $\alpha_k = \frac{\beta}{k+1}$ $\beta > \frac{1}{m}$ $\rightarrow$ strong convexity parameter

Claim: $E(f(x_n) - f(x^*)) \leq \frac{C}{n+1}$

where $C = \max \left\{ \frac{\beta^2 L \sigma^2}{2(\beta m - 1)}, 2(f(x_1) - f(x^*)) \right\}$

Proof: by induction

Base case: holds trivially

Assume claim holds for " n "

Ref:

Bottou, Nocedal, Curtis

"Opt for large scale learning"

Survey, 2018

$$E(f(x_{n+1}) - f(x^*)) \leq (1 - \alpha_n) E(f(x_n) - f(x^*)) + \frac{\alpha_n^2 L \sigma^2}{2}$$

induction hypothesis $\rightarrow$

$$\leq \left(1 - \frac{\beta_m}{n+1}\right) \frac{C}{n+1} + \frac{\beta^2 L \sigma^2}{2(n+1)^2}$$

$$= \frac{(n+1 - \beta_m) C}{(n+1)^2} + \frac{\beta^2 L \sigma^2}{2(n+1)^2}$$

$$= \frac{n}{(n+1)^2} C - \frac{(\beta_{m-1}) C}{(n+1)^2} + \frac{\beta^2 L \sigma^2}{2(n+1)^2}$$

$$\frac{\beta^2 L \sigma^2}{2(n+1)^2} - C \frac{(\beta_{m-1})}{(n+1)^2} \leq 0 \quad \text{Since}$$

$$C = \max \left(\frac{\beta^2 L \sigma^2}{2(\beta_{m-1})}, 2(f(x_n) - f(x^*)) \right)$$

So,

$$E(f(x_{n+1}) - f(x^*)) \leq \frac{nC}{(n+1)^2} \leq \frac{C}{n+2}$$

SG for convex functions: unbiased gradients case

Problem: $\min_x f(x)$

SG: $x_{k+1} = x_k - \alpha_k g(x_k, \xi_k)$

Assume: (1) f is L -smooth

(2) **Assumption 7.1.** $\mathbb{E}_{\xi_k}[g(x_k, \xi_k)] = \nabla f(x_k)$ and $\mathbb{E}_{\xi_k}[\|g(x_k, \xi_k)\|_2^2] - \|\mathbb{E}_{\xi_k}[g(x_k, \xi_k)]\|_2^2 \leq \sigma^2$.

$$\|x_{k+1} - x^*\|^2 = \|x_k - x^*\|^2 - 2\alpha_k \langle g(x_k, \xi_k), x_k - x^* \rangle + \alpha_k^2 \|g(x_k, \xi_k)\|^2$$

$$\mathbb{E} \|x_{k+1} - x^*\|^2 \leq \mathbb{E} \|x_k - x^*\|^2 - 2\alpha_k \langle \nabla f(x_k), x_k - x^* \rangle + \alpha_k^2 (\sigma^2 + \|\nabla f\|^2)$$

Convex f : $f(x^*) \geq f(x_k) + \langle \nabla f(x_k), x^* - x_k \rangle$

$\|\nabla f(x_k)\|^2 \leq L \langle \nabla f(x_k), x_k - x^* \rangle \longrightarrow$ How? n.w.

$$(2\alpha_k - L\alpha_k^2)(\mathbb{E} f(x_k) - f(x^*)) \leq (\mathbb{E} \|x_k - x^*\|^2 - \mathbb{E} \|x_{k+1} - x^*\|^2) + \alpha_k^2 \sigma^2$$

$$\alpha_k (\mathbb{E} f(x_k) - f(x^*)) \leq \frac{1}{(2 - L\alpha_k)} (\mathbb{E} \|x_k - x^*\|^2 - \mathbb{E} \|x_{k+1} - x^*\|^2 + \alpha_k^2 \sigma^2)$$

$$\sum_{k=1}^N \alpha_k (E(f(x_k)) - f(x^*)) \leq$$

$$\sum_{k=1}^N \frac{1}{(2-L\alpha_k)} \left(E \|x_k - x^*\|^2 - E \|x_{k+1} - x^*\|^2 \right) + \sum_{k=1}^N \frac{\alpha_k^2 \sigma^2}{(2-L\alpha_k)}$$

Set $L\alpha_k \leq 1$

$$\sum_{k=1}^N \alpha_k (E(f(x_k)) - f(x^*))$$

$$\leq \|x_1 - x^*\|^2 - E \|x_{N+1} - x^*\|^2 - \left(\sum_{k=2}^N \left(\frac{1}{2-L\alpha_{k-1}} - \frac{1}{2-L\alpha_k} \right) E \|x_k - x^*\|^2 \right) + \sum_{k=1}^N \alpha_k^2 \sigma^2$$

Can be ignored

$$E f(x_N) - f(x^*) \leq \frac{1}{\sum \alpha_k} \left(\|x_1 - x^*\|^2 + \sum \alpha_k^2 \sigma^2 \right)$$

$$\alpha_k \triangleq \alpha$$

$$E f(x_N) - f(x^*) \leq \frac{1}{N\alpha} \left(\|x_1 - x^*\|^2 + N\alpha^2 \sigma^2 \right)$$

$$= \frac{\|x_1 - x^*\|^2}{N\alpha} + \alpha \sigma^2$$

$$\alpha = \frac{1}{\sqrt{N}} \Rightarrow E f(x_N) - f(x^*) \leq \frac{\|x_1 - x^*\|^2}{\sqrt{N}} + \frac{\sigma^2}{\sqrt{N}}$$

Lecture - 19

SG for 200 case

↳ convex

↳ strongly convex

→ project onto a convex set C with diameter D
($\|x-y\| \leq D$ $\forall x, y \in C$)

Algo:
$$x_{k+1} = \Pi(x_k - \alpha g(x_k, \xi_k))$$

Bias gradient:
$$\|E_k g - \nabla f\| \leq C_1 \delta^2$$

$$E \|g - E g\|^2 \leq C_2 / \delta^2$$

Assume: (i) f is L -smooth & convex

(ii) $\|\nabla f(x)\| \leq B$ $\Leftarrow$ Bounded gradients

$$\begin{aligned} \|x_{k+1} - x^*\|^2 &= \|x_k - \alpha g(x_k, \xi_k) - x^*\|^2 \\ &= \|x_k - x^*\|^2 - 2\alpha g^T(x_k - x^*) + \alpha^2 \|g\|^2 \end{aligned}$$

Let $\Delta_k = g - \nabla f$

$$\begin{aligned} E_k \|x_{k+1} - x^*\|^2 &\leq \|x_k - x^*\|^2 - 2\alpha \nabla f^T(x_k - x^*) - 2\alpha \Delta_k^T(x_k - x^*) \\ &\quad + \alpha^2 \left(\|E_k g\|^2 + \frac{C_2}{\delta^2} \right) \end{aligned}$$

$$\begin{aligned} &\leq \|x_k - x^*\|^2 - 2\alpha \nabla f^T(x_k - x^*) + 2\alpha C_1 \delta^2 \|x_k - x^*\| \\ &\quad + \alpha^2 \left(\|\nabla f\|^2 + 2\sqrt{2} C_1 \delta^2 \|\nabla f\| + \alpha C_1 \delta^4 + \frac{C_2}{\delta^2} \right) \end{aligned}$$

↳ (x)

Since $-\sum_{i=1}^d y_i \leq \|y\|_1$

$$\begin{aligned} \|E_k g\|^2 &= \|\nabla f + c_1 \delta^2 \mathbf{1}_d\|^2 \quad \text{vector of ones} \\ &\leq \|\nabla f\|^2 + c_1^2 \delta^4 d + 2c_1 \delta^2 \nabla f^T \mathbf{1}_d \\ &\leq \|\nabla f\|^2 + c_1^2 \delta^4 d + 2c_1 \delta^2 \|\nabla f\| \sqrt{d} \end{aligned}$$

Letting $z_k = x_k - x^*$ and using $\|\nabla f(x)\|^2 \leq L \nabla f(x)^T (x - x^*)$

(*) $\Rightarrow$

$$\begin{aligned} E_k z_{k+1}^2 &\leq z_k^2 - 2\alpha \nabla f^T z_k + 2\alpha c_1 \delta^2 \|z_k\|_1 \\ &\quad + \alpha^2 \left[L \nabla f^T z_k + 2\sqrt{d} c_1 \delta^2 L \|z_k\| + d c_1^2 \delta^4 + \frac{c_2}{\delta^2} \right] \end{aligned}$$

Convex f : $f(x^*) \geq f(x_k) + \langle \nabla f(x_k), x^* - x_k \rangle$

Using this for

$$\begin{aligned} E_k z_{k+1}^2 &\leq z_k^2 - (2\alpha - L\alpha^2) (f(x_k) - f(x^*)) \\ &\quad + 2\sqrt{d} \|z_k\| c_1 \delta^2 (\alpha + L\alpha^2) + \alpha^2 \left[d c_1^2 \delta^4 + \frac{c_2}{\delta^2} \right] \end{aligned}$$

() using $\|y\|_1 \leq \sqrt{d} \|y\|_2$

Rearranging,

$$\begin{aligned} &\alpha (f(x_k) - f(x^*)) \\ &\leq \frac{1}{2-L\alpha} \left[z_k^2 - E_k z_{k+1}^2 + 2\sqrt{d} \|z_k\| c_1 \delta^2 (\alpha + L\alpha^2) + \alpha^2 \left(d c_1^2 \delta^4 + \frac{c_2}{\delta^2} \right) \right] \end{aligned}$$

Setting $\alpha \leq \frac{1}{L}$ & summing over $k=1 \dots N$ taking total expectations,

$$\sum_{k=1}^N \alpha \mathbb{E} (f(x_k) - f(x^*))$$

$$\leq \sum_{k=1}^N \left(\mathbb{E} z_k^2 - \mathbb{E} z_{k+1}^2 \right) + 2\sqrt{\alpha} \sum_{k=1}^N \mathbb{E} \|z_k\| C_1 \delta^2 (\alpha + L\alpha^2)$$

$$+ \sum_{k=1}^N \alpha^2 \left(d C_1 \delta^2 + \frac{C_2}{\delta^2} \right)$$

$$\leq z_1^2 + 2\sqrt{\alpha} D N C_1 \delta^2 (\alpha + L\alpha^2) + N\alpha^2 \left(d C_1 \delta^4 + \frac{C_2}{\delta^2} \right)$$

$$\mathbb{E} f(x_R) - f(x^*) \quad x_k \xrightarrow[\text{random}]{\text{unif. at}} \{x_1, \dots, x_N\}$$

$$\leq \frac{1}{N\alpha} \left[z_1^2 + 2\sqrt{\alpha} D N C_1 \delta^2 (\alpha + L\alpha^2) + \alpha^2 \left(d C_1 \delta^4 + \frac{C_2}{\delta^2} \right) \right]$$

Setting $\alpha = \min\left(\frac{1}{L}, \frac{1}{N^{2/3}}\right)$, $\delta = \frac{1}{N^{1/6}}$

$$\mathbb{E} f(x_R) - f(x^*) \leq \frac{\text{const}}{N^{1/3}}$$

The claim follows. $\square$

Remark 5.8. In contrast to the constant step size case handled previously, with a diminishing step size, we have a bound that vanishes as $m \rightarrow \infty$. However, the step size choice requires the knowledge of the strong convexity parameter μ , while the constant step size case in Theorem 5.4 did not assume such information. On a related note, it is possible to obtain a bound of $O(1/\sqrt{m})$ with a step size choice that does not require the knowledge of μ , and more importantly, with a bound that does not scale inversely with μ . Such a bound may be preferable for ill-conditioned problems, where μ is very small. The reader is referred to (Nemirovski *et al.*, 2009) for the details.

5.3.2 SG with biased gradient information

As before, we consider the update iteration in (5.21). Unlike the previous section, where we assumed unbiased gradient estimates (i.e., the condition A5.1 holds), here the estimate $\hat{\nabla} f(\theta_k)$ is a biased approximation to the gradient of the objective function f at θ_k .

As in the asymptotic analysis in Section 4.1.2, the biased gradient estimate $\hat{\nabla} f(\theta_k)$ can be decomposed as follows:

$$\begin{aligned}\hat{\nabla} f(\theta_k) &= \nabla f(\theta_k) + \beta_k + \eta_k, \text{ where} \\ \beta_k &= E \left[\hat{\nabla} f(\theta_k) \mid \mathcal{F}_k \right] - \nabla f(\theta_k), \\ \eta_k &= \hat{\nabla} f(\theta_k) - E \left[\hat{\nabla} f(\theta_k) \mid \mathcal{F}_k \right],\end{aligned}\tag{5.29}$$

where $\mathcal{F}_k$ is a σ -field generated by $\{\theta_i, i \leq k\}$. In the above, β_k is the bias in the gradient estimate and η_k , $n \geq 0$, is a martingale difference sequence.

Using a simultaneous perturbation-based gradient estimate implies $\beta_k = O(\delta_k^2)$, where δ_k is the perturbation parameter used in forming the estimate (see Chapter 3 for several examples). While the bias goes down as δ_k^2 , the variance of the gradient estimate scales inversely with δ_k^2 . This has been formalized earlier in assumptions A5.2–A5.3.

We now present a non-asymptotic bound in expectation for the SG algorithm (5.21) with inputs from a biased gradient oracle that satisfies the aforementioned assumptions.

$$\theta_{k+1} = \theta_k - a_k \hat{\nabla} f(\theta_k)$$

136

Non-asymptotic analysis of stochastic gradient algorithms

f is μ -strongly convex

Proposition 5.1. Suppose the objective function f is L -smooth (see Definition 3.1), and assumptions A5.2–A5.3 hold. Then, we have

$$\begin{aligned} \mathbb{E} \|\theta_{m+1} - \theta^*\|^2 &\leq \underbrace{2 \exp(-2\mu\Gamma(m)) \|\theta_0 - \theta^*\|^2}_{\text{initial error}} \\ &\quad + \underbrace{2 \sum_{k=1}^n a_k^2 \exp(-2\mu(\Gamma(m) - \Gamma_k)) c_1^2 \delta_k^4}_{\text{bias error}} \\ &\quad + \underbrace{\sum_{k=1}^n a_k^2 \exp(-2\mu(\Gamma(m) - \Gamma_k)) c_2 \delta_k^{-2}}_{\text{sampling error}}, \end{aligned} \quad (5.30)$$

where $\Gamma(k) := \sum_{i=1}^k a_i$.

$$\int_0^1 f''(\theta^* + \lambda(\theta_m - \theta^*)) (\theta_m - \theta^*) d\lambda = f'(\theta_m)$$

Proof. Let $z_m = \theta_m - \theta^*$ denote the error at time instant n of the algorithm (5.21). Using $\nabla f(\theta^*) = 0$, we have

$$\left(\int_0^1 \nabla^2 f(\theta^* + \lambda(\theta_m - \theta^*)) d\lambda \right) z_m = \nabla f(\theta_m).$$

Using the fact above, we arrive at a recursion for z_m from (5.29). Letting

$J_m := \int_0^1 \nabla^2 f(\theta^* + \lambda(\theta_m - \theta^*)) d\lambda$, we have

$$\begin{aligned} z_{m+1} &= (I - a(m)J_m)z_m - a(m)(\beta_m + \eta_m) \\ &= \Pi_m z_0 - \sum_{k=1}^n a(k) \Pi_m \Pi_k^{-1} (\beta_k + \eta_k), \end{aligned}$$

where $\Pi_m := \prod_{k=1}^n (I - a(k)J_k)$.

By the conditional Jensen's inequality, we obtain

$$(\mathbb{E}_m \|z_{m+1}\|)^2 \leq \mathbb{E}_m(\langle z_m, z_m \rangle)$$

$$\begin{aligned} &= \mathbb{E}_m \left(\|\Pi_m z_0\|^2 + \left\| \sum_{k=1}^n a(k) \Pi_m \Pi_k^{-1} \beta_k \right\|^2 + \left\| \sum_{k=1}^n a(k) \Pi_m \Pi_k^{-1} \eta_k \right\|^2 \right. \\ &\quad \left. - 2 \left\langle \Pi_m z_0, \sum_{k=1}^n a(k) \Pi_m \Pi_k^{-1} \beta_k \right\rangle - 2 \left\langle \Pi_m z_0, \sum_{k=1}^n a(k) \Pi_m \Pi_k^{-1} \eta_k \right\rangle \right) \end{aligned}$$

shorthand for $\prod_{j=2}^n (I - a(j)J_j)$

$$\left\langle \sum_{k=1}^n a(k) \Pi_m \Pi_k^{-1} \beta_k, \sum_{k=1}^n a(k) \Pi_m \Pi_k^{-1} \eta_k \right\rangle \quad (5.31)$$

$$\begin{aligned} &\leq 2 \|\Pi_m z_0\|^2 + 2 \sum_{k=1}^n a(k)^2 \|\Pi_m \Pi_k^{-1}\|^2 c_1^2 \delta_k^4 \\ &\quad + \sum_{k=1}^n a(k)^2 \|\Pi_m \Pi_k^{-1}\|^2 \mathbb{E} \|\eta_k\|^2. \end{aligned} \quad (5.32)$$

For the last inequality, we have used the following facts: (i) η_k is a martingale difference implying the last two cross terms in (5.31) are zero; (ii) $\beta_k \leq c_1 \delta_k^2$ from A5.3; and (iii) Cauchy-Schwarz inequality for the first cross term in (5.31).

Now, we bound each of the square terms in (5.32) separately. Since the objective is strongly convex, we have that $\|I - a(m)J_m\| \leq \exp(-\mu a(m))$. Hence,

$$\begin{aligned} \|\Pi_m \Pi_k^{-1}\|_2 &= \left\| \prod_{j=k+1}^n (I - a_j J_j) \right\|_2 \\ &\leq \prod_{j=k+1}^n \|(1 - a_j \mu)I - a_j(J_j - \mu I)\|_2 \\ &\leq \prod_{j=k+1}^n \|(1 - a_j \mu)I\|_2 \leq \prod_{j=k+1}^n (1 - a_j \mu) \\ &\leq \exp(-\mu(\Gamma(m) - \Gamma(k))). \end{aligned} \quad (5.33)$$

$\Gamma(m) = \sum_{j=1}^m a(j)$

From A5.3, we can infer that the second moment of the martingale difference is bounded above by c_2/δ_k^2 . The main claim now follows by plugging the bound on η_m and (5.33) into (5.32). $\square$

By specializing the result in the proposition above, we derive a non-asymptotic bound of the order $O(1/\sqrt{m})$.

Theorem 5.6. (Biased gradients and strongly convex objec-

tive) Let $a(k) = c/k$ and $\delta_k = \delta_0/k^\delta$. Then,

$$\begin{aligned} \mathbb{E} \|\theta_m - \theta^*\| &\leq \frac{\sqrt{2} \|\theta_0 - \theta^*\|}{m^{\mu c}} + \frac{\sqrt{2} c c_1 \delta_0^2}{\sqrt{2\mu c - 4\delta - 1}} m^{-\frac{1}{2} - 2\delta} \\ &\quad + \frac{\sqrt{c_2 c}}{\delta_0 \sqrt{2\mu c + 2\delta - 1}} m^{\delta - \frac{1}{2}}. \end{aligned}$$

Remark 5.9. Choosing $\delta = 0$, one can obtain a bound of the order $O(m^{-1/2})$ for simultaneous perturbation schemes that lead to biased gradient estimates, and this bound matches the corresponding bound with unbiased gradient information up to constant factors. Contrast this with the difference in rates between biased and unbiased gradient information for the non-convex and convex cases in the previous sections.

Remark 5.10. Using L -smoothness of f and $\nabla f(\theta^*) = 0$, we have

$$\mathbb{E}[f(\theta_m)] - f(\theta^*) \leq \frac{L}{2} \mathbb{E} \|\theta_m - \theta^*\|^2 = O\left(\frac{1}{m}\right).$$

Proof. Bounding a sum by an integral, we obtain

$$\exp(-\mu \Gamma(m)) \leq \exp(-\mu c \ln m) = m^{-\mu c}.$$

Plugging $a(k) = c/k$ and $\delta_k = \delta_0/k^\delta$ into the bias error term in (5.30), we obtain

$$\begin{aligned} \sum_{k=1}^m a(k)^2 \exp(-2\mu(\Gamma(m) - \Gamma_k)) c_1^2 \delta_k^4 &\leq \sum_{k=1}^m \frac{c^2}{k^2} m^{-2\mu c} k^{2\mu c} c_1^2 \frac{\delta_0^4}{m^{4\delta}} \\ &\leq c^2 m^{-2\mu c} c_1^2 \delta_0^4 \sum_{k=1}^m k^{2\mu c - 4\delta - 2} \\ &\leq \frac{c^2 c_1^2 \delta_0^4}{(2\mu c - 4\delta - 1)} m^{-1 - 4\delta}. \end{aligned}$$

Along similar lines, the sampling error term in (5.30) can be upper-bounded as follows:

$$\sum_{k=1}^m a(k)^2 \exp(-2\mu(\Gamma(m) - \Gamma_k)) \frac{c_2}{\delta_k^2} \leq \frac{c^2 c_2}{\delta_0^2 (2\mu c - 4\delta - 1)} m^{-1 + 2\delta}.$$

□