

Lecture 21 - catalyst

Thursday, April 1, 2021 13:28

today: catalyst $\rightarrow$ accelerate
non-convex opt.
submodular opt.

catalyst algorithm [Lin Maillard & Harchaoui NeurIPS 2015]

"meta-algorithm" : outer loop which uses a linearly convergent alg in inner loop to get overall acceleration $\rightarrow$

main idea: use the accelerated proximal point algorithm

with approximation inner loop of prox operator

proximal point alg. : is proximal gradient with $f=0$

$$\boxed{w_{t+1} = \text{prox}_{\gamma}^{\Omega}(w_t)} \quad (\text{to solve } \min_w \Omega(w))$$

catalyst alg. (for μ -strongly $F(w)$)

let $q \triangleq \frac{\mu}{1 + \frac{1}{\gamma}}$

(γ is algorithmic parameter)

repeat:

$$w_{t+1} \approx \arg \min_w \underbrace{F(w) + \frac{1}{2\gamma} \|w - z_t\|_2^2}_{\triangleq G_t(w)}$$

s.t. $(G_t(w_{t+1}) - \min_w G_t(w)) \leq \epsilon_t$ to be specified

$$z_t = \text{prox}_{\gamma}^{F(\cdot)}(z_t)$$

use linear loop optimization [e.g. SAGA or AFW] with warmstart

$$z_{t+1} = w_{t+1} + \beta_{t+1} (w_{t+1} - w_t)$$

"extrapolation" / "momentum"

[accelerated Nesterov trick piece]

β_{t+1} is found using fancy equations so that everything works

• solve for α_{t+1} in eq: $\alpha_{t+1}^2 = (1 - \alpha_{t+1}) \alpha_t^2 + q \alpha_{t+1}$

$$\beta_{t+1} \triangleq \frac{\alpha_t (1 - \alpha_t)}{\alpha_t^2 + \alpha_{t+1}} \quad (\text{pick } \alpha_{t+1} \in]0, 1[)$$

catalyst trick: use γ & ϵ_t
 s.t. overall # of inner loop calls
 gives an overall acceleration

with clever analysis of warm starting

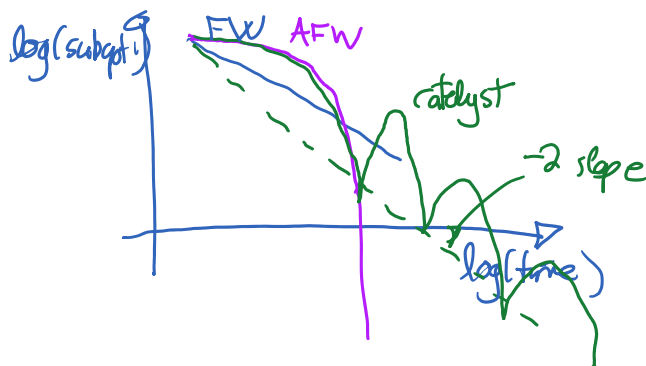
acceleration results:

if inner loop alg. has convergence exp $(-\frac{\tilde{\mu}}{L}t)$ strong convexity G_L/w $\tilde{\mu} \geq \mu + \frac{1}{8}$
 then with correct constants for γ & ϵ_t

(μ -strongly convex F) linear rate $\mu = \frac{1}{K}$ becomes with catalyst $\frac{1}{\sqrt{K}}$ for catalyst
 (F convex case) $O(\frac{1}{t})$ on F becomes $O(\frac{1}{\sqrt{t}})$

result: we can get (theory) accelerated SAGA
 " SVRG
 " AFW
 etc.

issue: catalyst is not adaptive to local strong convexity
and finicky for choice of γ, ϵ_t, μ etc...



(4h15)

non-convex optimization

recall: FW with line search on f non-convex

$$\min_{s \leq t} g(w_s) \leq O\left(\frac{1}{\sqrt{t}}\right)$$

FW gap

convex: $\mathbb{E} f(w_t) - f^* \leq \epsilon$

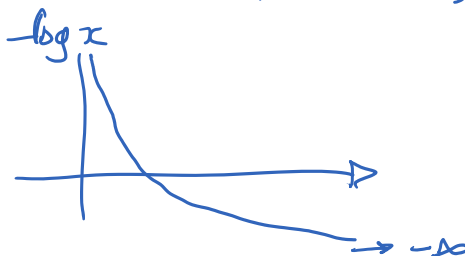
$$GD \rightarrow \frac{1}{\epsilon} \quad \text{Nesterov} \frac{1}{\sqrt{\epsilon}}$$

$$\text{non-convex: } \mathbb{E} \|\nabla f(w_t)\|_2^2 \leq \epsilon$$

blew: if f is μ -strongly convex

$$\Rightarrow f(w_t) - f^* \leq \frac{1}{2\mu} \|\nabla f(w_t)\|_2^2$$

note: $\|\nabla f(w_t)\|$ small $\nRightarrow f(w_t) - f^*$ is small
when f is not strongly convex



* can get a $O(\frac{1}{\epsilon})$ rate for gradient descent:

$$f(w) \leq f(w_t) + \langle \nabla f(w_t), w - w_t \rangle + \frac{L}{2} \|w - w_t\|^2 \quad \forall w$$

$$w_{t+1} = w_t - \frac{1}{L} \nabla f(w_t)$$

(L -smoothness of f
but no need of
convexity)

$$\Rightarrow f(w_{t+1}) \leq f(w_t) - \frac{1}{2L} \|\nabla f(w_t)\|^2$$

if f^* is finite

$$\Rightarrow f(w_t) - f^* \leq f(w_0) - f^* - \frac{1}{2L} \sum_{t=0}^{t-1} \|\nabla f(w_t)\|^2$$

$$\Rightarrow \sum_t \|\nabla f(w_t)\|^2 \leq 2L (f(w_0) - f^*)$$

$$\Rightarrow t \cdot \min_{s \leq t} \|\nabla f(w_s)\|^2 \leq 2L (f(w_0) - f^*)$$

$$\text{ie. } \boxed{\min_{s \leq t} \|\nabla f(w_s)\|^2 \leq \frac{2L}{t} (f(w_0) - f^*)}$$

Faster nonconvex optimization via VR

(Reddi, Hefny, Sra, Póczos, Smola, 2016; Reddi et al., 2016)

Algorithm	Nonconvex (Lipschitz smooth)
SGD	$O\left(\frac{1}{\epsilon^2}\right)$
GD	$O\left(\frac{n}{\epsilon}\right)$
SVRG	$O\left(n + \frac{n^{2/3}}{\epsilon}\right)$
SAGA	$O\left(n + \frac{n^{2/3}}{\epsilon}\right)$
MSVRG	$O\left(\min\left(\frac{1}{\epsilon^2}, \frac{n^{2/3}}{\epsilon}\right)\right)$

Remarks

$$\mathbb{E}[\|\nabla g(\theta_t)\|^2] \leq \epsilon$$

New results for convex case too; additional nonconvex results
For related results, see also (Allen-Zhu, Hazan, 2016)

20

Linear rates for nonconvex problems

$$\min_{\theta \in \mathbb{R}^d} g(\theta) = \frac{1}{n} \sum_{i=1}^n f_i(\theta)$$

The Polyak-Łojasiewicz (PL) class of functions

$$g(\theta) - g(\theta^*) \leq \frac{1}{2\mu} \|\nabla g(\theta)\|^2$$

(Polyak, 1963); (Łojasiewicz, 1963)

Linear rates for nonconvex problems

$$g(\theta) - g(\theta^*) \leq \frac{1}{2\mu} \|\nabla g(\theta)\|^2 \quad \Bigg| \quad \mathbb{E}[g(\theta_t) - g^*] \leq \epsilon \quad \text{😎}$$

Algorithm	Nonconvex	Nonconvex-PL
SGD	$O\left(\frac{1}{\epsilon^2}\right)$	$O\left(\frac{1}{\epsilon^2}\right)$
GD	$O\left(\frac{n}{\epsilon}\right)$	$O\left(\frac{n}{2\mu} \log \frac{1}{\epsilon}\right)$
SVRG	$O\left(n + \frac{n^{2/3}}{\epsilon}\right)$	$O\left(\left(n + \frac{n^{2/3}}{2\mu}\right) \log \frac{1}{\epsilon}\right)$
SAGA	$O\left(n + \frac{n^{2/3}}{\epsilon}\right)$	$O\left(\left(n + \frac{n^{2/3}}{2\mu}\right) \log \frac{1}{\epsilon}\right)$
MSVRG	$O\left(\min\left(\frac{1}{\epsilon^2}, \frac{n^{2/3}}{\epsilon}\right)\right)$	—

Variant of **nc-SVRG** attains this fast convergence!

(Reddi, Hefny, Sra, Póczos, Smola, 2016; Reddi et al., 2016) 22

Submodular optimization

Submodularity is an analog of convexity for tractability of set functions

(combinatorial opt.)

$$F: 2^V \rightarrow \mathbb{R}$$

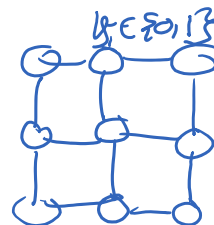
$V = \{1, \dots, d\}$ is "ground set"

$2^V = \{V \rightarrow \{0, 1\}\} = \text{set of all subsets of } V$

concrete example:

Ising model

$$E(y) = \sum_i \Theta_i y_i - \sum_i \sum_{j \in \text{neighbor of } i} G_{ij} y_i y_j$$



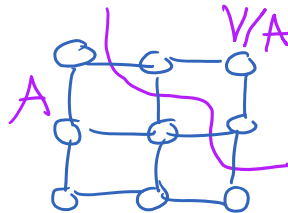
$$A_y = \{i : y_i = 1\}$$

$$F(A_y) =$$

when $G_{ij} > 0$, $\Rightarrow E(y)$ is "attractive potential" submodular

MRF here is "associative Markov network" AMN

can minimize $E(y)$ (or $F(A_y)$) by using "graph cut algorithm"



equivalence with min cost network flow problem

F is submodular $\Leftrightarrow F(A) + F(B) \geq F(A \cap B) + F(A \cup B) \quad \forall A, B \subseteq V$

$\Leftrightarrow$ function $A \mapsto F(A \cup \{k\}) - F(A)$ is non-increasing for all k

$$\text{i.e. } F(A \cup \{k\}) - F(A) \leq F(B \cup \{k\}) - F(B) \quad \text{if } B \subseteq A$$

"diminishing return property"

$\Rightarrow$ intuitively, that greedy alg. are not "too bad" for maximization

* $F(A) \stackrel{\Delta}{=} g(|A|)$
↑ cardinality

if g is concave $\Rightarrow$ then F is submodular

* Link with convexity $\rightarrow$ Loewy extension (cls. set.) $A \subseteq V$

$\leftarrow$...

* embed sets as corners of hypercube in dimension d $v(A) = \mathbb{1}_A \in \{0,1\}^d$

Boas extension f extends $F(\cdot)$ from corners to entire hypercube using convex interpolation

(piecewise linear fct. on $[0,1]^d$)

$$f(w) = F(A) \text{ when } w = v(A)$$

$$\text{let's say } w = \sum_i \alpha_i v_i \Rightarrow f(w) = \sum_i \alpha_i F(A_i)$$

$\downarrow$
 $v(A_i)$

F is submodular $\Leftrightarrow$ Boas extension f is convex

* can write $f(w) = \max_{S \in B(F)} \langle s, w \rangle$

"base polytope"

← this can be computed efficiently using greedy alg

(LMO over $B(F)$ is efficient)

$$\min_{A \subseteq V} F(A) = \min_{w \in [0,1]^d} \left(\max_{S \in B(F)} \langle s, w \rangle \right)$$

$\underbrace{\hspace{1cm}}_{f(w)}$

→ use projected subgradient method

$$\partial f(w) = \arg \max_{S \in B(F)} \langle s, w \rangle$$

* with l_2 -regularization, use duality to get a smooth obj

$$\min_{S \in B(F)} \frac{1}{2} \|s\|^2$$

→ use "min-norm pt." alg
variant of FCFW alg. * SOTA for submodular opt.