

today: • finish calibration
• start convex optimization

optimization normalization for calibration function

Setup let $s(x, \tilde{y})$ be of the form $F\theta(x)$ $\theta(x) \in \mathbb{R}^r$

$$s(x, \cdot) = F\theta(x) \quad \theta_j(\cdot) \in \mathcal{H} \subseteq \text{RKHS}$$

$\underbrace{\theta_j(\cdot)}_{\text{optimization variables}}$

$\in \mathbb{R}^k$

$$\mathcal{J}(\theta) = \mathbb{E}_{(x,y) \sim p} \mathcal{J}(x, y; \theta)$$

projection $(x^{(t)}, y^{(t)}) \stackrel{\text{iid}}{\sim} p$

run projected gradient SGD on $\mathcal{J}(\theta)$ i.e. $\theta^{(k+1)} = \mathcal{P}_{\mathcal{H}}(\theta^{(k)} - \gamma \nabla_{\theta} \mathcal{J}(x^{(t)}, y^{(t)}; \theta^{(k)}))$

ball of radius D around $\mathbf{0}$

$$\nabla_{\theta} \mathcal{J}(x^{(t)}, y^{(t)}; \theta) = F^T \nabla_s \mathcal{J}(x^{(t)}, y^{(t)}; s) \Phi(x^{(t)})^T$$

$r \times \infty$

feature map of $\mathcal{H}$
 $k(x^{(t)}, \cdot)$

$$\theta^{(k)} = F^T \sum_t \alpha_t \Phi(x^{(t)})^T$$

$r \times \infty$ $\theta^{(k)}(x) \rightarrow \langle \Phi(x^{(t)}), \Phi(x) \rangle$

Convergence result:
(Thm. 5)

if $\|\theta^*\|_{\mathcal{H}} \leq D$ & $\mathcal{J}$ is convex and differentiable

and if $\mathbb{E}_{(x,y) \sim p} \|\nabla_{\theta} \mathcal{J}(x, y; \theta)\|_{\mathcal{H}}^2 \leq M^2$

then averaged projected SGD with step-size $\gamma = \frac{2D}{M\sqrt{n}}$ gives

$$\mathbb{E}[\mathcal{J}(\bar{\theta}^{[n]})] - \mathcal{J}(\theta^*) \leq \frac{2DM}{\sqrt{n}}$$

Thm. 6 learning complexity

let G^* minimize $L(\theta)$ with $\|\theta^*\|_{HS} \leq D$

choosing $n \geq \frac{4D^2M^2}{H(\epsilon)^2}$ implies $\mathbb{E}[L(\theta^{(n)})] \leq L(\theta^*) + \epsilon$

define a meaningful side

in the paper: we compute $D, M, H(\epsilon)$ for specific losses l and the quadratic l to get sample complexity

* moral here:

* some losses are harder than others (worst case sample complexity)
[0-1 loss is difficult in general "no bonus of a loss"]

* have linked computation to statistical performance in consistency framework
→ convex surrogate losses

* (non-parametric analysis) could handle dependence on x using RKHS

Caveats:

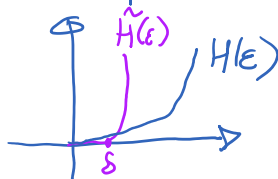
* distribution free result (i.e. worst case over all distributions)

→ still need more theory?

(e.g. role of $p(y|x)$ or other surrogates?)

models $\|\theta^*\|_{HS} \leq D$ constraint
could be big for "bad p "
[→ no free lunch]

* follow-up: inconsistent surrogate loss with computational/statistical advantages [Neurips 2018]



14/20

part II: convex optimization

motivation: $\min_w \frac{\lambda \|w\|^2}{2} + \frac{1}{n} \sum_{i=1}^n \ell(x^{(i)}, y^{(i)}; w)$

convex surrogate loss

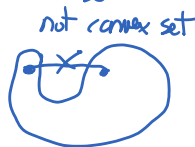
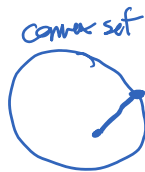
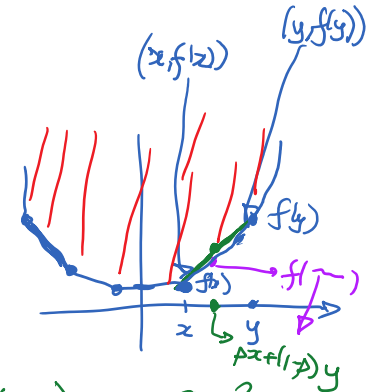
convex analysis recap:

$f: \mathbb{R}^d \rightarrow \mathbb{R}$

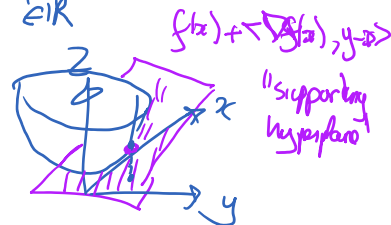
$$f: \mathbb{R}^d \rightarrow \mathbb{R}$$

f is convex $\Leftrightarrow$

$$f(\underbrace{px + (1-p)y}_{\text{convex combination between } x \text{ \& } y}) \leq pf(x) + (1-p)f(y) \quad \forall x, y \in \text{dom}(f)$$



$$\text{epigraph}(f) \triangleq \{(x, y) : y \geq f(x)\}$$



$\Rightarrow$ If f is differentiable at x and convex

$$\Rightarrow f(y) \geq f(x) + \langle \nabla f(x), y - x \rangle \quad \forall y$$

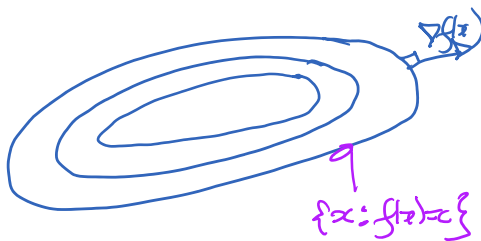
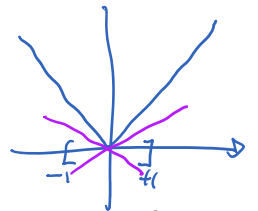
(suppose f is convex)

subdifferential

subgradient v of f at x : $v \in \partial f(x)$

$$\Leftrightarrow \forall y \in \text{dom}(f), \quad f(y) \geq f(x) + \langle v, y - x \rangle$$

$$|x| = \max\{x, -x\}$$



$$\partial f(x) = \text{conv}\{\nabla f(x_1), \nabla f(x_2)\}$$

when $f(x) = \max_i f_i(x)$ where f_i is differentiable

$$\partial f(x) = \text{conv}\{\nabla f_{i^*}(x) : i^* \in \arg\max_i f_i(x)\}$$

(Darskin's theorem)

Darskin's theorem : https://en.wikipedia.org/wiki/Darskin%27s_theorem

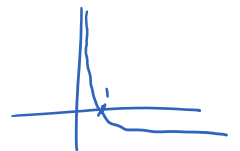
Clarke's subdifferential $\rightarrow$ nice gen. to non-convex

$$f: \mathbb{R}^d \rightarrow \mathbb{R} \cup \{\pm\infty\}$$

$$\text{dom}(f) \triangleq \{x \in \mathbb{R}^d : f(x) < \infty\}$$

e.g. $-\log x$ is convex

$$\text{dom}(f) = \{x \in \mathbb{R} : x > 0\}$$



$$\Rightarrow \min_x f(x) = \min_{x \in \text{dom}(f)} f(x)$$

Some standard assumptions:

$$f \text{ is } \mu\text{-strongly convex} \Leftrightarrow f(y) \geq f(x) + \underbrace{\langle \nabla f(x), y-x \rangle}_{\substack{\langle v, y-x \rangle \\ \text{for any } v \in \partial f(x)}} + \frac{\mu}{2} \|y-x\|^2$$

μ -strongly convex $\Rightarrow$ strong convexity constant

$\forall x, y \in \text{dom}(f)$

$$f \text{ is } \mu\text{-strongly convex} \Leftrightarrow f - \frac{\mu}{2} \|\cdot\|_2^2 \text{ is convex}$$

$$f \text{ is } L\text{-smooth} \text{ i.e. } f \text{ is } L\text{-Lipschitz cts. gradient } \forall x \text{ (with respect to norm } \|\cdot\|)$$

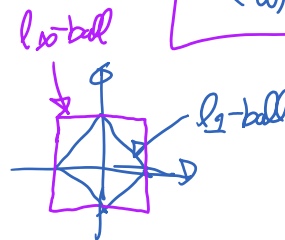
$$\Leftrightarrow \|\nabla f(x) - \nabla f(y)\|_* \leq L \|x-y\| \quad \forall x, y$$

$$(\|\cdot\|_p)_* = \|\cdot\|_q$$

$$\text{where } \frac{1}{p} + \frac{1}{q} = 1$$

$$p=2 \Rightarrow q=2$$

$$p=1 \Rightarrow q=\infty$$



"dual norm"

$$\|w\|_* \triangleq \sup_{\|v\| \leq 1} \langle w, v \rangle$$

Generalized C-S.

$$\langle w, v \rangle \leq \|w\|_* \|v\|$$

Fundamental descent lemma:

when ∇f is L -Lipschitz (lemma holds even if f is not convex)

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

$$* \quad f(\underbrace{x - \gamma \nabla f(x)}_{y_x}) \leq f(x) - \underbrace{\gamma \langle \nabla f(x), \nabla f(x) \rangle}_{\|\nabla f(x)\|_2^2} + \frac{L}{2} \gamma^2 \|\nabla f(x)\|_2^2$$

$$= f(x) - \underbrace{[\gamma(1 - \frac{\gamma L}{2})]}_{\geq 0} \|\nabla f(x)\|_2^2$$

$$\geq 0 \Leftrightarrow 0 < \gamma \leq \frac{2}{L}$$

$$\gamma \Leftrightarrow \left| 0 < \gamma < \frac{1}{L} \right|$$

→ minimize RHS with respect to γ

gives $\gamma^* = \frac{1}{L}$

$$f(y_{\gamma^*}) \leq f(x) - \frac{1}{2L} \|\nabla f(x)\|_2^2$$

proof intuition of descent lemma:

think of 2nd order Taylor expansion

integral form of remainder

$$f(y) = f(x) + \langle \nabla f(x), y-x \rangle + \frac{1}{2} \int_{\gamma=0}^1 \langle y-x, \underbrace{H(x+\gamma(y-x))}_{\text{Hessian of } f} \rangle y-x \, d\gamma$$

f L -smooth
 $\frac{1}{2}$ twice diff.

$\Rightarrow$ $\sup$ ϵ -value of H in absolute value $\leq L$

$$v^T H v \leq \lambda_{\max}(H) \|v\|_2^2$$

$$\leq L \|y-x\|^2$$