

today: • finish calibration
• start convex optimization

optimization normalization for calibration function $\Theta \sim \mathbb{R}^r \times \mathbb{D}$
 setup let $s(x, y)$ be of the form $F\Theta(x)$ $\Theta(x) \in \mathbb{R}^r$

$$s(x, \cdot) = F\Theta(x)$$

$$\Theta(\cdot) \in \mathcal{H} \leftarrow \text{RKHS}$$

optimization variables

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

projection $(x^{(t)}, y^{(t)}) \sim \rho$

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

ball of radius D around 0

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

$\mathbb{R}^r \times \mathbb{D}$

$\in \mathbb{R}^k$

$$\Theta^{(t)} = F^T \sum_t \alpha_t(\cdot) \Phi(x^{(t)})^T$$

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

$$\Theta^{(t)}(x) \rightarrow \langle \Phi(x^{(t)}), \Phi(x) \rangle$$

$$= F^T \sum_t \alpha_t(\cdot) \underbrace{\langle \Phi(x^{(t)}), \Phi(x) \rangle}_{k(x^{(t)}, x)}$$

$$\|\Theta\|_{\mathcal{H}}^2 = \sum_{j=1}^r \|\Theta_j\|_{\mathcal{RKHS}}^2$$

convergence result:
 (thm. 5)

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

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

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

$$\Theta^{[n]} = \frac{1}{n} \sum_{t=1}^n \Theta^{(t)}$$

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

thm. 6 learning complexity

let Θ^* minimize $L(\Theta)$ with $\|\Theta^*\|_{\mathcal{H}} \leq D$

choosing $n \geq \frac{4D^2M^2}{\epsilon^2}$
 $\rightarrow \frac{1}{\epsilon^2}$

implies $\mathbb{E}[L(\Theta^{[n]})] \leq L(\Theta^*) + \epsilon$

choosing $n \geq \frac{4D^2 M^2}{\tilde{H}(\epsilon)^2}$

implies $\mathbb{E}[L(\hat{\theta}^{[n]})] \leq L(\theta^*) + \epsilon$

offers a meaningful scale

in the paper: we compute $D^2 M^2$ & $\tilde{H}(\epsilon)$ for specific losses l and the quadratic $\int l^2$ to get sample complexity

⊕ moral here:

* some losses l are harder than others (worst case sample complexity)

[0-1 loss is difficult in general]
"too harsh of a loss"

* have linked computation to statistical performance in consistency framework

↳ convex surrogate loss

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

reveals:

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

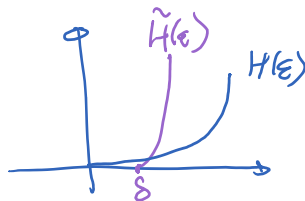
→ still need more theory?

(e.g. rule $p(y|x)$)
or other surrogates

model $\| \theta \|_{\text{HS}} \leq D$ constraint

could be big for "bad"
[→ no free lunch]

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



15/18

part II: convex optimization

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

convex surrogate loss

convex analysis recap:

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

f is convex $\Leftrightarrow$

$\Leftrightarrow$ epigraph(f) is convex

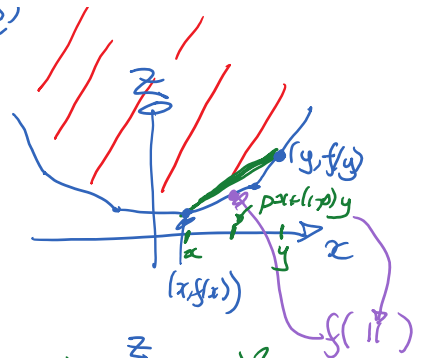
/// $\neq$ ///

f is convex $\Leftrightarrow$

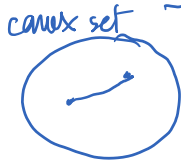
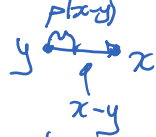
$$f(\underbrace{px + (1-p)y}_{\text{convex combination between } x \text{ \& } y}) \leq pf(x) + (1-p)f(y) \quad \forall x, y$$

convex combination
between x & y

$$\& p \in [0, 1]$$



$$y + p(x - y)$$

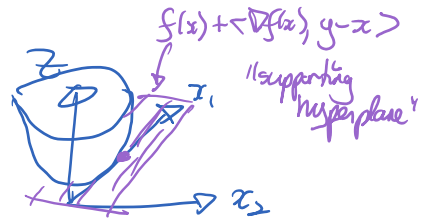


$$\text{epigraph}(f) \triangleq \{(x, y) : y \geq f(x)\}$$

$$x \in \mathbb{R}^d, y \in \mathbb{R}$$

* f is differentiable at x and convex

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

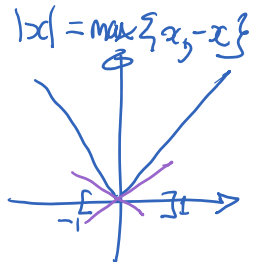


(suppose f is convex)

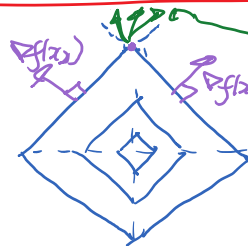
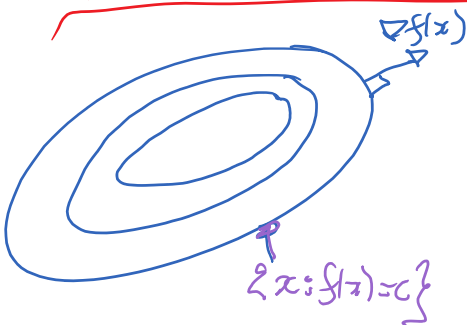
subdifferentiable

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

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



$$\partial |x|_{x=0} = [-1, 1]$$



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

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

$$\partial f(x) = \text{conv}\{\nabla f_i(x) : i \in \arg\max_j f_j(x)\}$$

(Danskin's thm.)

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

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

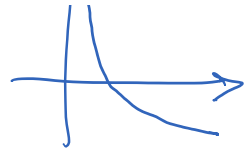
$$f : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{\pm\infty\}$$

"extended reals"

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



minimizing f over $\mathbb{R}^n$ is equivalent to



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

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

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 convexity constant

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

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

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

"dual norm"

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

$\Leftrightarrow$

generalized C.S.

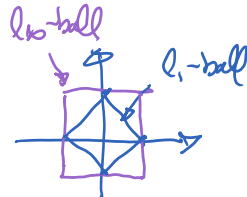
$$\langle w, v \rangle \leq \|w\|_* \|v\|$$

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

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

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

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



Fundamental descent lemma:

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

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

apply lemma

with $y = y_8$

$$* \quad f(\underbrace{x - \gamma \nabla f(x)}_{y_8}) \leq f(x) - \gamma \underbrace{\langle \nabla f(x), \nabla f(x) \rangle}_{\|\nabla f(x)\|^2} + \frac{L}{2} \gamma^2 \|\nabla f(x)\|^2$$

$$= f(x) - \left[\gamma \left(1 - \frac{\gamma L}{2} \right) \right] \|\nabla f(x)\|^2$$

$$> 0 \Leftrightarrow \boxed{0 < \gamma < \frac{2}{L}}$$

→ minimize BHS with respect to γ

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

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

proof intuition of descent Lemma:

think of 2nd order Taylor expansion

integral form of the remainder
↓

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

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

f is L -smooth & twice diff. $\Rightarrow$ top e-value of $H \leq L$
in absolute value

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

↑
in absolute value