

today: • logistic regression
• numerical optimization

logistic regression:

setup: binary classification $\mathcal{Y} = \{0, 1\}$ $X \in \mathbb{R}^d$

generative model motivation:

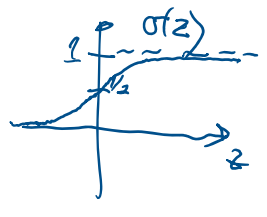
suppose only assumption is there exists pdf's (densities) in $\mathbb{R}^d$
for each class conditionals:

$$p(x|Y=1) \quad \& \quad p(x|Y=0)$$

$$\begin{aligned} p(Y=1|X=x) &= \frac{p(Y=1, X=x)}{p(Y=1, X=x) + p(Y=0, X=x)} p(X=x) \\ &= \frac{1}{1 + \frac{p(Y=0, X=x)}{p(Y=1, X=x)}} = \frac{1}{1 + \exp(-f(x))} \end{aligned}$$

where $f(x) \triangleq \log \underbrace{\frac{p(X=x|Y=1)}{p(X=x|Y=0)}}_{\text{class-conditional ratio}} + \log \underbrace{\frac{p(Y=1)}{p(Y=0)}}_{\text{prior odds ratio}}$
"log odds"

in general $p(Y=1|X=x) = \sigma(f(x))$ where $\sigma(z) \triangleq \frac{1}{1 + \exp(-z)}$ "sigmoid function"



some properties of $\sigma(z)$: $\sigma(-z) = 1 - \sigma(z)$ $[\sigma(z) + \sigma(-z) = 1]$
 $\frac{d}{dz} \sigma(z) = \sigma(z) \sigma(-z) = \sigma(z) (1 - \sigma(z))$

⊛ to motivate linear logistic regression

consider class conditional in the exponential family scalar fct. $\rightarrow$ log partition fct.

$$p(x|n) \triangleq \underbrace{h(x)}_{\text{"canonical parameter"}} \exp(\underbrace{n^T T(x)}_{\text{"sufficient statistics"}} - \underbrace{A(n)}_{\text{normalizer}})$$

these specify the "flat" exponential family

Gaussian:

$$p(x) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{x^2}{2}\right)$$

Gaussian:

$$\log p(x|\mu, \sigma^2) = -\frac{1}{2} \log 2\pi\sigma^2 - \underbrace{\frac{(x-\mu)^2}{2\sigma^2}}_{\substack{\frac{1}{2\sigma^2} \exp(-\frac{(x-\mu)^2}{2\sigma^2}) \\ \text{variance}}} - \left(\frac{x^2}{2\sigma^2} - \frac{x\mu}{\sigma^2} + \frac{\mu^2}{2\sigma^2} \right)$$

$$\text{let } T(x) = \begin{bmatrix} -x^2/2 \\ x \end{bmatrix}$$

$$n(\mu, \sigma^2) = \begin{bmatrix} 1/\sigma^2 \\ \mu/\sigma^2 \end{bmatrix}$$

$$A(n) = \frac{1}{2} \log(2\pi\sigma^2) + \frac{\mu^2}{2\sigma^2}$$

$$p(x|Y=1) = p(x|m_1)$$

$$p(x|Y=0) = p(x|m_0)$$

$$\text{log odds } f(x) = \log \frac{\overbrace{p(x|m_1)}^{p(x|Y=1)}}{\underbrace{p(x|m_0)}_{p(x|Y=0)}} + \log \frac{\overbrace{p(Y=1)}^{\pi}}{\underbrace{p(Y=0)}_{1-\pi}}$$

$$= (m_1 - m_0)^T T(x) + A(m_0) - A(m_1) + \log\left(\frac{\pi}{1-\pi}\right)$$

$$\triangleq w^T \phi(x)$$

$$\text{where } w = \begin{pmatrix} m_1 - m_0 \\ A(m_0) - A(m_1) + \log\left(\frac{\pi}{1-\pi}\right) \end{pmatrix} \quad \phi(x) = \begin{pmatrix} T(x) \\ 1 \end{pmatrix}$$

get logistic regression model?

$$p_w(Y=1|X=x) = \sigma(w^T \phi(x))$$

"feature map"

decision boundary

$$\mathbb{I} \{ w^T \phi(x) > 0 \}$$

exercise to reader: try argument above with $p(x|y) = N(x|\mu_y, \Sigma_y)$ (hwk 2)

• if $\Sigma_0 = \Sigma_1$, then $\phi(x) = \begin{pmatrix} x \\ 1 \end{pmatrix} \rightarrow$ "linear boundary"

• otherwise, $\phi(x) = \begin{pmatrix} -xx^T \\ x \\ 1 \end{pmatrix}$

optional

quadratic decision boundary

linear logistic regression model

linear logistic regression model

$$p(y=1|x, w) = \sigma(w^T x) \quad Y = \{0, 1\} \quad \text{decision rule: } \{ \{ w^T x > 0 \} \}$$

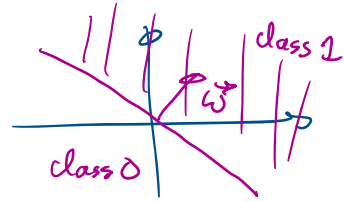
$$p(y=0|x, w) = 1 - \sigma(w^T x) = \sigma(-w^T x)$$

$$[if Y = \{-1, 1\} \text{ "Rademacher R.V."}]$$

$$\text{encode } p(y|x) = \sigma(y w^T x)$$

$$Y|X=x \text{ is Bernoulli}(\sigma(w^T x))$$

$$p(y|x) = \sigma(w^T x)^y (1 - \sigma(w^T x))^{1-y}$$



given $(x_i, y_i)_{i=1}^n$ iid

max-conditional log-likelihood to estimate w_{MCL}

$$l(w) = \sum_{i=1}^n \log p(y_i | x_i, w) = \sum_{i=1}^n y_i \log \sigma(w^T x_i) + (1 - y_i) \log (1 - \sigma(w^T x_i))$$

$$\nabla_w \sigma(w^T x) = x [\sigma(w^T x) (1 - \sigma(w^T x))] \quad \text{let } v_i \triangleq w^T x_i$$

$$\nabla_w l(w) = \sum_{i=1}^n x_i \left[\frac{\sigma(v_i) (1 - \sigma(v_i)) y_i}{\sigma(v_i) (1 - \sigma(v_i))} - \frac{(1 - y_i)}{\sigma(v_i) (1 - \sigma(v_i))} \right]$$

$$= \sum_{i=1}^n x_i \left[y_i [\sigma(v_i) + \sigma(v_i)] - \sigma(v_i) \right]$$

$$\boxed{\nabla l(w) = \sum_{i=1}^n x_i [y_i - \sigma(w^T x_i)]}$$

need to use numerical methods

solve for $\nabla l(w) = 0 \Rightarrow$ need to solve transcendental equation

$$\text{because } \frac{1}{1 + \exp(-w^T x_i)} = 0$$

↑ solve for this

contrast to linear regression

$$\nabla l(w) = \sum_{i=1}^n x_i [y_i - w^T x_i]$$

↑ in linear

numerical optimization

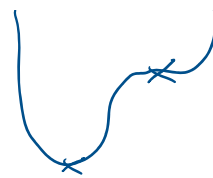
want to minimize $f(w)$ (unconstrained) [suppose f is differentiable]
s.t. $w \in \mathbb{R}^d$

1) gradient descent (1st order method)

→ start...

1) gradient descent (1st order method)

start w_0
 iterate: $w_{t+1} = w_t - \alpha_t \nabla f(w_t)$
 stopping criterion: if $\|\nabla f(w_t)\|^2 < \delta$
 then stop



e.g. $\delta = 10^{-6}$

note: if f is μ -strongly convex $\Rightarrow f(w_t) - \min_w f(w) \leq \frac{1}{2\mu} \|\nabla f(w_t)\|^2$
 $(\Leftrightarrow f(w) - \frac{\mu}{2} \|w\|^2$ is convex in w)
 if $\|\nabla f(w_t)\|^2 \leq \delta \Rightarrow \text{subopt.}(w_t) \leq \frac{\delta}{2\mu}$

step-size rules

a) constant step-size: $\alpha_t = \frac{1}{L}$ Lipschitz continuity constant for ∇f
 i.e. $\|\nabla f(w) - \nabla f(w')\| \leq L \|w - w'\|_2$

b) decreasing step-size rule
 $\alpha_t = \frac{c}{t}$ constant
 usually want $\sum_t \alpha_t = +\infty$
 $\sum_t \alpha_t^2 < \infty$

[this is more common for stochastic optimization]
 $\forall w, w' \in \mathbb{R}^d$

e.g. $f(w) = \mathbb{E}_{\xi} g(w, \xi)$

$w_{t+1} = w_t - \alpha_t \nabla_w g(w, \xi_t)$

e.g. $g(w, \xi) = \ell(y_i, h_w(x_i))$ for ERM
 $\xi = (x_i, y_i)$

c) choose α_t by "line search": $\min_{\gamma \in \mathbb{R}} f(w_t + \gamma d_t)$
 costly in general
 direction of update (e.g. $-\nabla f(w_t)$)

$\hookrightarrow$ instead do approximate search
 e.g. Armijo line search

[see Boyd's book]

15h38

Newton's method (2nd order method)

motivation: minimizing a quadratic approximation

Δ Hessian
 $[H(w_t)]_{ij} = \frac{\partial^2 f(w_t)}{\partial w_i \partial w_j}$

Taylor expansion: $f(w) \approx f(w_t) + \nabla f(w_t)^T (w - w_t) + \frac{1}{2} (w - w_t)^T H(w_t) (w - w_t)$

Taylor expansion
at w_t

$$f(w) = f(w_t) + \nabla f(w_t)^T (w - w_t) + \frac{1}{2} (w - w_t)^T H(w_t) (w - w_t) + O(\|w - w_t\|^3)$$

$$\hat{= Q_t(w)}$$

+ $O(\|w - w_t\|^3)$
Taylor's remainder

$$= Q_t(w) + O(\|w - w_t\|^3)$$

quadratic model approximation

$w_{t+1} \rightarrow$ obtain by minimizing $Q_t(w)$

$$\nabla_w Q_t(w) \stackrel{\text{want}}{=} 0$$

$$\nabla f(w_t) + H(w_t) (w - w_t) \stackrel{\text{want}}{=} 0$$

$$\Rightarrow w - w_t = H^{-1}(w_t) \nabla f(w_t)$$

$$w_{t+1} = w_t - H^{-1}(w_t) \nabla f(w_t)$$

Newton's update

$$d_t = H^{-1} \nabla f$$

$\downarrow$ inverse Hessian
 $\rightarrow O(d^3)$ time
and $O(d^2)$ space

$$H d_t = \nabla f$$

Note for hwk 2: solve for $H_t d_t = \nabla f(w_t)$
 $\uparrow$ for d_t

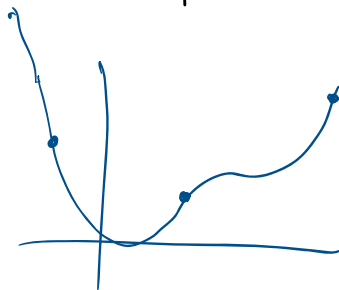
$$\min_d \|H_t d - \nabla f(w_t)\|^2$$

use `numpy.linalg.lstsq` to solve this

$\rightarrow$ much more numerically stable

Clamped Newton: you add a step-size to stabilize Newton's method

$$w_{t+1} = w_t - \underbrace{\alpha_t}_{\text{step-size}} H^{-1}(w_t) \nabla f(w_t)$$



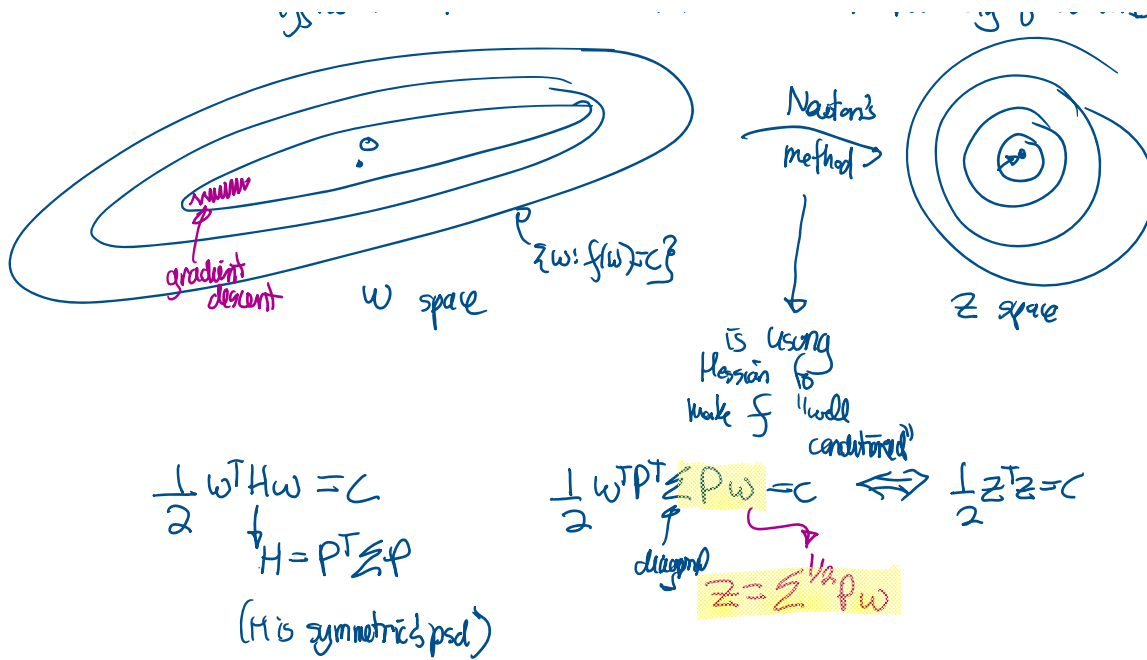
why Newton's method?

- much faster convergence in # iterations vs gradient descent
- affine invariant $\Rightarrow$ method is invariant to rescaling of variables



Next...





exercise to the reader:

$$z_{t+1} = z_t - \gamma \nabla f(z_t)$$

$$\iff w_{t+1} = w_t - \gamma H^{-1} \nabla f(w_t)$$

Big data logistic regression

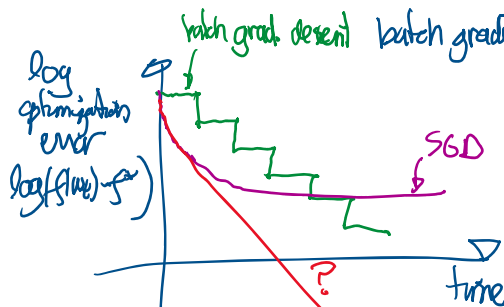
- big $d \Rightarrow$ cannot do $O(d^2)$ or $O(d^3)$ operations $\Rightarrow$ first order methods
- if n is large, you cannot use batch methods $\nabla f(w) = \frac{1}{n} \sum_{i=1}^n \nabla f_i(w)$ (full) batch gradient gradient of loss on some data pt $O(nd)$

instead you can use "incremental gradient method"

e.g. stochastic gradient descent (SGD) : $w_{t+1} = w_t - \gamma_t \nabla f_{i_t}(w_t)$

where i_t is picked unif. at random

SGD $\rightarrow$ cheap updates, but slower convergence per iteration, but faster convergence



→ SAG: stochastic average gradient

G.D. $w_{t+1} = w_t - \gamma_t \frac{1}{n} \sum_{i=1}^n \nabla f_i(w_t)$

SAG [2012] $w_{t+1} = w_t - \gamma_t \frac{1}{n} \sum_{i=1}^n \nabla f_i(w_t) + \text{memory}$

yes variance reduced methods!

SAGA [2012] $w_{t+1} = w_t - \gamma_t \frac{1}{n} \sum_{i=1}^n V_i$ memory
 where $V_i = \nabla f_i(w_t)$
 at each t , update $V_{it} = \nabla f_i(w_t)$

SAGA [2014] $w_{t+1} = w_t - \gamma_t \left(\nabla f_{it}(w_t) + \frac{1}{n} \sum_{j=1}^n V_j - V_{it} \right)$ $\frac{1}{n} \sum_{j=1}^n V_j$ $\frac{1}{n} \sum_{j=1}^n V_j - V_{it}$ $\nabla f_{it}(w_t)$
 (default method for large scale logistic regression in Scikit Learn)
 variance reduction correction

SVRG

Skipped part: Newton's method on logistic regression = IRLS:

Newton's method for logistic regression: IRLS

recall for $l(w)$: $\nabla l(w) = \sum_{i=1}^n x_i [y_i - \sigma(w^T x_i)]$

$H(l(w)) = -\sum_{i=1}^n x_i x_i^T \sigma(w^T x_i) \sigma(w^T x_i)$

$V^T H V = -\sum_{i=1}^n \underbrace{(V^T x_i)(x_i^T V)}_{(x_i^T V)^2 \geq 0} \underbrace{\sigma(-) \sigma(-)}_{\neq 0}$ $V^T H V \leq 0 \forall V \in \mathbb{R}^D$
 i.e. HJO
 i.e. concave fct.

notation: recall

$X = \begin{pmatrix} -x_1^T \\ \vdots \end{pmatrix}$
 $n \times d$

let $\mu_i \triangleq \sigma(w^T x_i) \in]0,1[$

$\nabla l(w) = \sum_i x_i [y_i - \mu_i] = X^T (y - \mu)$

Hessian $= -\sum_i x_i x_i^T \mu_i (1 - \mu_i) = -X^T D(w) X$
 $D_{ii} \triangleq \mu_i (1 - \mu_i)$
 diagonal matrix

$\mu_t \triangleq \begin{pmatrix} \vdots \\ \sigma(w_t^T x_i) \\ \vdots \end{pmatrix}$

Newton is maximizing instead of minimizing

Newton's update: $w_{t+1} = w_t - (-X^T D_t X)^{-1} X^T (y - \mu_t)$

$= (X^T D_t X)^{-1} [X^T D_t X w_t + X^T (y - \mu_t)]$

$w_{t+1} = (X^T D_t X)^{-1} [X^T D_t z_t]$ where $z_t \triangleq X w_t + D_t^{-1} (y - \mu_t)$

Newton is a rank-1 update to the Hessian

this is a solution to a "weighted least square problem"

$$\min_w \| D^{1/2} (\bar{z}_0 - Xw) \|_2^2$$

$\nwarrow$ weights $\searrow$ new target

$$\sum_i \frac{(z_i - w^T x_i)^2}{d_{ii}}$$

compare with
Gaussian noise
model for least
square

$$\sum_i \frac{(y_i - w^T x_i)^2}{2\sigma_i^2}$$

Newton's method for logistic regression

= iterated reweighted least squares (IRLS)