

Lecture 8 - logistic regression & optimization

Monday, September 30, 2024 12:33 PM

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 (densities) in $\mathbb{R}^d$
for each class conditionals

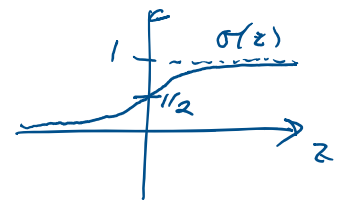
$$p(x|Y=1) \text{ \& } p(x|Y=0)$$

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

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

in general

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



some properties of $\sigma(z)$: $\sigma(-z) = 1 - \sigma(z)$ $[\sigma(z) + \sigma(-z) = 1]$

$$\frac{d\sigma(z)}{dz} = \sigma(z)\sigma(-z) = \sigma(z)(1 - \sigma(z))$$

⊛ to motivate linear logistic regression

consider class conditionals in the exponential family

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

scalar fct. $\rightarrow$ log partition fct.
 $\rightarrow$ normalise

"canonical parameter"

"sufficient statistics"

these specify the "flat" exponential family

Gaussian:

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

$\downarrow$
 $\frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$

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

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

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

$$p(x|Y=1) = p(x|\eta_1)$$

$$p(x|Y=0) = p(x|\eta_0)$$

$$\text{log odds } f(x) = \log \frac{\overbrace{p(x|\eta_1)}^{\triangle \pi}}{\underbrace{p(x|\eta_0)}_{p(x|Y=0)}} + \log \frac{\overbrace{p(Y=1)}^{\triangle \pi}}{\underbrace{p(Y=0)}_{1-\pi}}$$

$$= (\eta_1 - \eta_0)^T T(x) + A(\eta_0) - A(\eta_1) + \log\left(\frac{\pi}{1-\pi}\right) \\ \triangleq w^T \phi(x)$$

$$\text{where } w = \begin{pmatrix} \eta_1 - \eta_0 \\ A(\eta_0) - A(\eta_1) + \log\left(\frac{\pi}{1-\pi}\right) \end{pmatrix} \quad \phi(x) = \begin{pmatrix} T(x) \\ 1 \end{pmatrix}$$

get logistic regression model

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

"feature map"

decision boundary
 $\{x \mid w^T \phi(x) > 0\}$

exercise to the reader: try argument 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 in x "

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

Linear logistic regression model

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

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

$$[\text{if } Y = \{ \pm 1 \} \text{ "Radmacher R.V."}]$$

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

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

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

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

max conditional log-likelihood to estimate $\hat{w}_{ML}$

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

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

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

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

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

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

$$\text{because } \sum_i \frac{1}{1 + \exp(-x_i)} (\dots) = 0$$

↑ solve for this

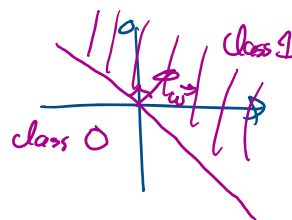
contrast to

$$\text{linear regression } \nabla l(w) = \sum_{i=1}^n x_i [y_i - \underbrace{w^T x_i}_{\text{linear in } w}]$$

13h28

numerical optimization

want to minimize $f(w)$ (unconstrained) [suppose f is differentiable]



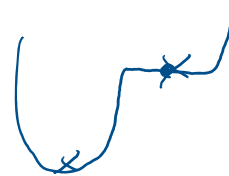
minimization optimization

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

1) gradient descent (1st order method)

start w_0
iterate: $w_{t+1} = w_t - \gamma_t \nabla f(w_t)$

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: $\gamma_t = \frac{1}{L}$ Lipschitz continuity constant for ∇f

$$\text{i.e. } \|\nabla f(w) - \nabla f(w')\|_2 \leq L \|w - w'\|_2 \quad \forall w, w' \in \mathbb{R}^d$$

b) decreasing step-size rule

$$\gamma_t = \frac{c}{t} \leftarrow \text{constant}$$

usually want $\sum_t \gamma_t = \infty$

$$\sum_t \gamma_t^2 < \infty$$

[this is more common for stochastic optimization]

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

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

$$\text{e.g. } g(w, \xi) = f(y_i, h_w(x_i)) \text{ for ERM} \\ \xi = (x_i, y_i)$$

c) choose γ_t by "line search": $\min_{\gamma \in \mathbb{R}} f(w_t + \gamma \nabla f(w_t))$

costly in general

$\hookrightarrow$ instead do approximate search

e.g. Armijo line search [see Boyd's book]

Newton's method (2nd order method)

method: minimizing a quadratic approximation

$$\frac{1}{2} \text{ Hessian} \\ [H(w)]_{ij} = \frac{\partial^2 f(w)}{\partial w_i \partial w_j}$$

Newton's method (2. order method)

method: Minimizing a quadratic approximation

Taylor expansion at w_t

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

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

Taylor's remainder

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

$$\nabla_w Q_t(w) = 0$$

want

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

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

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

Newton's update

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

linear space
 $\leadsto O(d^3)$ time and $O(d^2)$ space

$$H d_t = \nabla f$$

Note for hwk 2: solve for d_t in $H d_t = \nabla f(w_t)$

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

use `numpy.linalg.lstsq` to solve this

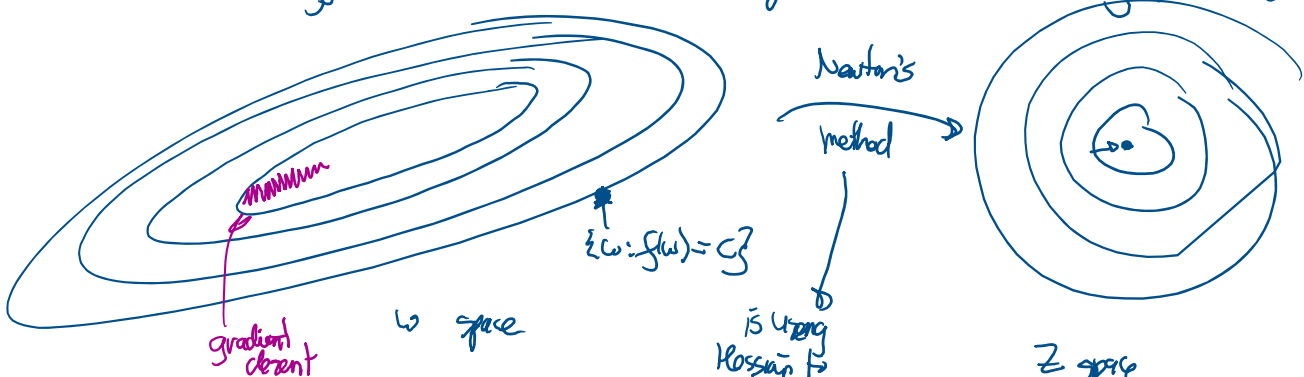
$\leadsto$ much more numerically stable

damped 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's convergence is invariant to rescaling of variables



ω space
 gradient descent

is using
Kossov's
make f "well
conditioned"

z space

$$\frac{1}{2} \omega^T H \omega = C \quad \text{diagonal}$$

$$\downarrow$$

$$H = P^T \tilde{Z} P$$

(H is symmetric psd)

$$\frac{1}{2} \omega^T P^T \tilde{Z} P \omega = C \Leftrightarrow \frac{1}{2} z^T z = C$$

$$z \triangleq \tilde{Z}^{1/2} P \omega$$

exercise to reader: $z_{t+1} = z_t - \gamma \nabla \tilde{f}(z_t)$

$$\Leftrightarrow \omega_{t+1} = \omega_t - \gamma H^{-1} \nabla f(\omega_t)$$

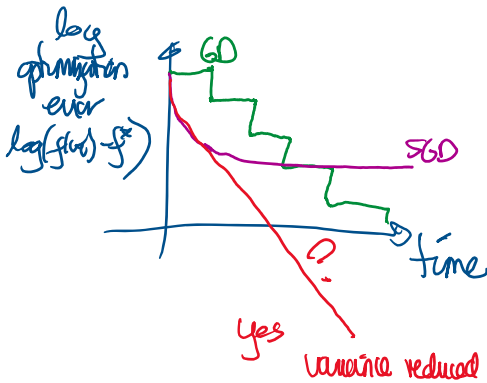
Big data logistic regression

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

instead you can use "incremental gradient method"

eg stochastic gradient descent (SGD): $\omega_{t+1} = \omega_t - \gamma_t \nabla f_{i_t}(\omega_t)$
 where i_t is picked unif at random

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



→ SAG: stochastic average gradient

$$GD: \omega_{t+1} = \omega_t - \gamma_t \frac{1}{n} \sum_{i=1}^n \nabla f_i(\omega_t)$$

$$SAG \quad \omega_{t+1} = \omega_t - \gamma_t \frac{1}{n} \sum_{i=1}^n v_i \quad \text{memory}$$

[2012]

$$\text{where } v_i = \nabla f_i(\omega_{old})$$

at each t , update $v_{i_t} = \nabla f_{i_t}(\omega_t)$

SAGA [2014]

$$\omega_{t+1} = \omega_t - \gamma_t \left(\nabla f_{i_t}(\omega_t) + \underbrace{\frac{1}{n} \sum_{j=1}^n v_j - v_{i_t}}_{\text{SAG/n}} \right)$$

(default method for large logistic regression in Scikit-learn)

variance reduction core

SVRG

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

Newton's method for logistic regression: IRLS

recall for $l(w)$: $\nabla l(w) = \sum_{i=1}^n x_i [y_i - \sigma(w^T x_i)]$
 $H(l(w)) = -\sum_{i=1}^n x_i x_i^T \sigma(w^T x_i) \sigma'(w^T x_i)$

$$v^T H v = -\sum_{i=1}^n \underbrace{(v^T x_i)(x_i^T v)}_{(x_i^T v)^2 \geq 0} \underbrace{\sigma(-) \sigma'(-)}_{\neq 0} \quad v^T H v \leq 0 \quad \forall v \in \mathbb{R}^d$$

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

$$\mu_t \triangleq \begin{pmatrix} \sigma(w_t^T x_1) \\ \vdots \\ \sigma(w_t^T x_n) \end{pmatrix}$$

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 $z_t \triangleq X w_t + D_t^{-1} (y - \mu_t)$

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

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

$\nwarrow$ weights $\searrow$ new target

$$\sum_i \frac{(z_i - w^T x_i)^2}{d_{ii}^{-1}} \quad \text{compare with Gaussian noise model for least square} \quad \sum_i \frac{(y_i - w^T x_i)^2}{\sigma_i^2}$$

Newton's method for logistic regression

= Iterated reweighted least squares (IRLS)