

today: - logistic regression
- numerical optimization

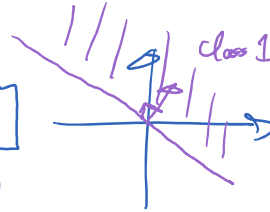
logistic regression model:

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

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

[if $\mathcal{Y} = \{\pm 1\}$ "Rademacher R.v."]

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



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

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

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

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

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

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

$$\begin{aligned} \nabla_w l(w) &= \sum_{i=1}^n x_i \left[\frac{\sigma(v_i) \sigma(-v_i)}{\sigma(v_i)} y_i - \frac{(1-y_i)}{\sigma(-v_i)} \right] \\ &= \sum_i x_i \left[y_i [\underbrace{\sigma(v_i) + \sigma(-v_i)}_1] - \sigma(v_i) \right] \end{aligned}$$

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

need to use numerical methods

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

contrast to linear regression

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

linear w

because $\frac{1}{1 + \exp(w^T x_i)} (\dots) = 0$
solve for this

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)

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

($\Leftrightarrow f(w) - \frac{\mu}{2} \|w\|^2$ is convex in w)

$\Rightarrow$ if $\|\nabla f(w_t)\|^2 < \delta$

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

step-size rules:

a) constant step-size: $\gamma_t = \frac{1}{L}$ $\leftarrow$ Lipschitz continuity constant for ∇f

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

b) decreasing step-size rule

[this is more common for stochastic optimization]

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

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

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

usually want: $\sum_t \gamma_t = \infty$
 $\sum_t \gamma_t^2 < \infty$

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

$\xi = (x_i, y_i)$

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

d_t is direction of update $\rightarrow \nabla f(w_t)$
 costly in general

$\rightarrow$ instead do approximate search

e.g. Armijo line search

[see Boyd's book]

15h50

Newton's method (2nd order method)

metawork: minimizing a quadratic approximation

Taylor expansion at w_t : $f(w) \approx f(w_t) + \nabla f(w_t)^T (w - w_t) + \frac{1}{2} (w - w_t)^T H(w_t) (w - w_t) + \mathcal{O}(\|w - w_t\|^3)$
 $\triangleq Q_t(w)$
 $= Q_t(w) + \mathcal{O}(\|w - w_t\|^3)$
 Taylor's remainder

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

Quadratic model approximation

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

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

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

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

$$w_{t+1} = w_t - H^{-1}(w_t) \nabla f(w_t) \quad \text{Newton's update}$$

↳ inverse Hessian

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

↔

→ $O(d^3)$ time
and $O(d^2)$ space

$$H d_t = \nabla f$$

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

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

use numpy, kindly, lotsq to solve this

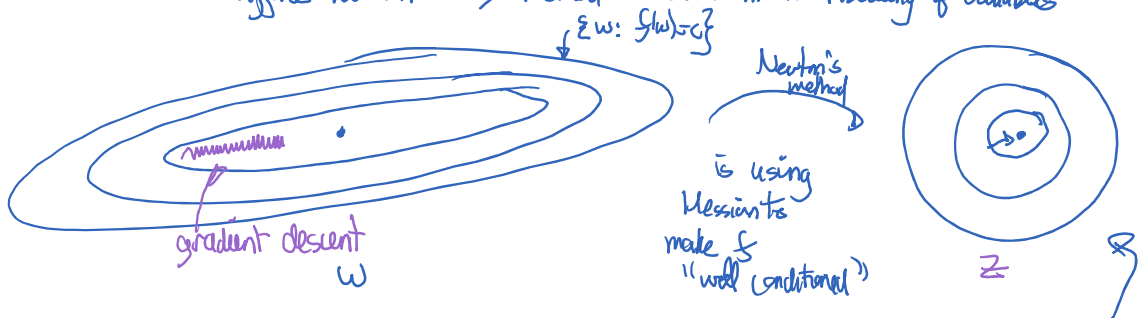
→ much more numerically stable

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



$$\frac{1}{2} w^T H w = c$$

$$\downarrow$$

$$H = P^T \Sigma P$$

(H is symmetric & posd)

$$\frac{1}{2} w^T P^T \Sigma P w = c \Leftrightarrow \frac{1}{2} z^T z = c$$

challenge

$$z = \Sigma^{-1/2} P w$$

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

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

Newton's method for logistic regression: IRLS

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

$$H(l(w)) = -\sum_{i=1}^n x_i x_i^T \sigma(w^T x_i) \sigma'(w^T x_i)$$

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

i.e. $H \preceq 0$
i.e. concave f.d.

Newton is maximizing
instead of minimizing

notation: recall

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

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

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

$$\text{Hessian} = -\sum_i x_i x_i^T \mu_i (1 - \mu_i) = -X^T D(w) X$$

$D_{ii} \triangleq \mu_i (1 - \mu_i)$
diagonal matrix

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

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 $\nwarrow$ new target

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

Newton's method for logistic regression

= iterated reweighted least squares (IRLS)

Big data logistic regression

• big $d \Rightarrow$ cannot do $O(d^2)$ or $O(d^3)$ operations $\Rightarrow$ first order methods

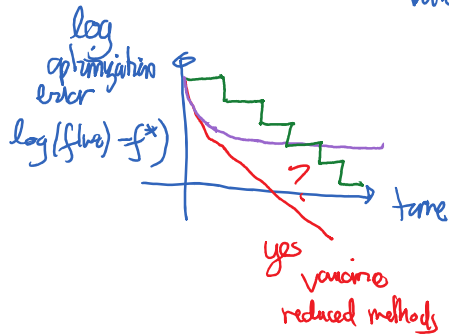
- big $d \Rightarrow$ cannot do $O(d^2)$ or $O(d^3)$ operations $\Rightarrow$ first order methods
- if n is large, you cannot use batch methods $\nabla f(w) = \frac{1}{n} \sum_{i=1}^n \nabla f_i(w)$ ^{gradient of scale data point} $O(nd)$
_{batch gradient}

instead you use "incremental gradient methods"

e.g. stochastic gradient descent (SGD) : $w_{t+1} = w_t - \eta_t \nabla f_{i_t}(w_t)$
 where i_t is picked unif at random $O(d)$

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

batch gradient $\rightarrow$ expensive updates, but faster convergence



SAG: stochastic average gradient

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

$$SAG : w_{t+1} = w_t - \eta_t \frac{1}{n} \sum_{i=1}^n v_i \leftarrow \text{memory} \quad \text{where } v_i = \nabla f_i(w_{old})$$

at each t_i update $v_{i_t} \triangleq \nabla f_{i_t}(w_{i_t})$

$$SAGA : w_{t+1} = w_t - \eta_t \left(\nabla f_{i_t}(w_t) + \frac{1}{n} \sum_{j=1}^n v_j - v_{i_t} \right)$$

(default method for large scale logistic regression in Sikit-Learn)

$\Rightarrow$ VRG

$\frac{1}{n} \sum_{j=1}^n v_j - v_{i_t}$
 $\nabla f_{i_t}(w_t)$
 Variance reduction correction