

## Computational Aspects

of MLE:

If the likelihood equations are non-linear fctns. of the parameters, then usually can't get a "closed-form" solution to these first-order conditions. Use a numerical approximation to the solution.

Lots of different techniques - illustrate with "Methods of Descent". (Max. log L equivalent to min.  $-\log L$ .)

Exactly same problem encountered with non-linear least squares.

General approach:  $\hat{\theta}_1 = \hat{\theta}_0 + s \cdot d(\hat{\theta}_0)$

$\hat{\theta}_0$  = initial value

$s$  = step-length (scalar,  $> 0$ )

$d(\cdot)$  = direction vector

Usually  $d(\cdot)$  depends on gradient vector (or maybe Hessian matrix), at  $\theta$ . Sometimes  $s = s(\text{Hessian})$ , too.

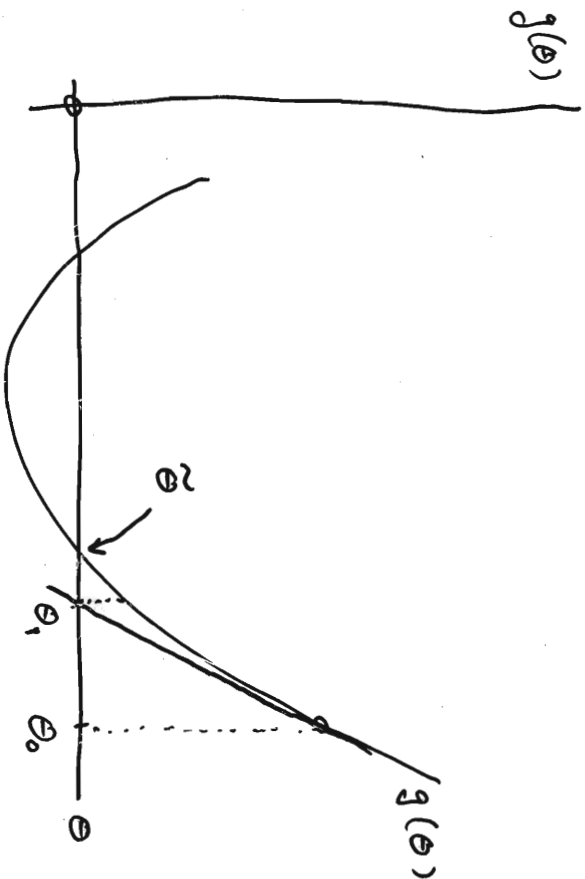
A specific member of family of descent methods is Newton-Raphson algorithm.

First, let's look at a simple geometric example (scalar case). Let  $g(\theta) = -\frac{\partial \log L(\theta)}{\partial \theta}$ . We want to find  $\hat{\theta}$

s.t.  $g(\hat{\theta}) = 0$ . Of course, also need to consider 2nd-order condition.

N.B.: Max. log L  $\Leftrightarrow$  min.  $-\log L = f$ .

Start at an arbitrary point,  $\theta_0$ , move to  $\theta_1$ ,  $\theta_2$  using algorithm, until convergence.



How do we get from  $\theta_0$  to  $\theta_1$ ?

$$\begin{aligned} \text{Slope of tangent} &= \left( \frac{g(\theta_0)}{\theta_0 - \theta_1} \right) \\ &= H(\theta_0) \quad \left[ = \frac{\partial g(\theta)}{\partial \theta} \right] \end{aligned}$$

$$\text{So, } \theta_1 = \theta_0 - g(\theta_0) / H(\theta_0)$$

More generally:

$$\theta_{n+1} = \theta_n - H^{-1}(\theta_n) g(\theta_n)$$

[ In this form, also o.k. for vector,  $\theta$ . ]

Clearly, this is just one possible approach where does this come from, mathematically? (Knowing this will help us see when it may work well, or badly.)

Let  $f(\underline{\theta}) = -\log L(\underline{\theta})$ . Objective is to min.  $f(\underline{\theta})$ . Let's approximate  $f(\cdot)$  using a 2nd-order Taylor's series expansion

$$\begin{aligned} f(\underline{\theta}) &\approx f(\underline{\tilde{\theta}}) + (\underline{\theta} - \underline{\tilde{\theta}})' \left( \frac{\partial f(\underline{\tilde{\theta}})}{\partial \underline{\theta}} \right) \Big|_{\underline{\theta} = \underline{\tilde{\theta}}} \\ &\quad + \frac{1}{2} (\underline{\theta} - \underline{\tilde{\theta}})' \left[ \frac{\partial^2 f(\underline{\tilde{\theta}})}{\partial \underline{\theta} \partial \underline{\theta}} \right] \Big|_{\underline{\theta} = \underline{\tilde{\theta}}} (\underline{\theta} - \underline{\tilde{\theta}}) \end{aligned}$$

(Approximation is valid in neighbourhood of  $\underline{\tilde{\theta}}$ , the value that minimizes  $f(\cdot)$ )

Re-writing this:

$$\begin{aligned} f(\underline{\theta}) &\approx f(\underline{\tilde{\theta}}) + (\underline{\theta} - \underline{\tilde{\theta}})' g(\underline{\tilde{\theta}}) \\ &\quad + \frac{1}{2} (\underline{\theta} - \underline{\tilde{\theta}})' H(\underline{\tilde{\theta}}) (\underline{\theta} - \underline{\tilde{\theta}}) \end{aligned}$$

Differentiate:

$$\begin{aligned} \left[ \frac{\partial f(\underline{\theta})}{\partial \underline{\theta}} \right] &\approx 0 + g(\underline{\tilde{\theta}}) \\ &+ \left(\frac{1}{2}\right)(2) H(\underline{\tilde{\theta}})(\underline{\theta} - \underline{\tilde{\theta}}) \\ &= H(\underline{\tilde{\theta}})(\underline{\theta} - \underline{\tilde{\theta}}) \end{aligned}$$

(because  $g(\underline{\tilde{\theta}}) = 0$ .)

Re-arranging:

$$\begin{aligned} \underline{\tilde{\theta}} &= \underline{\theta} - H^{-1}(\underline{\tilde{\theta}}) \left[ \frac{\partial f(\underline{\theta})}{\partial \underline{\theta}} \right] \\ &= \underline{\theta} - H^{-1}(\underline{\tilde{\theta}}) g(\underline{\theta}) \end{aligned}$$

So, begin with  $\underline{\theta} = \underline{\theta}_0$  & iterate:

$$\begin{aligned} \underline{\theta}_1 &= \underline{\theta}_0 - H^{-1}(\underline{\theta}_1) g(\underline{\theta}_0) \\ \underline{\theta}_2 &= \underline{\theta}_1 - H^{-1}(\underline{\theta}_2) g(\underline{\theta}_1) \\ &\vdots \\ \underline{\theta}_{n+1} &= \underline{\theta}_n - H^{-1}(\underline{\theta}_{n+1}) g(\underline{\theta}_n) \end{aligned}$$

or, approximately —

$$\underline{\theta}_{n+1} = \underline{\theta}_n - H^{-1}(\underline{\theta}_n) g(\underline{\theta}_n)$$

(i) Stop if  $\left| \frac{\theta_{n+1}^i - \theta_n^i}{\theta_n^i} \right| < \epsilon_i$

for  $i=1, 2, \dots, K$ .

(ii) Algorithm corresponds to case of  $s=1$ .

(iii)  $d(\underline{\theta}_n) = -H^{-1}(\underline{\theta}_n) g(\underline{\theta}_n)$

(iv) Algorithm will fail if  $H(\cdot)$  singular.

(v) If  $H(\cdot)$  becomes n.d., we'll

locate a max. of  $f(\cdot)$  & hence a min. of  $\log L$ . So, need to ensure that  $H(\underline{\tilde{\theta}})$  is p.d.

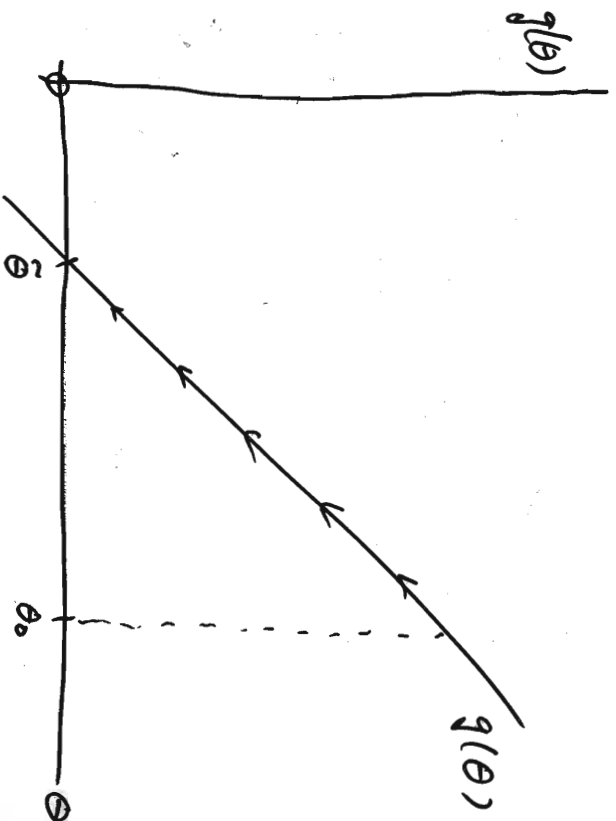
(vi) Algorithm may locate a local min.

(depending on initial value), or may oscillate & never converge.

(vii) Algorithm should work well if

$\log L(\underline{\theta})$  is approximately quadratic, at least in neighbourhood of max.

In fact, if  $f(\theta) = -\log L(\theta)$  is exactly quadratic, then convergence is immediate. (In this case,  $g(\theta)$  is linear.)



Need to experiment with different initial values to ensure we locate a global max. of  $\log L$ .

Example:

$$f(\theta) = \theta^3 - 2\theta^2 - 1$$

We can locate turning points analytically —

$$f'(\theta) = 3\theta^2 - 4\theta = \theta(3\theta - 4) = 0$$

$$\Rightarrow \theta = 0, \frac{4}{3}.$$

$$f''(\theta) = 6\theta - 4 \quad ; \quad \begin{cases} f''(0) = -4 & (\text{max.}) \\ f''(\frac{4}{3}) = 4 & (\text{min.}) \end{cases}$$

Now, let's see how the algorithm goes:

Start at  $\theta_0 = 1$ :

$$g(\theta_0) = f'(\theta_0) = -4 \quad ; \quad H(\theta_0) = f''(\theta_0) = 2$$

$$\text{So, } \theta_1 = \theta_0 - H(\theta_0)^{-1} g(\theta_0) = 1 - (-\frac{1}{2}) = 1.5$$

$$g(\theta_1) = \frac{3}{4} \quad ; \quad H(\theta_1) = 5$$

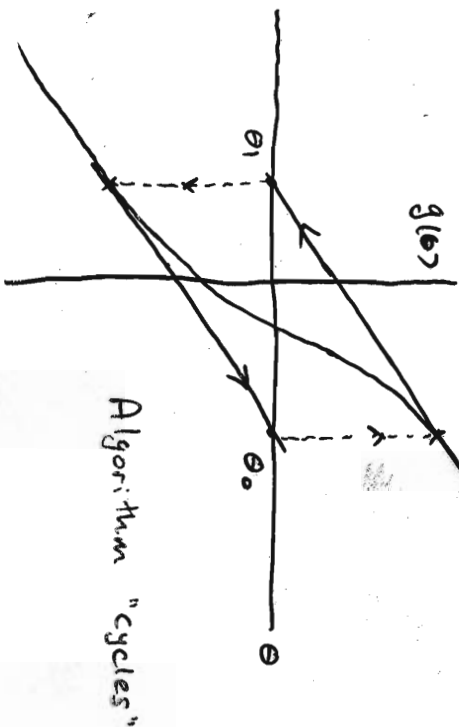
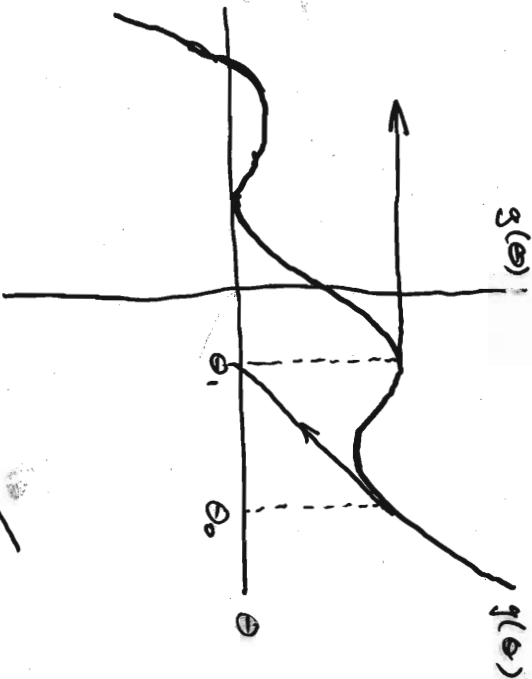
$$\text{So, } \theta_2 = 1.5 - (\frac{3/4}{5}) = 1.35$$

Rapidly get to  $\tilde{\theta} = 1.333$ .

What happens if we choose  $\theta_0 = -1$ ?

Things that can go wrong:

$H(\cdot)$  singular  
at  $\theta_1$ .



Algorithm "cycles"

"Concentrating" the Likelihood Fcn.

Sometimes the parameter vector falls into 2 natural parts:  $\theta = (\theta_1, \theta_2)$ .

So,  $L(\theta | y) = L(\theta_1, \theta_2 | y)$ . In addition, often the MLE for  $\theta_2$  can be written as a fcn. of the MLE for  $\theta_1$ :

$$\tilde{\theta}_2 = f(\tilde{\theta}_1). \text{ In this case,}$$

$$L(\theta_1, \theta_2 | y) = L[\theta_1, f(\theta_1) | y] \\ = L_c[\theta_1 | y]$$

where  $L_c(\cdot)$  is the concentrated likelihood fcn. The point is that we can then maximize  $L_c(\cdot)$  with respect to just  $\theta_1$ .

If we can get  $\tilde{\theta}_2 = f(\tilde{\theta}_1)$  analytically this means that the numerical maximization has to be done only w.r.t.  $\theta_1$  & not  $\theta_2$ , which simplifies matters significantly.

1.1.

N.B.: "Concentrating" the L.F. is not always an option - worth keeping in mind, though.

Example: (concentrating is unnecessary, but it still works)

$$y = X\beta + \varepsilon; \quad \varepsilon \sim N[0, \sigma^2 I]$$

$$p(y|\beta, \sigma) = \left(\frac{1}{\sigma\sqrt{2\pi}}\right)^n \exp\left[-\frac{1}{2\sigma^2}(y-X\beta)'(y-X\beta)\right]$$

$$\log L(\beta, \sigma|y) = -\frac{n}{2} \log \sigma^2 - \frac{n}{2} \log 2\pi - \frac{1}{2\sigma^2} (y-X\beta)'(y-X\beta)$$

$$(\partial \log L / \partial \sigma^2) = \frac{n}{2\sigma^2} + \frac{1}{2\sigma^4} (y-X\beta)'(y-X\beta) = 0$$

$$\Rightarrow \hat{\sigma}^2 = \frac{1}{n} (y-X\beta)'(y-X\beta)$$

$$\log L_c(\beta|y) = -\frac{n}{2} \log \left[ \frac{1}{n} (y-X\beta)'(y-X\beta) \right]$$

$$-\frac{n}{2} \log 2\pi - \frac{n}{2}$$

$$(\partial \log L_c / \partial \beta) = \frac{-\frac{n}{2}}{\frac{1}{n} (y-X\beta)'(y-X\beta)} \cdot \left[ \frac{\partial}{\partial \beta} (y-X\beta)'(y-X\beta) \right]$$

$$= 0$$

$$\Rightarrow \frac{\partial}{\partial \beta} [ (y-X\beta)'(y-X\beta) ] = 0$$

$$\Rightarrow \hat{\beta} = (X'X)^{-1} X'y$$

$$\text{so } \hat{\sigma}^2 = \frac{1}{n} (y-X\hat{\beta})'(y-X\hat{\beta})$$

1.2

So, "concentrating" the L.F. certainly works. Imagine that we knew how to differentiate w.r.t. a scalar, but not a vector. Then, we could have obtained  $L_c(\beta|y)$  above & then would have had to  $\max_{(\beta)} \{L_c(\beta|y)\}$  w.r.t.  $\beta$ .

That is, maximize

$$L_c(\beta|y) = -\frac{n}{2} \log \left[ \frac{1}{n} (y-X\beta)'(y-X\beta) \right] - \frac{n}{2} \log(2\pi) - \frac{n}{2}$$

$$\text{or, } \max_{\beta} \left\{ -\frac{n}{2} \log (y-X\beta)'(y-X\beta) \right\} \text{ w.r.t. } \beta.$$

We'll encounter situations where we can maximize partly analytically, but need to use numerical algorithm for the rest of problem. "Concentrating" will reduce dimension of space for numerical optimization.

Inequality Constraints:

Sometimes we'll want to impose an inequality constraint on one or more of our MLE's. For example:  $\tilde{\theta}_1 > 0$ ;  $0 < \tilde{\theta}_2 < 1$ . Most packages don't allow this. However, there are some tricks we can use to impose such constraints.

Example: We want  $\tilde{\theta}_1 \geq 0$ . So, replace  $\theta_1$  with  $(\phi_1)^2$ , which will be  $\geq 0$ . No constraint is placed on  $\phi_1$ ;  $\tilde{\theta}_1 = \tilde{\phi}_1^2$ .

Example: We want  $0 \leq \theta_2 \leq 1$ . Replace  $\theta_2$  by  $1/(1 + \phi_2)$ , or replace  $\theta_2$  by  $\exp(\phi)/(\exp(\phi) + 1)$  & then  $\tilde{\theta}_2 = 1/(1 + \tilde{\phi}^2)$ , or  $1/(1 + \exp(\tilde{\phi}))$  will satisfy the constraint.