

# Hypothesis Testing in a Likelihood Framework

## I Testing Simple hypotheses

Let  $X_1, \dots, X_n$  be iid with pdf  $f(x_i; \theta)$

and assume  $\theta$  is a scalar.

The hypothesis to be tested is simple & two sided:

$$H_0: \theta = \theta_0 \quad \text{vs.} \quad H_a: \theta \neq \theta_0$$

Statistical Test: Decision rule based on data to either reject  $H_0$  or not reject  $H_0$ .

### Likelihood Ratio Statistic

Consider the likelihood ratio

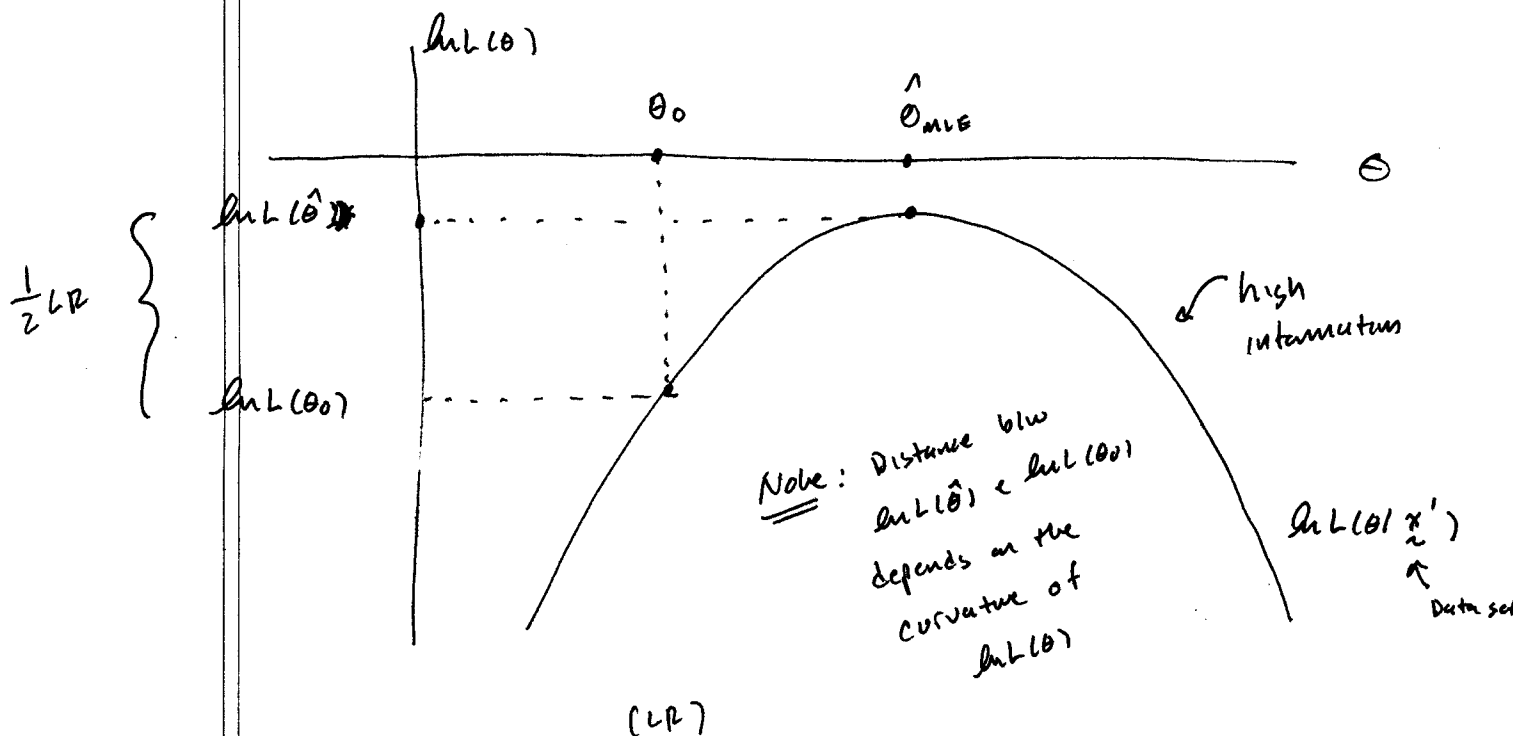
$$\lambda = \frac{L(\theta_0; \underline{x})}{L(\hat{\theta}_{MLE}; \underline{x})} \quad \left( = \frac{\max_{\text{s.t. } \theta = \theta_0} L(\theta; \underline{x})}{\max_{\substack{\theta \\ \underline{x} \text{ unrestricted}}} L(\theta; \underline{x})} \right)$$

= ratio of "restricted" to unrestricted likelihoods.

By construction  $0 < \lambda \leq 1$ .

If  $H_0: \theta = \theta_0$  is true then  $\lambda \approx 1$  and

if  $H_0: \theta \neq \theta_0$  is not true then  $\lambda < 1$



The LR statistic  $\lambda$  is a simple transformation of  $\lambda$

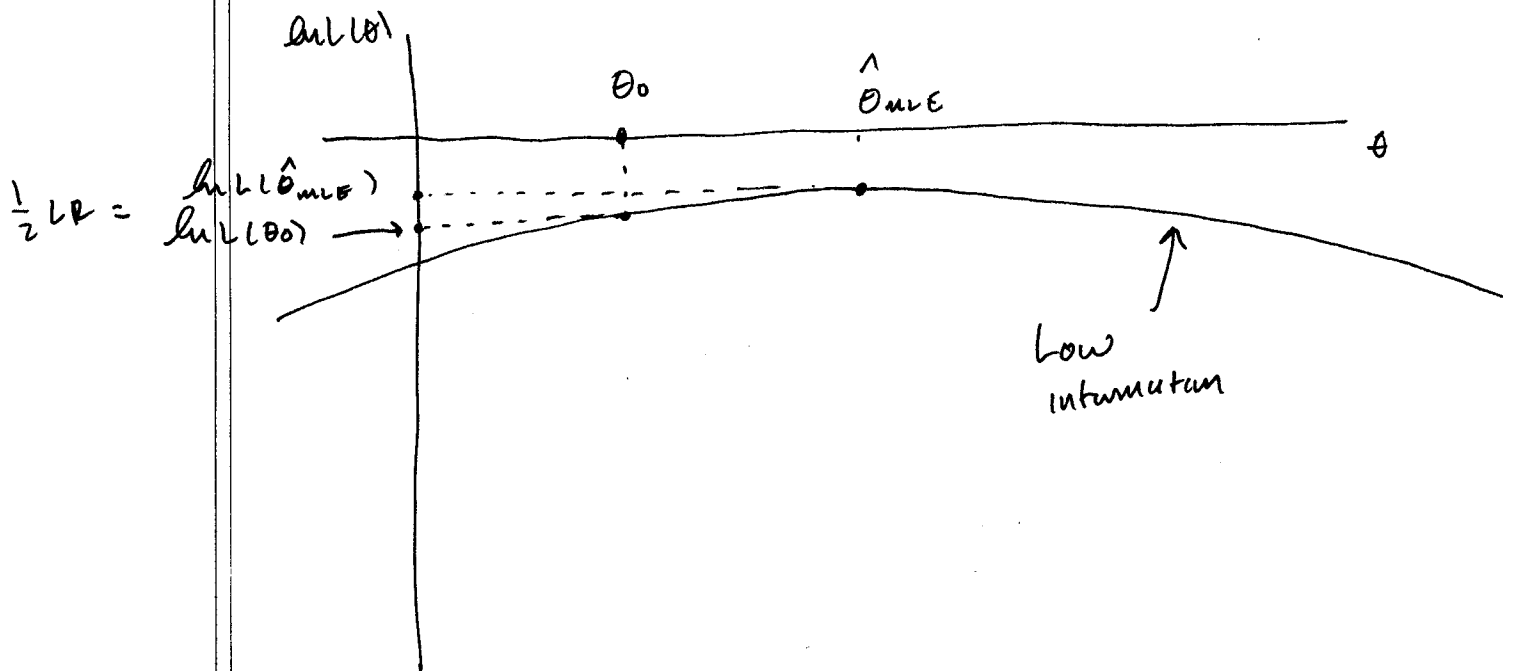
s.t. LR is large if  $H_0: \theta = \theta_0$  is not true and

is small if  $H_0: \theta = \theta_0$  is true :

$$LR = -2 \ln \lambda$$

$$= -2 \left[ \ln L(\theta_0; \tilde{x}) - \ln L(\hat{\theta}_{MLE}; \tilde{x}) \right]$$

From the diagram, notice that the distance between  $\ln L(\hat{\theta}_{MLE})$  and  $\ln L(\theta_0)$  depends on the curvature of  $\ln L(\theta)$  near  $\hat{\theta}_{MLE}$ . If the curvature is sharp (i.e. high information) then  $LR$  will be large for  $\theta_0$  values away from  $\hat{\theta}_{MLE}$ . If, however, the curvature of  $\ln L(\theta)$  is flat (i.e. low information) as depicted below



then  $LR$  will be small for  $\theta_0$  values away from  $\hat{\theta}_{MLE}$ .

### Result (no proof)

Under very general regularity conditions, if  $H_0: \theta = \theta_0$  is true then

$$LR = -2 \ln \lambda \xrightarrow{d} \chi^2(1)$$

$\nearrow$   
 $\chi^2$  random variable  
with 1 d.f.

Note: The # of d.f. is based on the # of restrictions being tested.

We reject  $H_0: \theta = \theta_0$  at, say, 5% level if  $LR > \chi^2_{0.95}(1) = 95\%$  quantity of  $\chi^2(1)$  ~~is~~ distrib

### Wald Statistic

The Wald Statistic is based directly on the asymptotic distribution of  $\hat{\theta}_{MLE}$ :

$$\hat{\theta}_{MLE} \overset{\Delta}{\sim} N(\theta, I(\theta; \tilde{x})^{-1})$$

$\nearrow$   
Sample information matrix

Where  $I(\theta; \underline{x})$  is estimated using  $\hat{\theta}_{MLE}$ :

$$\hat{\text{Var}}(\hat{\theta}_{MLE}) = I(\hat{\theta}_{MLE}; \underline{x})^{-1} = \text{estimated asymptotic variance of } \hat{\theta}_{MLE}$$

Given these results, we know that the t-ratio

$$\frac{\hat{\theta}_{MLE} - \theta_0}{\sqrt{\hat{\text{Var}}(\hat{\theta}_{MLE})}} = (\hat{\theta}_{MLE} - \theta_0) \sqrt{I(\hat{\theta}_{MLE}; \underline{x})}$$

assuming  $H_0: \theta = \theta_0$  is true.

is asymptotically  $N(0, 1)$ . ~~Therefore~~ Using the

CMT it follows that the square of the t-ratio

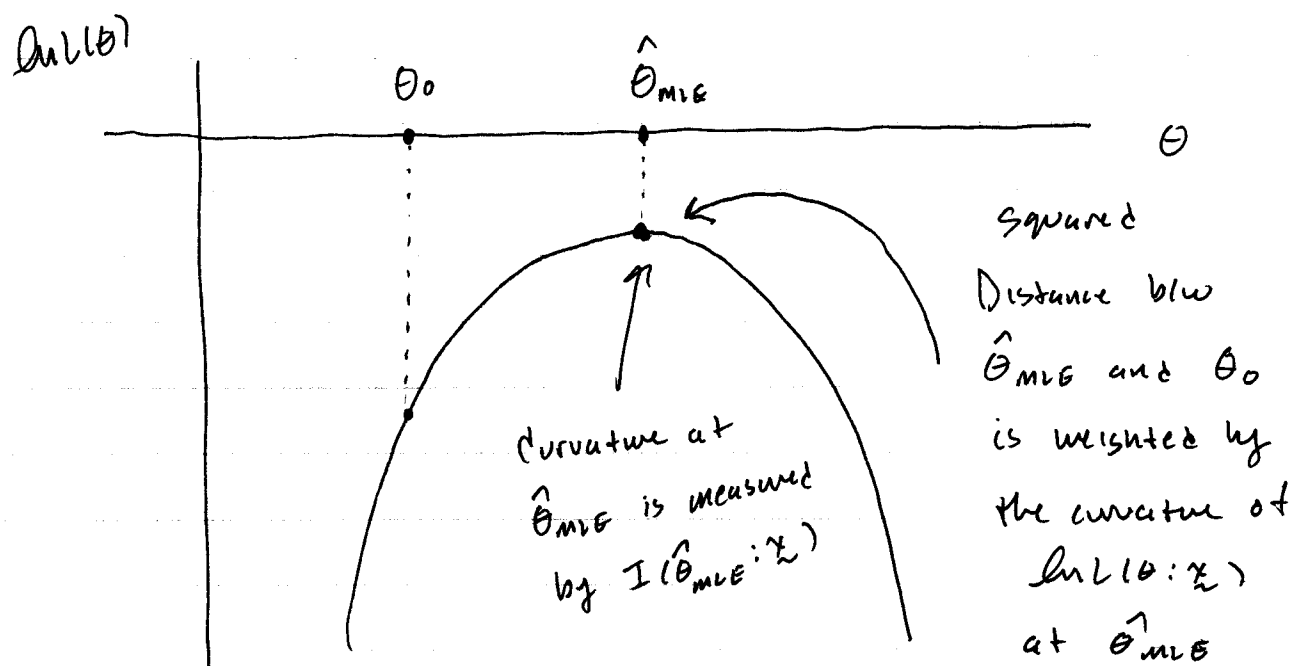
is asymptotically  $\chi^2(1)$ . The Wald statistic is defined

to be simply the square of this t-ratio:

$$\text{Wald} = \frac{(\hat{\theta}_{MLE} - \theta_0)^2}{\hat{\text{Var}}(\hat{\theta}_{MLE})} = (\hat{\theta}_{MLE} - \theta_0)^2 \cdot I(\hat{\theta}_{MLE}; \underline{x})$$

Result: Under  $H_0$ ,  $\text{Wald} \xrightarrow{d} \chi^2(1)$

The intuition behind the Wald statistic is illustrated in the following figure



If the curvature at  $\hat{\theta}_{MLE}$  is sharp then  $I(\hat{\theta}_{MLE}; \mathbf{X})$  is big and the squared distance  $(\hat{\theta}_{MLE} - \theta_0)^2$  gets blown up when constructing the Wald statistic.

If the curvature at  $\hat{\theta}_{MLE}$  is low then  $I(\hat{\theta}_{MLE}; \mathbf{X})$  is small and the squared distance  $(\hat{\theta}_{MLE} - \theta_0)^2$  does not get blown up when constructing the Wald statistic.

## LM Statistic

With MLE,  $\hat{\theta}_{MLE}$  solves the first order condition

$$0 = \frac{\partial \ln L(\hat{\theta}_{MLE} | \underline{x})}{\partial \theta} = S(\hat{\theta}_{MLE}; \underline{x})$$

If  $H_0: \theta = \theta_0$  is true then we should expect that

$$0 \approx \frac{\partial \ln L(\theta_0; \underline{x})}{\partial \theta} = S(\theta_0; \underline{x})$$

and if  $H_0: \theta = \theta_0$  is not true then

$$0 \neq \frac{\partial \ln L(\theta_0; \underline{x})}{\partial \theta} = S(\theta_0; \underline{x})$$

The LM test is based on how far  $S(\theta_0; \underline{x})$  is from zero.

Recall the following properties of  $S(\theta; \underline{x})$ . If  $H_0: \theta = \theta_0$  is true then

$$(i) E\{S(\theta_0: \underline{x})\} = 0$$

$$(ii) \text{var}(S(\theta_0: \underline{x})) = I(\theta_0: \underline{x})$$

Further, it can be shown that

$$S(\theta_0: \underline{x}) \overset{A}{\sim} N(0, I(\theta_0: \underline{x}))$$

The LM statistic for testing  $H_0: \theta = \theta_0$  is then

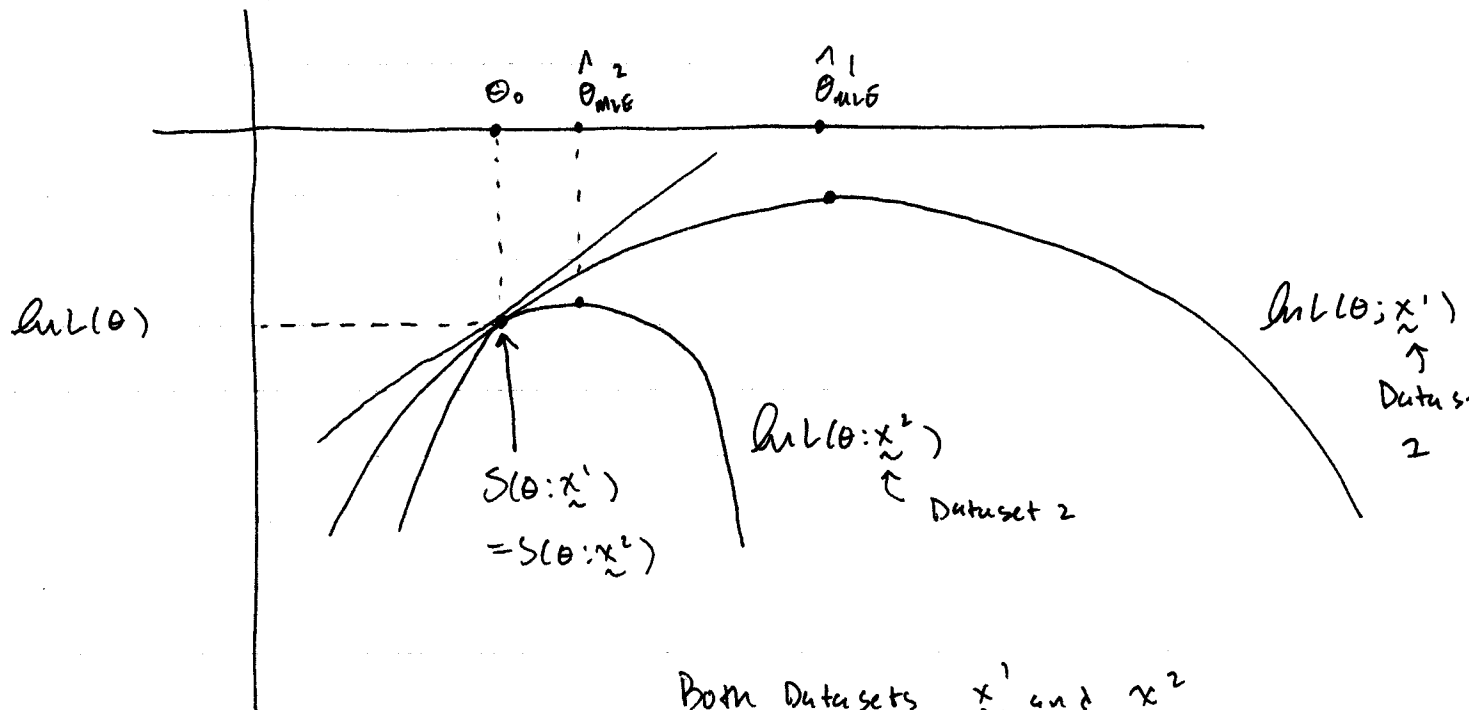
$$LM = \frac{S(\theta_0: \underline{x})^2}{I(\theta_0: \underline{x})} = S(\theta_0: \underline{x})^2 \cdot I(\theta_0: \underline{x})^{-1}$$

← Note: The LM statistic is based on the "restricted" estimate of  $\theta$

Result: Under  $H_0: \theta = \theta_0$

$$LM \xrightarrow{d} \chi^2(1)$$

The basic intuition behind the LM statistic is given in the following diagram=



Both Datasets  $\tilde{x}^1$  and  $\tilde{x}^2$   
have the same score values  
at  $\theta_0$ :

$$S(\theta_0; \tilde{x}^1) = S(\theta_0; \tilde{x}^2)$$

$$LM^1 = S(\theta_0; \tilde{x}^1)^2 \cdot I(\theta_0; \tilde{x}^1)^{-1}$$

$$LM^2 = S(\theta_0; \tilde{x}^2)^2 \cdot I(\theta_0; \tilde{x}^2)^{-1}$$

Here  $LM^2 < LM^1$  since  $I(\theta_0; \tilde{x}^2)^{-1} > I(\theta_0; \tilde{x}^1)^{-1}$

i.e. the Dataset that produces the largest curvature

at  $\theta_0$  indicates that  $\theta_0$  is closer to  $\hat{\theta}_{MLE}$ .

## Test statistics for General Nonlinear hypotheses

---

Let  $X_1, \dots, X_n$  be iid with pdf  $f(x; \theta)$

where  $\theta \in \mathbb{R}^k$ . Consider testing the general non-linear hypothesis

$$H_0: g(\theta) = 0 \quad \text{vs.} \quad g(\theta) \neq 0$$

$\underset{\sim}{\theta} \quad \underset{\sim}{\theta} \quad \underset{J \times 1}{\sim}$

Here

$$g(\theta) = \begin{pmatrix} g_1(\theta) \\ g_2(\theta) \\ \vdots \\ g_J(\theta) \end{pmatrix}$$

Scalar functions  
= continuous and 2 times differentiable

and  $g(\theta) = 0$  places  $J \times k$  total restrictions on  $\theta$

Example  $X \sim N(\mu, \sigma^2)$ ,  $\theta = (\mu, \sigma^2)'$

$$H_0: \mu = \mu_0 \quad \text{vs.} \quad H_1: \mu \neq \mu_0$$

$$\Leftrightarrow H_0: g(\theta) = 0 \quad \text{vs.} \quad H_1: g(\theta) \neq 0$$

$$\text{with } g(\theta) = \mu - \mu_0.$$

Define

$$G(\theta) = \frac{\partial g(\theta)}{\partial \theta'} =$$

$$= \begin{bmatrix} \frac{\partial g_1(\theta)}{\partial \theta_1} & \frac{\partial g_1(\theta)}{\partial \theta_2} & \dots & \frac{\partial g_1(\theta)}{\partial \theta_K} \\ \vdots & \vdots & & \vdots \\ \frac{\partial g_J(\theta)}{\partial \theta_1} & \frac{\partial g_J(\theta)}{\partial \theta_2} & \dots & \frac{\partial g_J(\theta)}{\partial \theta_K} \end{bmatrix}$$

and assume  $\text{rank}(G) = J < K$ .

LR Test

$$\lambda = \max_{\theta} L(\theta; \underline{x}) \quad \text{s.t.} \quad g(\theta) = 0$$

---


$$\max_{\theta} L(\theta; \underline{x})$$

$$= \frac{L(\tilde{\theta}_{MLE}; \underline{x})}{L(\hat{\theta}_{MLE}; \underline{x})}$$

} Requires both  
restricted &  
unrestricted  
estimates!

where  $\tilde{\theta}_{MLE}$  = restricted MLE computed  
subject to  $g(\tilde{\theta}) = \tilde{a}$ ; i.e.

$$g(\tilde{\theta}_{MLE}) = 0$$

and  $\hat{\theta}_{MLE}$  = unrestricted MLE - In general,  $g(\hat{\theta}_{MLE}) \neq 0$

Then

$$LR = -2 \ln \lambda = -2 \left[ \ln L(\tilde{\theta}_{MLE}; \mathbf{x}) - \ln L(\hat{\theta}_{MLE}; \mathbf{x}) \right]$$

and under  $H_0$

$$LR \xrightarrow{d} \chi^2(J)$$

where  $J$  = # of restrictions under  $H_0$ .

### LM Test

Idea: If  $H_0: g(\tilde{\theta}) = 0$  is true then

$$S(\tilde{\theta}_{MLE}; \mathbf{x}) \approx 0$$

and

$$I(\tilde{\theta}_{MLE}; \mathbf{x}) \text{ consistently estimates } \text{Var}(S(\tilde{\theta}_{MLE}; \mathbf{x}))$$

Then

$$LM = \underset{1 \times K}{S(\tilde{\theta}_{MLE}; \tilde{x})}' \underset{K \times K}{I(\tilde{\theta}_{MLE}; \tilde{x})}^{-1} \underset{K \times 1}{S(\tilde{\theta}_{MLE}; \tilde{x})}$$

and under  $H_0$

$$LM \xrightarrow{d} \chi^2(J)$$

### Wald Test

Here the Wald test is based on the asymptotic normality of  $\hat{\theta}_{MLE}$  and the "delta method":

$$\hat{\theta}_{MLE} \underset{\sim}{\overset{A}{\sim}} N(\theta, I(\hat{\theta}_{MLE}; \tilde{x})^{-1})$$

$$\underset{\sim}{g}(\hat{\theta}_{MLE}) \underset{J \times 1}{\overset{A}{\sim}} N\left(g(\theta), \underbrace{g(\hat{\theta}_{MLE}) I(\hat{\theta}_{MLE}; \tilde{x})^{-1} g(\hat{\theta}_{MLE})}_{J \times J}\right)$$

$$\begin{aligned} \text{Wald} &= \underset{1 \times J}{g(\hat{\theta}_{MLE})}' \left[ \underset{J \times K}{g(\hat{\theta}_{MLE})} \underset{K \times K}{I(\hat{\theta}_{MLE}; \tilde{x})}^{-1} \underset{K \times J}{g(\hat{\theta}_{MLE})} \right]^{-1} \underset{J \times 1}{g(\hat{\theta}_{MLE})} \\ &\xrightarrow{d} \chi^2(J) \quad \text{under } H_0: g(\theta) = \underset{\sim}{0}. \end{aligned}$$