

Bayesian Interpretations of Heteroskedastic Consistent Covariance Estimators Using the Informed Bayesian Bootstrap

Dale J. Poirier

University of California, Irvine

September 1, 2008

Abstract

This paper provides Bayesian rationalizations for White's *heteroskedastic consistent* (HC) covariance estimator and various modifications of it. An informed Bayesian bootstrap provides the statistical framework.

1. Introduction

Often researchers find it useful to recast frequentist procedures in Bayesian terms, so as to get a clearer understanding of how and when the procedures work. This paper takes this approach with respect to the consistent, in the face of heteroskedasticity of *unknown* form, covariance matrix estimator proposed by White (1980) for the case of stochastic regressors, building on the fixed regressor analysis of Eicker (1963, 1967), and the misspecified maximum likelihood analysis of Huber (1967).

2. Preliminaries

This section reviews a Bayesian conjugate analysis of a multinomial sampling distribution. Let $\mathbf{z}$ denote a discrete random variable belonging to one of J categories indexed ($j = 1, 2, \dots, J$) with associated probabilities $\text{Prob}(\mathbf{z} = j | \boldsymbol{\theta}) = \theta_j$ ($j = 1, 2, \dots, J$), where $\boldsymbol{\theta} = [\theta_1, \theta_2, \dots, \theta_J]' \in \Theta^J$ and Θ^J denotes the unit simplex in $\Re^J$. n independent draws $\mathbf{z} = [\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_n]'$ on $\mathbf{z}$ imply the multinomial likelihood function

$$\mathcal{L}(\boldsymbol{\theta}; \mathbf{z}) = \left(\frac{n!}{n_1! n_2! \dots n_J!} \right) \prod_{j=1}^J \theta_j^{n_j} \propto \prod_{j=1}^J \theta_j^{n_j}, \quad (1)$$

where $N_j = n_j$, the number of z_j ($i = 1, 2, \dots, n$) equal to j ($j = 1, 2, \dots, J$), are sufficient statistics for θ . The natural conjugate prior for θ is the Dirichlet distribution $D_J(\mathbf{y})$ with density

$$f^D(\theta|\mathbf{y}) = \frac{\Gamma\left(\sum_{j=1}^J \mathbf{y}_j\right)}{\prod_{j=1}^J \Gamma(\mathbf{y}_j)} \prod_{j=1}^J \theta_j^{\mathbf{y}_j-1}, \quad \theta \in \Theta^J, \quad (2)$$

given $\mathbf{y}_j > \mathbf{0}$ ($j = 1, 2, \dots, J$) and $\mathbf{y} = [\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_J]'$.¹ Without loss of generality assume the first m categories have $N_j = n_j > 0$. It follows from Bayes' Theorem that the posterior distribution of θ is the Dirichlet distribution $D_J(\bar{\mathbf{v}})$, where $\bar{\mathbf{v}} = [\bar{\mathbf{v}}_1, \bar{\mathbf{v}}_2, \dots, \bar{\mathbf{v}}_J]'$, $\bar{\mathbf{v}}_j = \mathbf{y}_j + \mathbf{n}_j$ ($j = 1, 2, \dots, m$), and $\bar{\mathbf{v}}_j = \mathbf{y}_j$ ($j = m+1, m+2, \dots, J$). The marginal mass function of z is the multinomial-Dirichlet mass function

$$f^{MD}(z|\mathbf{y}) = \frac{\Gamma\left(\sum_{j=1}^J \mathbf{y}_j\right)}{\Gamma\left(\mathbf{n} + \sum_{j=1}^J \mathbf{y}_j\right)} \prod_{j=1}^m \frac{\Gamma(\mathbf{y}_j + \mathbf{n}_j)}{\Gamma(\mathbf{y}_j)}. \quad (3)$$

The predictive mass function for a new observation z_{N+1} is

$$\text{Prob}(z_{N+1} = j|z) = E(\theta_j|z) = \frac{\bar{\mathbf{v}}_j}{\sum_{i=1}^J \bar{\mathbf{v}}_i} \quad (j = 1, 2, \dots, J). \quad (4)$$

Sampling from a Dirichlet distribution can be accomplished by simulating J independent gamma random variables $\mathbf{g} = [\mathbf{g}_1, \mathbf{g}_2, \dots, \mathbf{g}_J]'$ with means/variances $\mathbf{v}_j$ ($j = 1, 2, \dots, J$) and unit scale parameters, i.e., with densities

$$f_{\gamma}(\mathbf{g}_j|\mathbf{v}_j) = [\Gamma(\mathbf{v}_j)]^{-1} \mathbf{g}_j^{\mathbf{v}_j-1} \exp(-\mathbf{g}_j), \quad \mathbf{g}_j > 0, \mathbf{v}_j > 0, \quad (j = 1, 2, \dots, J), \quad (5)$$

and then forming

$$\theta_j = \frac{\mathbf{g}_j}{\mathbf{g}_1 + \mathbf{g}_2 + \dots + \mathbf{g}_J} \quad (j = 1, 2, \dots, J). \quad (6)$$

¹ Note: $E(\theta|\mathbf{y}) = (\mathbf{1}_J' \mathbf{y})^{-1} \mathbf{y}$ and $\text{Var}(\theta|\mathbf{y}) = [(\mathbf{1}_J' \mathbf{y})^2 (\mathbf{1}_J' \mathbf{y} + 1)]^{-1} [(\mathbf{1}_J' \mathbf{y}) \text{diag}\{\mathbf{y}\} - \mathbf{y} \mathbf{y}']$.

The resulting θ_j constitutes one draw from the Dirichlet distribution $\mathbf{D}_j(\mathbf{v})$. When $\mathbf{v} = \mathbf{y}$, the draw θ is from the prior density $f_D(\theta|\mathbf{y})$; when $\mathbf{v} = \bar{\mathbf{v}}$, the draw θ is from the posterior density $f_D(\theta|\bar{\mathbf{v}})$.

3. Bayesian Bootstrap

Interpret the multinomial categories of $\mathbf{z}$ as J distinct values taken by a $(K+1) \times 1$ vector $\mathbf{z} = [\mathbf{y}, \mathbf{x}']'$, where $\mathbf{y}$ is scalar and $\mathbf{x}$ is $K \times 1$.² In other words, the J categories are reinterpreted as the support points for the joint distribution of $\mathbf{z}$, i.e., $\mathbf{z} = j$ is reinterpreted as $\mathbf{z} = \mathbf{z}_j$ where $\mathbf{z}_j = [\mathbf{y}_j, \mathbf{x}_j']'$. Suppose m of the J support points for $\mathbf{z}$ are observed in a sample of size $n \geq m$ and $\tilde{J}$ others are not observed. Denote the m observed points by $\mathbf{z}^{(1)} = [\mathbf{y}^{(1)}, \mathbf{X}^{(1)}]'$, where $\mathbf{y}^{(1)}$ is $m \times 1$ and $\mathbf{X}^{(1)}$ is $m \times K$, and denote the $\tilde{J}$ unobserved support points by $\mathbf{z}^{(2)} = [\mathbf{y}^{(2)}, \mathbf{X}^{(2)}]'$, where $\mathbf{y}^{(2)}$ is $\tilde{J} \times 1$ and $\mathbf{X}^{(2)}$ is $\tilde{J} \times K$. Note that $N_j = 0$ ($j = m + 1, m + 2, \dots, m + \tilde{J}$). Similarly, partition θ , $\mathbf{g}$, $\mathbf{y}$, and $\bar{\mathbf{v}}$ into m and $\tilde{J}$ components corresponding to $\mathbf{z}^{(1)}$ and $\mathbf{z}^{(2)}$: $\theta = [\theta^{(1)'}, \theta^{(2)'}]'$, $\mathbf{g} = [\mathbf{g}^{(1)'}, \mathbf{g}^{(2)'}]'$, $\mathbf{y} = [\mathbf{y}^{(1)'}, \mathbf{y}^{(2)'}]'$, and $\bar{\mathbf{v}} = [\bar{\mathbf{v}}^{(1)'}, \bar{\mathbf{v}}^{(2)'}]'$. The **standard case** corresponds to data that are distinct and exhaust the support: $J = n = m$, $N_1 = N_2 = \dots = N_J = 1$, $\tilde{J} = 0$, and $\mathbf{z}^{(2)}$ is null. See Chamberlain and Imbens (2003), Hahn (1977), Lancaster (2003; 2004, Section 3.4), and Ragusa (2007). The **general case** is any non-standard case.

The Bayesian bootstrap of Rubin (1981) begins by selecting a parameter/functional of interest, say, $\beta = \beta(\theta)$. For example, in the context of a linear model, suppose interest focuses on the moment condition $E_{\mathbf{z}|\theta}[\mathbf{x}(\mathbf{y} - \mathbf{x}'\beta)] = \mathbf{0}_K$.³ Then $\beta = \beta(\theta)$ are the coefficients associated with the linear projection of $\mathbf{y}$ on $\mathbf{x}$: $\beta = [E_{\mathbf{z}|\theta}(\mathbf{x}'\mathbf{x})]^{-1} E_{\mathbf{z}|\theta}(\mathbf{x}'\mathbf{y})$, where

$$E_{\mathbf{z}|\theta}(\mathbf{x}'\mathbf{x}) = \mathbf{X}' \text{diag}\{\theta\} \mathbf{X} = \mathbf{X}^{(1)'} \text{diag}\{\theta^{(1)}\} \mathbf{X}^{(1)} + \mathbf{X}^{(2)'} \text{diag}\{\theta^{(2)}\} \mathbf{X}^{(2)}, \quad (7)$$

$$E_{\mathbf{z}|\theta}(\mathbf{x}'\mathbf{y}) = \mathbf{X}' \text{diag}\{\theta\} \mathbf{y} = \mathbf{X}^{(1)'} \text{diag}\{\theta^{(1)}\} \mathbf{y}^{(1)} + \mathbf{X}^{(2)'} \text{diag}\{\theta^{(2)}\} \mathbf{y}^{(2)}, \quad (8)$$

with $\mathbf{X} = [\mathbf{X}^{(1)'}, \mathbf{X}^{(2)'}]'$, $\mathbf{y} = [\mathbf{y}^{(1)'}, \mathbf{y}^{(2)'}]'$, and $\text{diag}\{\cdot\}$ denotes a diagonal matrix with diagonal

²Chamberlain and Imbens (2003, p. 12) note that because J can be arbitrarily large and our data are measured with finite precision, the finite support assumption is arguably not restrictive.

³Chamberlain and Imbens (2003, p. 13) discuss how this approach can be extended to the overidentified case in which the number of moment conditions exceeds the dimension of β .

elements equal to the given argument. β has the *weighted least squares (WLS)* representation^{4,5}

$$\beta = \beta(\theta) = [X' \text{diag}\{\theta\} X]^{-1} X' \text{diag}\{\theta\} y. \quad (9)$$

Note that the “parameter” β depends on both $z^{(1)} = [y^{(1)}, X^{(1)}]$ and $z^{(2)} = [y^{(2)}, X^{(2)}]$, as well as θ .

The prior and posterior distributions for θ in Section 2 induce prior and posterior distributions for β via (9), albeit not particularly tractable, because β is a nonlinear function of the θ . The limiting (as $\underline{y} \rightarrow 0_j$) “noninformative” improper prior together with the standard case ($\tilde{J} = 0$ and $z^{(2)}$ null) have received most of the attention. This is unfortunate. Because the number of hyperparameters in $\underline{y}$ is J , which is greater than or equal to the sample size n , an *informative* prior is warranted. Unfortunately, this necessitates eliciting a large number of prior hyperparameters.

This paper introduces an informative prior for θ with two properties: all elements in $\underline{y}$ are positive implying prior density (2) is proper, and unobserved support points (fictitious observations) $z^{(2)} = [y^{(2)}, X^{(2)}]'$ (together with $\underline{y}^{(2)}$) are introduced to reflect prior beliefs about β . The resulting combination is defined to be an *informed Bayesian bootstrap*.

The multinomial probabilities θ are not the parameters of interest. The functional $\beta(\theta)$ is the focus of attention. Choosing $\underline{y}$, so that the implied prior for β is sensible, is a challenging task. An obvious reference case is the *symmetric prior* where all elements in $\underline{y}$ are the same. Then θ_j ($j = 1, 2, \dots, J$) are identically distributed, and the prior and posterior for β are centered over $b = (X'X)^{-1}X'y$.

For example, consider the standard case and *symmetric prior* in which θ_j ($j = 1, 2, \dots, J$) are identically distributed with

$$\underline{y}_j = c \quad (j = 1, 2, \dots, J), \quad c > 0. \quad (10)$$

Then $E(\theta_j) = J^{-1}$ ($j = 1, 2, \dots, J$), regardless of c , and $\text{Var}(\theta_j) = (J - 1)[J^2(Jc + 1)]^{-1}$ which is decreasing in c . As $c \rightarrow \infty$, the elements of θ_j ($j = 1, 2, \dots, J$) become degenerate random variables equal to J^{-1} .

⁴Since the denominator in (6) is common across j , Lancaster (2003) notes that β also has the WLS representation $\beta = \beta(g) = [X' \text{diag}\{g\} X]^{-1} X' \text{diag}\{g\} y$. This leads to a slightly different approximation than discussed in the text.

⁵Gilstein and Learner (1983) characterize the set of WLS estimates as the union of the bounded convex polytopes formed by the hyperplanes on which residuals are zero.

Because $\bar{\mathbf{v}} = (\mathbf{c} + 1)\mathbf{1}_J$, $\mathbf{1}_J' \bar{\mathbf{v}} = J(\mathbf{c} + 1)$, the posterior mean of θ is the same as the prior mean J^{-1} , however, the posterior variance $\text{Var}(\theta_j) = (J - 1)[J^2(J\{c+1\} + 1)]^{-1}$ is smaller than the prior variance. In other words, the hyperparameter c only affects the prior and posterior dispersions. Setting $c = .005$ corresponds to a fairly diffuse proper prior, whereas, the uniform case $c = 1$ corresponds to a fairly tight prior. Jeffreys' prior corresponds to $c = .50$.

Note the standard case assumption in the preceding example is important for the posterior analysis. Suppose only m of the possible J support points are observed. Let $\mathbf{n}^{(1)} = [\mathbf{n}_1, \mathbf{n}_2, \dots, \mathbf{n}_m]'$. Then $\bar{\mathbf{v}} = [\bar{\mathbf{v}}^{(1)'} , \bar{\mathbf{v}}^{(2)'}]'$, where $\bar{\mathbf{v}}^{(1)} = \mathbf{y}^{(1)} + \mathbf{n}^{(1)}$ and $\bar{\mathbf{v}}^{(2)} = \mathbf{y}^{(2)}$. In this case, the posterior mean of θ consists of two different subvectors $E(\theta^{(1)}|\mathbf{z}) = \frac{\bar{\mathbf{v}}^{(1)}}{\mathbf{1}_m' \bar{\mathbf{v}}^{(1)} + \mathbf{1}_J' \mathbf{y}^{(2)}} = \frac{c\mathbf{1}_m + \mathbf{n}^{(1)}}{cJ + \mathbf{n}}$ and $E(\theta^{(2)}|\mathbf{z}) = \frac{\bar{\mathbf{v}}^{(2)}}{\mathbf{1}_m' \bar{\mathbf{v}}^{(1)} + \mathbf{1}_J' \mathbf{y}^{(2)}} = \frac{c\mathbf{1}_J}{cJ + \mathbf{n}}$. In other words, the symmetry of the prior does **not** carry over to the posterior. Heterogeneous priors are introduced in Section 6.

4. HC Estimation

Lancaster (2003) showed the equivalence to order J^{-1} of the covariance matrix of the bootstrap distribution of the OLS and the Eicker/White/Huber heteroskedasticity robust covariance matrix estimate. Furthermore, he showed that the covariance matrices of the Efron (1979) bootstrap and the Bayesian bootstrap of Rubin (1981) are identical to $O(J^{-1})$.

Lancaster uses a limiting (as $\mathbf{y} \rightarrow \mathbf{0}_N$) “noninformative” improper prior in the standard case ($\tilde{\mathbf{J}} = 0$ and $\mathbf{z}^{(2)}$ is null). Extending Lancaster (2003) to the **general case** in which $\tilde{\mathbf{J}} > 0$ and $\mathbf{z}^{(2)}$ is non-null, this paper approximates the posterior distribution of $\beta = \beta(\theta)$ in (9) by taking a Taylor series expansion around the posterior mean $\bar{\theta} \equiv E(\theta|\mathbf{z}) = \bar{\mathbf{v}}/(\mathbf{1}_J' \bar{\mathbf{v}})$, where $\mathbf{1}_J$ is a $J \times 1$ vector of ones. This leads to the approximation

$$\begin{aligned}
 \beta^*(\theta) &= \beta(\bar{\theta}) + \left[\frac{\partial \beta(\theta)}{\partial \theta'} \right]_{\theta=\bar{\theta}} (\theta - \bar{\theta}) \\
 &= \beta(\bar{\theta}) + [\mathbf{y} - \mathbf{X}\beta(\bar{\theta})]' \otimes (\mathbf{X}' \text{diag}\{\bar{\theta}\} \mathbf{X})^{-1} \mathbf{X}' \left[\frac{\partial \text{vec}(\text{diag}\{\theta\})}{\partial \theta'} \right]_{\theta=\bar{\theta}} (\theta - \bar{\theta}) \\
 &= \beta(\bar{\theta}) + [\mathbf{u}(\bar{\theta})' \otimes \mathbf{Q}(\bar{\theta})'] \begin{bmatrix} \mathbf{e}_1 \mathbf{e}_1' \\ \vdots \\ \mathbf{e}_J \mathbf{e}_J' \end{bmatrix} (\theta - \bar{\theta}) \\
 &= \beta(\bar{\theta}) + \mathbf{R}(\bar{\theta})' (\theta - \bar{\theta}),
 \end{aligned} \tag{11}$$

where $\mathbf{e}_j$ is a $J \times 1$ vector with the j^{th} element equal to unity and all other elements equal to zero, $\mathbf{Q}(\bar{\boldsymbol{\theta}}) = \mathbf{X}(\mathbf{X}' \text{diag}\{\bar{\boldsymbol{\theta}}\} \mathbf{X})^{-1}$, $\mathbf{u}(\bar{\boldsymbol{\theta}}) = \mathbf{y} - \mathbf{X}\boldsymbol{\beta}(\bar{\boldsymbol{\theta}})$, and

$$\mathbf{R}(\bar{\boldsymbol{\theta}}) = [\mathbf{e}_1 \mathbf{e}_1', \dots, \mathbf{e}_J \mathbf{e}_J'] [\mathbf{u}(\bar{\boldsymbol{\theta}}) \otimes \mathbf{Q}(\bar{\boldsymbol{\theta}})] = \text{diag}\{\mathbf{u}(\bar{\boldsymbol{\theta}})\} \mathbf{Q}(\bar{\boldsymbol{\theta}}). \quad (12)$$

The posterior $\boldsymbol{\theta} | \mathbf{z} \sim \mathbf{D}_J(\bar{\mathbf{v}})$ implies that the exact posterior mean and variance of (11) are

$$\mathbf{E}[\boldsymbol{\beta}^*(\boldsymbol{\theta}) | \mathbf{z}] = \boldsymbol{\beta}(\bar{\boldsymbol{\theta}}), \quad (13)$$

$$\begin{aligned} \mathbf{V}^{\text{HC}} &= \text{Var}[\boldsymbol{\beta}^*(\boldsymbol{\theta}) | \mathbf{z}] \\ &= \mathbf{R}(\bar{\boldsymbol{\theta}})' \text{Var}(\boldsymbol{\theta} | \mathbf{z}) \mathbf{R}(\bar{\boldsymbol{\theta}}) \\ &= (\mathbf{I}_J' \bar{\mathbf{v}} + 1)^{-1} \mathbf{Q}(\bar{\boldsymbol{\theta}})' \text{diag}\{\mathbf{u}(\bar{\boldsymbol{\theta}})\} [\text{diag}\{\bar{\boldsymbol{\theta}}\} - \bar{\boldsymbol{\theta}} \bar{\boldsymbol{\theta}}'] \text{diag}\{\mathbf{u}(\bar{\boldsymbol{\theta}})\} \mathbf{Q}(\bar{\boldsymbol{\theta}}). \end{aligned} \quad (14a)$$

Since the first-order condition for WLS implies $(\mathbf{X}' \text{diag}\{\bar{\boldsymbol{\theta}}\} \mathbf{X})^{-1} \mathbf{X}' \text{diag}\{\bar{\boldsymbol{\theta}}\} \mathbf{u}(\bar{\boldsymbol{\theta}}) = \mathbf{0}_K$, and because

$$\begin{aligned} \text{diag}\{\mathbf{u}(\bar{\boldsymbol{\theta}})\} [\text{diag}\{\bar{\boldsymbol{\theta}}\} - \bar{\boldsymbol{\theta}} \bar{\boldsymbol{\theta}}'] \text{diag}\{\mathbf{u}(\bar{\boldsymbol{\theta}})\} &= \text{diag}\left\{[\bar{\theta}_1 [\mathbf{u}(\bar{\boldsymbol{\theta}})]^2, \dots, \bar{\theta}_J [\mathbf{u}(\bar{\boldsymbol{\theta}})]^2]'\right\} \\ &\quad - \text{diag}\{\bar{\boldsymbol{\theta}}\} \mathbf{u}(\bar{\boldsymbol{\theta}}) \mathbf{u}(\bar{\boldsymbol{\theta}})' \text{diag}\{\bar{\boldsymbol{\theta}}\}. \end{aligned} \quad (15)$$

it follows that (14a) simplifies to

$$\begin{aligned} \mathbf{V}^{\text{HC}} &= (\mathbf{I}_J' \bar{\mathbf{v}} + 1)^{-1} \mathbf{Q}(\bar{\boldsymbol{\theta}})' \text{diag}\left\{[\bar{\theta}_1 [\mathbf{u}(\bar{\boldsymbol{\theta}})]^2, \dots, \bar{\theta}_J [\mathbf{u}(\bar{\boldsymbol{\theta}})]^2]'\right\} \mathbf{Q}(\bar{\boldsymbol{\theta}}) \\ &= (\mathbf{I}_J' \bar{\mathbf{v}} + 1)^{-1} (\mathbf{X}' \text{diag}\{\bar{\boldsymbol{\theta}}\} \mathbf{X})^{-1} [\mathbf{X}' \text{diag}\{\bar{\theta}_1 [\mathbf{u}(\bar{\boldsymbol{\theta}})]^2, \dots, \bar{\theta}_J [\mathbf{u}(\bar{\boldsymbol{\theta}})]^2\} \mathbf{X}] (\mathbf{X}' \text{diag}\{\bar{\boldsymbol{\theta}}\} \mathbf{X})^{-1} \quad (14b) \\ &= \frac{\mathbf{I}_J' \bar{\mathbf{v}}}{(\mathbf{I}_J' \bar{\mathbf{v}} + 1)} (\mathbf{X}' \text{diag}\{\bar{\mathbf{v}}\} \mathbf{X})^{-1} [\mathbf{X}' \mathbf{U}^{\text{HC}} \mathbf{X}] (\mathbf{X}' \text{diag}\{\bar{\mathbf{v}}\} \mathbf{X})^{-1}, \end{aligned}$$

where $\mathbf{U}^{\text{HC}} \equiv \mathbf{U}^{\text{HC}}(\bar{\mathbf{v}}) \equiv \text{diag}\{\bar{v}_1 [\mathbf{u}_1(\bar{\mathbf{v}})]^2, \dots, \bar{v}_J [\mathbf{u}_J(\bar{\mathbf{v}})]^2\}$. To the extent that $\boldsymbol{\beta}^*(\boldsymbol{\theta})$ in (11) is a good approximation to $\boldsymbol{\beta}(\boldsymbol{\theta})$ in (9), (13) and (14b) should be good approximations to the posterior mean and variance of $\boldsymbol{\beta}$.

To add some transparency to (14b), consider symmetric prior (10) and the standard case. Then, the posterior mean is $\bar{\boldsymbol{\theta}} \equiv \mathbf{E}(\boldsymbol{\theta} | \mathbf{z}) = \mathbf{J}^{-1} \mathbf{I}_J$, $\boldsymbol{\beta}(\bar{\boldsymbol{\theta}}) = \mathbf{b}$, and (14b) simplifies to

$$\mathbf{V}^{\text{HC}} = \left(\frac{J}{J(c+1)+1} \right) \mathbf{V}^{\text{HC0}}, \quad (16)$$

where $\mathbf{V}^{\text{HC0}}$ is the White/Eicker robust covariance matrix

$$\mathbf{V}^{\text{HC0}} = (\mathbf{X}' \mathbf{X})^{-1} \mathbf{X}' \mathbf{U}^{\text{HC0}} \mathbf{X} (\mathbf{X}' \mathbf{X})^{-1}, \quad (17)$$

$U^{HC0} \equiv \text{diag}\{(u_1^{\text{OLS}})^2, \dots, (u_J^{\text{OLS}})^2\}$ and $u^{\text{OLS}} = y - Xb$. As $c \rightarrow 0$ (Lancaster's case), $V^{HC} \rightarrow [J/(J+1)]V^{HC0}$. Therefore, the noninformative symmetric prior provides a Bayesian justification for using the OLS estimate b together with V^{HC0} in approximate Bayesian inference for β .

5. Choice of $z^{(2)} = [y^{(2)}, X^{(2)}]$

$z^{(j)} = [y^{(j)}, X^{(j)}]$ ($j = 1, 2$) play similar roles in defining $\beta = \beta(\theta)$ in (9) and transforming prior beliefs about θ into prior beliefs about β . Approximate the **prior** distribution of $\beta(\theta)$ by taking a Taylor series expansion around the prior mean $\underline{\theta} \equiv E(\theta) = \underline{v}/(\underline{1}_J' \underline{v})$. Then analogous to (13), the prior mean of β can be approximated by

$$\begin{aligned} \beta(\underline{\theta}) &= \beta(\underline{v}) = [X' \text{diag}\{\underline{\theta}\} X]^{-1} X' \text{diag}\{\underline{\theta}\} y \\ &= [X' \text{diag}\{\underline{v}\} X]^{-1} X' \text{diag}\{\underline{v}\} y \\ &= [X^{(1)'} \text{diag}\{\underline{v}^{(1)}\} X^{(1)} + X^{(2)'} \text{diag}\{\underline{v}^{(2)}\} X^{(2)}]^{-1} [X^{(1)'} \text{diag}\{\underline{v}^{(1)}\} y^{(1)} + X^{(2)'} \text{diag}\{\underline{v}^{(2)}\} y^{(2)}]. \end{aligned} \quad (18)$$

One goal in choosing $z^{(2)} = [y^{(2)}, X^{(2)}]$ may be to locate (18) at an *a priori* likely value. For example, suppose $\tilde{J} = n$ and once again consider symmetric prior (10). Then all the diagonal matrices in (18) cancel. Often an interesting prior for public consumption centers all the slopes in (18) over zero. Choose the first column of $X^{(2)}$ equal to $-\underline{1}_J$, the remaining columns of $X^{(2)}$ equal to the corresponding columns of $X^{(1)}$, and $y^{(2)} = -y^{(1)}$. Then $\beta(\underline{\theta}) = [\bar{y}, 0_{K-1}']'$.

Next consider approximate posterior mean (13). Because

$$X' \text{diag}\{\bar{v}\} X = X' \text{diag}\{\underline{v}\} X + X^{(1)'} \text{diag}\{n^{(1)}\} X^{(1)}, \quad (19)$$

$$X' \text{diag}\{\bar{v}\} y = X' \text{diag}\{\underline{v}\} y + X^{(1)'} \text{diag}\{n^{(1)}\} y^{(1)}, \quad (20)$$

(13) has the representation

$$\begin{aligned} \beta(\bar{\theta}) &= [X' \text{diag}\{\bar{\theta}\} X]^{-1} X' \text{diag}\{\bar{\theta}\} y \\ &= [X' \text{diag}\{\bar{v}\} X]^{-1} X' \text{diag}\{\bar{v}\} y \\ &= [X' \text{diag}\{\underline{v}\} X + X^{(1)'} \text{diag}\{n^{(1)}\} X^{(1)}]^{-1} [X' \text{diag}\{\underline{v}\} y + X^{(1)'} \text{diag}\{n^{(1)}\} y^{(1)}] \\ &= [X' \text{diag}\{\underline{v}\} X + X^{(1)'} \text{diag}\{n^{(1)}\} X^{(1)}]^{-1} \\ &\quad \left[(X' \text{diag}\{\underline{v}\} X) \beta(\underline{\theta}) + (X^{(1)'} \text{diag}\{n^{(1)}\} X^{(1)}) \beta^{(1)}(n^{(1)}) \right], \end{aligned} \quad (21)$$

which is a matrix-weighted average of approximate prior mean (18) and the WLS estimate based on the observed $\mathbf{z}^{(1)} = [\mathbf{y}^{(1)}, \mathbf{X}^{(1)}]$:

$$\boldsymbol{\beta}^{(1)}(\mathbf{n}^{(1)}) = [\mathbf{X}^{(1)'} \mathbf{diag}\{\mathbf{n}^{(1)}\} \mathbf{X}^{(1)}]^{-1} \mathbf{X}^{(1)'} \mathbf{diag}\{\mathbf{n}^{(1)}\} \mathbf{y}^{(1)}. \quad (22)$$

These fictitious observations enter approximate prior mean (18) and posterior mean (21) through the $K(K+1)/2$ unique elements in $\mathbf{X}^{(2)'} \mathbf{diag}\{\mathbf{y}^{(2)}\} \mathbf{X}^{(2)}$, and the $K \times 1$ location vector $\mathbf{X}^{(2)'} \mathbf{diag}\{\mathbf{y}^{(2)}\} \mathbf{y}^{(2)}$ [or $\boldsymbol{\beta}^{(2)}(\mathbf{y}^{(2)})$]. The actual $\tilde{\mathbf{J}}$ ($K+1$) elements in $\mathbf{z}^{(2)}$ are *not* required to compute the approximation (18) or (21). This is no more onerous than the customary linear regression case.

6. Choice of $\mathbf{y}$

Symmetric prior (10) centers on OLS. When θ_j ($j = 1, 2, \dots, J$) are *not* identically distributed, then attention switches to WLS. But unfortunately, the not identically distributed case requires choice of J hyperparameters: $\mathbf{y}_j$ ($j = 1, 2, \dots, J$). This section considers extensions of White/Eicker/Huber, motivated by sampling properties of $\mathbf{V}^{\text{HC0}}$, that suggest choices for $\mathbf{y}_j$ ($j = 1, 2, \dots, J$).

The White/Eicker robust covariance matrix (17) has been criticized by numerous authors for its sampling distribution (e.g., it tends to underestimate variances and test statistics have the wrong size).⁶ Numerous suggestions have been made for adding weights to the diagonal elements of $\mathbf{U}^{\text{HC0}}$. These weights can be interpreted as a choice for $\bar{\mathbf{v}}_j$ ($j = 1, 2, \dots, J$) in $\mathbf{U}^{\text{HC}}(\bar{\mathbf{v}})$, from which values of prior hyperparameters $\mathbf{y}_j$ ($j = 1, 2, \dots, J$) can be backed out. The resulting $\mathbf{y}$ results in *proper* priors and a well-defined marginal mass function (3).

An early critic of (17), predating White (1980), was Hinkley (1977) who suggested using $\mathbf{b}$ and scaling-up $\mathbf{V}^{\text{HC0}}$ by $J/(J - K)$:

$$\mathbf{V}^{\text{HC1}} = [J/(J - K)] \mathbf{V}^{\text{HC0}}. \quad (23)$$

While this reminiscent of symmetric prior for the standard case, it is different in an important way: unlike (23), (16) scales $\mathbf{V}^{\text{HC0}}$ down, not up. Hence, a Bayesian justification of (23) is unclear.

Other modifications of $\mathbf{V}^{\text{HC0}}$ involve diagonal elements $\mathbf{h}_j$ ($j = 1, 2, \dots, J$) of the so-called “hat

⁶See Chesher and Jewitt(1987), Cribari-Neto et al. (2000), Cribari-Neto and Zarkos (1999, 2001), Godfrey (2006), and MacKinnon and White (1985).

matrix" $X(X'X)^{-1}X'$. Cribari-Neto and Zarkos (2001) showed that the presence of high leverage points, as measured by h_j ($j = 1, 2, \dots, J$), in the design matrix is more decisive for the finite-sample behavior of different HC robust covariance matrix estimators than the degree of heteroskedasticity.

Table 1 summarizes the effects of adding weights to the diagonal elements of U^{HC0} , using a common notation HC0-HC4 in the literature. Besides the original White/Eicker/Huber case (HC0) and Hinkley's (HC1), Table 1 contains the modifications by Horn, Horn, et al. (1975) (HC2), MacKinnon and White (1985) (HC3), and Cribari-Neto (2004) (HC4). In addition Table 1 contains three Bayesian analogs (denoted HC2a, HC3a, and HC4a) that give three different choices for $\mathbf{v}_j$ ($j = 1, 2, \dots, J$) that imply the same diagonal weights for U^{HC0} suggested by HC2-HC4. HC2a-HC4a, however, center inferences over $\beta(\bar{\mathbf{v}})$ (instead of $\mathbf{b}$) and use $\mathbf{u}(\bar{\boldsymbol{\theta}})$ instead of the OLS residuals.

7. Summary

The Bayesian bootstrap provides a Bayesian semiparametric analysis of the linear model not requiring a distributional assumption other than the multinomial sampling. Using the Bayesian bootstrap and a symmetric improper prior in the standard case, Lancaster (2003) showed the White/Eicker/Huber robust standard errors are reasonable asymptotic approximations to the posterior standard deviations around OLS. Using a proper (possibly heterogeneous) prior in the general case, this paper has extended Lancaster's results and linked the White/Eicker/Huber robust standard errors to posterior Bayesian analysis, but with two important differences: the posterior location for $\beta = \beta(\boldsymbol{\theta})$ is WLS not OLS, and the posterior covariance matrix (14b), while of a similar form to (16), is evaluated at WLS residuals instead of OLS residuals.⁷

Choice of $\mathbf{z}^{(2)} = [\mathbf{y}^{(2)}, \mathbf{X}^{(2)}]'$ and $\mathbf{v}$ is critical. Section 5 suggested choosing $\mathbf{z}^{(2)}$ to locate the prior over zero slopes. Section 6 drew on various modifications of the White/Eicker/Huber estimator to suggest choices for $\mathbf{v}$.

⁷Leamer (1991?) calls the treatment of heteroskedasticity that alters the size but not the location of confidence sets "**White-washing**." He argues an **increase** in the standard deviations after White-correction signals a potentially **large** change in the location of the confidence sets. On the other hand, if the White-correction leaves unchanged or **reduces** the standard errors, he argues the estimates are likely to be **insensitive** to reweighting.

References

- Chamberlain, G. and G. W. Imbens, 2003, "Nonparametric Applications of Bayesian Inference," *Journal of Business & Economic Statistics* 21, 12-18.
- Chesher, A. and I. Jewitt, 1987, "The Bias of a Heteroskedasticity Consistent Covariance Matrix Estimator," *Econometrica* 55, 1217-1222.
- Cribari-Neto, F., 2004, "Asymptotic Inference Under Heteroskedasticity of Unknown Form," *Computational Statistics & Data Analysis* 45, 215-233.
- Cribari-Neto, F., S. L. P. Ferrari, and G. M. Cordeiro, 2000, "Improved Heteroscedasticity-Consistent Covariance Matrix Estimators," *Biometrika* 87, 907-918.
- Cribari-Neto, F., S. G. Zarkos, 1999, "Bootstrap Methods for Heteroskedastic Regression Models: Evidence on Estimation and Testing," *Econometric Reviews* 18, 211-228.
- Cribari-Neto, F., S. G. Zarkos, 2001, "Heteroskedasticity-Consistent Covariance Matrix Estimation: White's Estimator and the Bootstrap," *Journal of Statistical Computing and Simulation* 68, 391-411.
- Cribari-Neto, F. and N. Galvão, 2003, "A Class of Improved Heteroskedasticity-Consistent Covariance Matrix Estimators," *Communications in Statistics: Theory and Methods* 32, 1951-1980.
- Efron, B., 1979, "Bootstrap Methods: Another Look at the Jackknife," *Annals of Statistics* 7, 1-26.
- Eicker, F., 1963, "Asymptotic Normality and Consistency of the Least Squares Estimators for Families of Linear Regressions," *Annals of Mathematical Statistics* 34, 447-456.

Eicker, F., 1967, "Limit Theorems for Regressions with Unequal and Dependent Errors," in *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability*, Vol. 1. Berkeley: University of California Press, 59-82.

Godfrey, L. G., 2006, "Tests for Regression Models with Heteroskedasticity of Unknown Form," *Computational Statistics & Data Analysis* 50, 2715-2733,

Hahn, J., 1977, "Bayesian Bootstrap of the Quantile Regression Estimator: A Large Sample Study," *International Economic Review* 38, 206–213.

Hinkley, D. V., 1977, "Jackknifing in Unbalanced Situations," *Technometrics* 19, 285-292.

Horn, S. D., Horn, R. A., Duncan, D. B., 1975, "Estimating Heteroskedastic Variances in Linear Models," *Journal of the American Statistical Association* 70, 380-385.

Huber P. J., 1967, "The Behavior of Maximum Likelihood Estimates Under Nonstandard Conditions," *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability*, Vol. I (Berkeley: University of California Press), 221–233.

Kauermann, G. and R. J. Carroll, 2001, "A Note on the Efficiency of Sandwich Covariance Matrix Estimation," *Journal of the American Statistical Association* 96, 1387-1396.

Lancaster, T., 2003, "A Note on Bootstraps and Robustness," unpublished manuscript, Brown University.

Lancaster, T., 2004, *An Introduction to Modern Bayesian Econometrics* (Oxford: Blackwell Publishing).

MacKinnon, J. and H. White (1985), "Some Heteroskedasticity-Consistent Covariance Matrix Estimators with Improved Finite Sample Properties," *Journal of Econometrics* 29, 305-325.

Ragusa, G., 2007, "Bayesian Likelihoods for Moment Condition Models," unpublished manuscript, University of California, Irvine.

Rubin, D., 1981, "The Bayesian Bootstrap," *Annals of Statistics* 9, 130-134.

White, H., 1980, "A Heteroskedasticity-Consistent Covariance Matrix Estimator and a Direct Test for Heteroskedasticity," *Econometrica* 48, 817-838.

Table 1: Choice of $\mathbf{v}$

	Authors	$\mathbf{v}_j$	$E[\boldsymbol{\beta}^*(\mathbf{g}) \mathbf{z}]$	$\mathbf{U}^{\text{HC}}$	$\mathbf{V}^{\text{HC}}$
HC0	White/Eicker/Huber	$\rightarrow 0$	b	$\text{diag}\{[\mathbf{u}_1^{\text{OLS}}]^2, \dots, [\mathbf{u}_J^{\text{OLS}}]^2\}$	$(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{U}^{\text{HC0}}\mathbf{X}(\mathbf{X}'\mathbf{X})^{-1}$
HC1	Hinkley		b	$\left(\frac{J}{J-K}\right)\text{diag}\{[\mathbf{u}_1^{\text{OLS}}]^2, \dots, [\mathbf{u}_J^{\text{OLS}}]^2\}$	$\left(\frac{J}{J-K}\right)\mathbf{V}^{\text{HC0}}$
HC2	Horn, et al.		b	$\text{diag}\left\{\frac{[\mathbf{u}_1^{\text{OLS}}]^2}{1-h_1}, \dots, \frac{[\mathbf{u}_J^{\text{OLS}}]^2}{1-h_J}\right\}$	$(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{U}^{\text{HC2}}\mathbf{X}(\mathbf{X}'\mathbf{X})^{-1}$
HC2a		$\frac{h_j}{1-h_j}$	$\beta(\bar{\mathbf{v}})$	$\text{diag}\left\{\frac{[\mathbf{u}_1(\bar{\mathbf{v}})]^2}{1-h_1}, \dots, \frac{[\mathbf{u}_J(\bar{\mathbf{v}})]^2}{1-h_J}\right\}$	$\kappa(\mathbf{X}'\text{diag}\{\bar{\mathbf{v}}\}\mathbf{X})^{-1}\mathbf{X}'\mathbf{U}^{\text{HC2a}}\mathbf{X}(\mathbf{X}'\text{diag}\{\bar{\mathbf{v}}\}\mathbf{X})^{-1}$
HC3	MacKinnon/White		b	$\text{diag}\left\{\frac{[\mathbf{u}_1^{\text{OLS}}]^2}{(1-h_1)^2}, \dots, \frac{[\mathbf{u}_J^{\text{OLS}}]^2}{(1-h_J)^2}\right\}$	$(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{U}^{\text{HC2}}\mathbf{X}(\mathbf{X}'\mathbf{X})^{-1}$
HC3a		$\frac{h_j(2-h_j)}{(1-h_j)^2}$	$\beta(\bar{\mathbf{v}})$	$\text{diag}\left\{\frac{[\mathbf{u}_1(\bar{\mathbf{v}})]^2}{(1-h_1)^2}, \dots, \frac{[\mathbf{u}_J(\bar{\mathbf{v}})]^2}{(1-h_J)^2}\right\}$	$\kappa(\mathbf{X}'\text{diag}\{\bar{\mathbf{v}}\}\mathbf{X})^{-1}\mathbf{X}'\mathbf{U}^{\text{HC3a}}\mathbf{X}(\mathbf{X}'\text{diag}\{\bar{\mathbf{v}}\}\mathbf{X})^{-1}$
HC4	Cribari-Neto		b	$\text{diag}\left\{\frac{[\mathbf{u}_1^{\text{OLS}}]^2}{(1-h_1)^{\delta_1}}, \dots, \frac{[\mathbf{u}_J^{\text{OLS}}]^2}{(1-h_J)^{\delta_J}}\right\}$	$(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{U}^{\text{HC3}}\mathbf{X}(\mathbf{X}'\mathbf{X})^{-1}$
HC4a		$\frac{1-(1-h_j)^{\delta_j}}{(1-h_j)^{\delta_j}}$	$\beta(\bar{\mathbf{v}})$	$\text{diag}\left\{\frac{[\mathbf{u}_1(\bar{\mathbf{v}})]^2}{(1-h_1)^{\delta_1}}, \dots, \frac{[\mathbf{u}_J(\bar{\mathbf{v}})]^2}{(1-h_J)^{\delta_J}}\right\}$	$\kappa(\mathbf{X}'\text{diag}\{\bar{\mathbf{v}}\}\mathbf{X})^{-1}\mathbf{X}'\mathbf{U}^{\text{HC4a}}\mathbf{X}(\mathbf{X}'\text{diag}\{\bar{\mathbf{v}}\}\mathbf{X})^{-1}$

Note: $\beta(\bar{\mathbf{v}}) = [\mathbf{X}'\text{diag}\{\bar{\mathbf{v}}\}\mathbf{X}]^{-1}\mathbf{X}'\text{diag}\{\bar{\mathbf{v}}\}\mathbf{y}$, $\bar{\mathbf{v}}_j = \mathbf{v}_j + \mathbf{n}_j$ ($j = 1, 2, \dots, m$), $\bar{\mathbf{v}}_j = \mathbf{v}_j$ ($j = m+1, m+2, \dots, J$), $\delta_j = \min\left\{4, \frac{h_j}{K/N}\right\}$ ($j = 1, 2, \dots, J$), and $\kappa = \mathbf{v}_J' \bar{\mathbf{v}} / (\mathbf{v}_J' \bar{\mathbf{v}} + 1)$.