

ESTIMATING STANDARD ERRORS FOR THE PARKS MODEL: CAN JACKKNIFING HELP?

by

W. Robert Reed and Rachel S. Webb*

Abstract

Non-spherical errors, namely heteroscedasticity, serial correlation and cross-sectional correlation are commonly present within panel data sets. These can cause significant problems for econometric analyses. The FGLS(Parks) estimator has been demonstrated to produce considerable efficiency gains in these settings. However, it suffers from underestimation of coefficient standard errors, oftentimes severe. Potentially, jackknifing the FGLS(Parks) estimator could allow one to maintain the efficiency advantages of FGLS(Parks) while producing more reliable estimates of coefficient standard errors. Accordingly, this study investigates the performance of the jackknife estimator of FGLS(Parks) using Monte Carlo experimentation. We find that jackknifing can -- in narrowly defined situations -- substantially improve the estimation of coefficient standard errors. However, its overall performance is not sufficient to make it a viable alternative to other panel data estimators.

JEL Categories: C23, C15

Keywords: Panel Data estimation, Parks model, cross-sectional correlation, jackknife, Monte Carlo.

Revised, November 23, 2010

*The contact author is Rachel S. Webb, Department of Economics & Finance, University of Canterbury, Private Bag 4800, Christchurch 8020, New Zealand; email: rachel309@gmail.com.

I. INTRODUCTION

Panel data commonly suffer from a variety of nonspherical error behaviours, including heteroscedasticity, serial correlation, and cross-sectional correlation. As is well known, the simultaneous occurrence of serial and cross-sectional correlation bedevils existing estimation procedures. The Parks model (Parks, 1967) remains the most commonly used estimation procedure for simultaneously handling cross-sectional and serial correlation. For example, the options available with the Stata command “xtgls” are all variations of the Parks model. Recent applications include Congleton and Bose (2010); Stallman and Deller (2010); Kebede, Kagochi, and Jolly (2010); and Roll, Schwartz, and Subrahmanyam (2009). A quick search of papers in Web of Science that cite Parks (1967) produces hundreds more. However, while FGLS(Parks) is consistent and asymptotically efficient, it can produce notoriously bad estimates of coefficient standard errors in finite samples.

The only other parametric estimator that simultaneously addresses both serial and cross-sectional correlation is Beck and Katz’s PCSE estimator (Beck and Katz, 1995). Beck and Katz (1995) propose a two-step estimator that they claim produces reliable standard error estimates at no cost to estimator efficiency when compared to FGLS(Parks). In a recent paper, Chen, Lin and Reed (2010) show that the latter claim does not generally hold. Specifically, the PCSE estimator compares poorly with FGLS(Parks) on efficiency grounds when data are characterized by both serial and cross-sectional correlation. There remains, therefore, a demand for an estimation procedure that produces both relatively efficient coefficient estimates and reliable standard errors.

This paper uses Monte Carlo experiments to study whether jackknifing the FGLS(Parks) estimator provides a solution to this problem. On the face of it, jackknifing would appear to be a promising avenue. As a result of increased computer processing speeds, jackknifing has become increasingly feasible (Breunig, 2002; Sunil, 2002). Further, it has been shown to reliably estimate coefficient standard errors in a variety of settings (Schucany, 1989; Jennrich, 2008). Potentially, jackknifing would allow one to maintain the efficiency advantages of FGLS(Parks) while producing more reliable estimates of coefficient standard errors.

While jackknifing with panel data characterized by both serial and cross-sectional correlation is not without its challenges (as we discuss below), it stands in contrast with bootstrapping. To date, no successful bootstrapping procedures have been developed for the Parks model. For example, block bootstrapping techniques have been developed for one-way clustering such as serial correlation or cross-sectional correlation (e.g., Cameron, Gelbach, and Miller, 2008). However, there are no block bootstrapping procedures that are valid for the simultaneous occurrence of both of these. One can resample “blocks” of observations, where the blocks are clusters based on groups or clusters based time, but one cannot do both. Relatedly, newly developed techniques exist for calculating robust standard errors with multi-way clustering such as both group and time (Cameron, Gelbach, and Miller, 2006), but these procedures do not allow cross-sectional and serial correlation to interact, as in the Parks model.¹ A further attraction of jackknifing is that it easily incorporates unbalanced panels.

¹ We explain this in further detail below.

Unfortunately, our Monte Carlo simulations find that while jackknifing can improve estimation of coefficient standard errors, its overall performance is not sufficient to make it a viable alternative to other panel data estimators.

II. THE PARKS ERROR STRUCTURE AND THE PROBLEM WITH ESTIMATING STANDARD ERRORS

The data generating process. This paper analyzes the following panel data problem. Let the DGP be represented as follows:

$$(1) \quad \mathbf{y} = [\mathbf{i} \quad \mathbf{x}] \begin{bmatrix} \beta_0 \\ \beta_x \end{bmatrix} + \boldsymbol{\varepsilon} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon},$$

where N and T are the number of cross-sectional units and time periods; β_0 and β_x are scalars; and $\mathbf{y}$, $\mathbf{i}$, $\mathbf{x}$, and $\boldsymbol{\varepsilon}$ are, respectively, $NT \times 1$ vectors of observations of the dependent variable, a constant term, observations of the exogenous explanatory variable, and unobserved errors, where $\boldsymbol{\varepsilon} \sim N(\mathbf{0}, \boldsymbol{\Omega}_{NT})$.

The $NT \times NT$ error variance-covariance matrix, $\boldsymbol{\Omega}_{NT}$, is structured according to the Parks model (Parks, 1967). It assumes (i) groupwise heteroscedasticity; (ii) first-order serial correlation; and (iii) time-invariant cross-sectional correlation.² This implies the following specification for $\boldsymbol{\Omega}_{NT}$:

$$(2) \quad \boldsymbol{\Omega}_{NT} = \boldsymbol{\Sigma} \otimes \boldsymbol{\Pi},$$

² In its most general form, the Parks model assumes groupwise, first-order serial correlation. In contrast, our experiments model the DGP with a common AR(1) parameter, ρ , that is the same across groups. We do this to facilitate comparison with previous Monte Carlo studies of this problem that have also assumed a common AR(1) parameter (cf., Chen, Lin, and Reed, 2010)

$$\text{where } \Sigma = \begin{bmatrix} \sigma_{11} & \sigma_{12} & \cdots & \sigma_{1N} \\ \sigma_{21} & \sigma_{22} & \cdots & \sigma_{2N} \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_{N1} & \sigma_{N2} & \cdots & \sigma_{NN} \end{bmatrix}, \mathbf{H} = \begin{bmatrix} 1 & \rho & \rho^2 & \cdots & \rho^{T-1} \\ \rho & 1 & \rho & \cdots & \rho^{T-2} \\ \rho^2 & \rho & 1 & \cdots & \rho^{T-3} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho^{T-1} & \rho^{T-2} & \rho^{T-3} & \cdots & 1 \end{bmatrix}^3.$$

The GLS estimators for β and $\text{var}(\hat{\beta})$ are given by the usual formulae:

$$\hat{\beta} = (X' \Omega_{NT}^{-1} X)^{-1} X' \Omega_{NT}^{-1} y \quad \text{and} \quad \text{Var}(\hat{\beta}) = (X' \Omega_{NT}^{-1} X)^{-1}.$$

In the case of Feasible Generalized Least Squares (FGLS), Ω_{NT} is replaced with $\hat{\Omega} = \hat{\Sigma} \otimes \hat{\Pi}$, so that

$$\text{Var}(\hat{\beta}_{FGLS}) = (X' \hat{\Omega}^{-1} X)^{-1}.$$

In other words, FGLS does not adjust coefficient standard errors for the additional uncertainty that arises from the fact that the elements of Ω_{NT} are unknown and must be estimated. This causes FGLS to underestimate coefficient standard errors. As there are a total of $\frac{N(N+1)}{2} + 1$ unique elements in Ω_{NT} , the degree of underestimation can be quite substantial.

III. JACKKNIFING THE FGLS(PARKS) ESTIMATOR

Let $\hat{\beta}$ be the FGLS(Parks) estimator given NT data points. Define $\hat{\beta}_i$ as the FGLS(Parks) estimate derived from dropping the i^{th} observation,

$$\hat{\beta}_i = (X' \hat{\Omega}_{NT-i}^{-1} X)^{-1} X' \hat{\Omega}_{NT-i}^{-1} y,$$

where X and y are the data observations corresponding to the $NT-1$ observations, and $\hat{\Omega}_{NT-i}$ is the estimate of the corresponding error variance-covariance matrix.

³ Note that cross-sectional and serial correlation “interact” in the error variance-covariance matrix of Equation (2). This is evidenced by the fact that all the elements in the $T \times T$, off-diagonal blocks are nonzero in the presence of serial correlation. In contrast, with two-way clustering of group and time effects, only the main diagonal of the off-diagonal blocks are nonzero.

The i th “pseudo-value” is defined by $\hat{\beta}_i^* = (NT)\hat{\beta} - (NT-1)\hat{\beta}_i$. The jackknife estimate of β is given by $\hat{\beta}^* = \frac{1}{NT} \sum_{i=1}^{NT} \hat{\beta}_i^*$, and the corresponding standard error for each of the elements of $\hat{\beta}^*$ is given by $s.e.(\hat{\beta}^*) = \sqrt{\frac{\sum_{i=1}^{NT} (\hat{\beta}_i^* - \hat{\beta}^*)^2}{NT(NT-1)}}$.

A complication arises when constructing $\hat{\Omega}_{NT-1}$. Not only must the values of ρ and the $\sigma_{\varepsilon,ij}$ s be re-estimated with the deletion of an observation, but $\hat{\Omega}$ now has dimensions $(NT-1) \times (NT-1)$. Let the deleted observation be indexed by it . For the t^{th} group, Π must be modified to account for the deleted t^{th} observation. To illustrate, if $T=5$ and $t=3$, Π_i becomes

$$\Pi_i = \begin{vmatrix} 1 & \rho & \rho^3 & \rho^4 \\ \rho & 1 & \rho^2 & \rho^3 \\ \rho^3 & \rho^2 & 1 & \rho \\ \rho^4 & \rho^3 & \rho & 1 \end{vmatrix}.$$

IV. DESCRIPTION OF THE MONTE CARLO EXPERIMENTS

As noted above, there are $\frac{N(N+1)}{2}$ unique σ_{ij} elements, and one unique value of ρ in Ω_{NT} . The Monte Carlo experiments require that population values be set for each of these parameters. In addition, a distribution must be determined for the explanatory variable, $\mathbf{x}$. An innovation of our study is that we set these parameters to match that of actual panel data.

Our artificial statistical environments consist of four families of data sets: (i) annual, U.S. state data and the level of real Per Capita Personal Income (PCPI); (ii) annual, U.S. state data and the growth of real PCPI; (iii) annual, international data and the level of real per capita Gross Domestic Product (GDP); and (iv) annual, international data and the growth of real per capita GDP.

We suppose that a researcher is interested in identifying the relationship between either the level or growth of the respective PCPI/GDP variables and a single explanatory variable. For the U.S. state income data, we use “tax burden” for the explanatory variable.⁴ Tax burden is defined as the ratio of state and local taxes over personal income and is commonly used as a measure of state tax rates (Helms, 1985; Wasylenko, 1997). The explanatory variable for the international GDP data is “government expenditure share,” measured by the share of government expenditures over GDP (Mankiw, 1995; Fölster and Henrekson, 2004).⁵

Each of the data set families consists of various-sized (balanced) data sets characterized by the number of cross-sectional units (N) and time periods (T). The idea is to set the underlying Monte Carlo parameter values so that the resulting, simulated data sets “look like” the kind of panel data that a researcher would encounter while estimating the relationship, say, between taxes and U.S. state PCPI levels, or between government expenditures and national GDP growth.

For example, to create an artificial statistical environment that is patterned after real data on U.S. income (PCPI) growth and taxes (tax burden) and, we start with 40 years of PCPI and tax burden data on 48 states (omitting Alaska and Hawaii), covering

⁴ PCPI data come from the Bureau of Economic Analysis. Tax data comes from the U.S. Census.

⁵ Real per capita GDP and government consumption data are taken from the Penn World tables, TABLE 6.1.

the period 1960-1999. A long time series is crucial for our approach because we want to have multiple observations for each element of the error covariance matrix. Most studies use time series where T is between 10 and 25 years. By having a data series substantially longer than that, we can sample multiple T -year, TSCS data sets in order to construct a “representative” error structure for a T -year, TSCS data set.

The first step consists of determining “representative” values for ρ and the σ_{ij} ’s. We begin by creating a sample using the first N states in our data set.⁶ Next, we choose the T -year period, 1960 to $(1960+T-1)$. We then estimate a fixed effects regression model for this sample, relating the dependent variable Y (= U.S. state PCPI) to a set of state fixed effects (D^j) and the explanatory variable X (= tax burden).

$$(3) \quad Y_{it} = \sum_{j=1}^N \alpha_j D_{it}^j + \alpha_{N+1} X_{it} + \text{error term}_{it},$$

where $i=1,2, \dots, N$; $t=1960, 1961, \dots, 1960+T-1$; and D^j is a state dummy variable that takes the value 1 for state j . Equation (3) is the basic “residual generating function.”

The residuals from this estimated equation are used to estimate ρ and the $\sigma_{u,ij}$ ’s in the usual manner, where the $\sigma_{u,ij}$ ’s are the covariances associated with the error term in the AR(1) equation, $\varepsilon_{it} = \rho \varepsilon_{i,t-1} + u_{it}$. Denote the associated estimates from this sample as

$$\hat{\rho}_i \text{ and } \hat{\Phi}_i = \begin{bmatrix} \hat{\sigma}_{u,11} & \hat{\sigma}_{u,12} & \cdots & \hat{\sigma}_{u,1N} \\ \hat{\sigma}_{u,21} & \hat{\sigma}_{u,22} & \cdots & \hat{\sigma}_{u,2N} \\ \vdots & \vdots & \ddots & \vdots \\ \hat{\sigma}_{u,N1} & \hat{\sigma}_{u,N2} & \cdots & \hat{\sigma}_{u,NN} \end{bmatrix}.$$

⁶ For example, since our data are organized alphabetically, the first five states would be Alabama, Arizona, Arkansas, California, and Colorado.

We repeat this process for every possible, T -contiguous year sample contained within the 40 years of data from 1960-1999 [i.e., 1960-(1960+ T -1), 1961-(1961+ T -1), 1962-(1962+ T -1), ..., (1999- T +1)-1999]. This produces a total of $40-T+1$ estimates of ρ and Φ , one for each T -contiguous year sample. We then average these to obtain “grand means” $\bar{\rho}$ and $\bar{\Phi}$. Our “representative” $NT \times NT$ error structure, Ω_{NT} , is then constructed as follows:

$$(4) \quad \Omega_{NT} = \bar{\Sigma} \otimes \bar{\Pi},$$

where

$$(5) \quad \bar{\Sigma} = \frac{1}{(1-\bar{\rho}^2)} \bar{\Phi},$$

and

$$(6) \quad \bar{\Pi} = \begin{bmatrix} 1 & \bar{\rho} & \bar{\rho}^2 & \cdots & \bar{\rho}^{T-1} \\ \bar{\rho} & 1 & \bar{\rho} & \cdots & \bar{\rho}^{T-2} \\ \bar{\rho}^2 & \bar{\rho} & 1 & \cdots & \bar{\rho}^{T-3} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \bar{\rho}^{T-1} & \bar{\rho}^{T-2} & \bar{\rho}^{T-3} & \cdots & 1 \end{bmatrix}.$$

This becomes the population error covariance matrix used for the associated Monte Carlo experiment. Note that every element of Ω_{NT} is based on error variance-covariance matrices estimated from actual panel data. In this sense, Ω_{NT} can be said to be “representative” of the kinds of error structures one encounters in “real world” data. “Real world” values of $\mathbf{x}$ are constructed similarly. Without loss of generality, we set $\beta_0 = \beta_x = 0$.

Given values for β_0 , β_x , ρ , the σ_{ij} ’s, and the distribution of $\mathbf{x}$, experimental observations are generated in the usual manner. Define $\mathbf{u}$ as an $NT \times I$ vector of standard

normal random variables. Define $\mathbf{Q}$ such that $\mathbf{Q}'\mathbf{Q} = \mathbf{\Omega}_{NT}$. Error terms are created by $\boldsymbol{\varepsilon} = \mathbf{Q}'\mathbf{u}$. These simulated errors are added to the deterministic component, $\beta_0 + \beta_x x_i$, to calculate stochastic observations of y_i , where $y_i = \beta_0 + \beta_x x_i + \varepsilon_i$, $i=1,2,\dots,NT$. Given an experimental data set of NT observations of (y_i, x_i) , we estimate $\hat{\boldsymbol{\beta}}$. We then perform the jackknifing procedure described above.

This procedure can be modified in a straightforward manner to conduct Monte Carlo experiments for alternative N and T values. We also employ a two-way fixed effects “residual generating function” (see Equation 3), where time dummy variables are also included. In the same way, we create artificial data for the other three data set families.

V. RESULTS AND DISCUSSION

The focus of our study is the “coverage rates” produced by the FGLS(Parks) and jackknife estimators, where the respective coverage rates are defined as the percent of 95% confidence intervals that contain the true population value of β_x . Coverage rates should be close to 95%.

Our main findings are:

1. The jackknife estimator can produce substantial improvements in coverage rates over FGLS(Parks).
2. Coverage rates for the jackknife estimator are unsatisfactory, except when $N=T$, and then only for some types of data.

TABLE 1 demonstrates the improvement that can come from jackknifing FGLS(Parks) estimates.

The numbers in the table represent the difference in coverage rates between FGLS(Parks) and the jackknife estimator. For example, using a population error variance-covariance matrix patterned after International GDP data (Level, Specification 1) and data sets of size $N=5$ and $T=5$, we find that FGLS(Parks) and the jackknife estimator produce coverage rates of 45.4 and 84.5 percent, respectively. Thus, the jackknife estimator has coverage rates that are 39.1 percentage points higher than the FGLS(Parks) estimator. It is the latter number that is reported in the table.

In general, the performance advantage of the jackknife estimator diminishes, and is sometimes reversed, as $\frac{T}{N}$ increases. This is primarily due to the better performance of FGLS(Parks). The last row of TABLE 1 averages the difference in coverage rates for values of N and T across the different population data sets. This generally confirms the observation that jackknifing results in greatest performance improvements when $N=T$.

To be a viable estimator, jackknifing should not only produce more reliable estimates of coefficient standard errors, but it should also have satisfactory coverage rates of its own. Unfortunately, TABLE 2 makes clear that this is not the case. Coverage rates are rarely close to 95 percent and are frequently less than 50 percent. When $N=T$, the jackknife estimator does slightly better. Overall, the coverage rates of the jackknife estimator compare poorly with alternative panel data estimators, such as the PCSE estimator (Beck and Katz, 1995).⁷

One disadvantage of our experimental methodology is that we do not directly control the values of cross-sectional and serial correlation. This is outweighed by the advantage of being able to measure estimator performance in simulated data

⁷ See Reed and Webb (2010) for coverage rates of the PCSE estimator using simulated data similar to that employed in this study.

environments patterned after the “real world.” The fact that the jackknife estimator performs poorly under these conditions eliminates it as a viable alternative to existing panel data estimators. Until a better approach is developed, the recommendation of Reed and Ye (2010) remains valid: Researchers should use FGLS(Parks) if the goal is estimator efficiency, and another estimator (e.g. the PCSE) if the concern is reliable hypothesis testing.

REFERENCES

- Beck, N., & Katz, J. N. (1995). What To Do (And What Not To Do) With Time-Series Cross-Section Data. *American Political Science Review* Vol. 89, No. 3, 634-647.
- Breunig, R. (2002). Bias Correction for Inequality Measures: An Application to China and Kenya. *Applied Economics Letters* Vol. 9, No. 12, 783-786.
- Cameron, A.C., Gelbach, J., and Miller, D.I. (2008). "Bootstrap-Based Improvements for Inference with Clustered Errors." *Review of Economics and Statistics* Vol. 90, No. 3, 414-427.
- Cameron, A.C., Gelbach, J., and Miller, D.I. (2006). "Robust Inference with Multi-Way Clustering." Cambridge, MA: NBER Technical Working Paper No. 327
- Chen, X., Lin, S., & Reed, W.R. (2010). A Monte Carlo Evaluation of the PCSE Estimator. *Applied Economics Letters* Vol. 17, No. 1, 7-10.
- Congleton, R.D., and Bose, F. (2010). "The Rise of The Modern Welfare State, Ideology, Institutions And Income Security: Analysis and Evidence." *Public Choice* Vol. 144, Nos. 3-4, 535-555.
- Fölster, S. and M. Henrekson. "Growth Effects of Government Expenditure and Taxation in Rich Countries." *European Economic Review* Vol. 45 (2001): 1501-1520.
- Helms, Jay L. "The Effect of State and Local Taxes on Economic Growth: A Time Series-Cross Section Approach." *The Review of Economics and Statistics* Vol. 67, No. 4 (1985): 574-582.
- Jennrich, R. I. (2008). Nonparametric Estimation of Standard Errors in Covariance Analysis Using the Infinitesimal Jackknife. *Psychometrika* Vol. 73, No. 4, 579-594.
- Kebede, E., Kagochi J., and Jolly C.M. (2010). "Energy Consumption And Economic Development In Sub-Sahara Africa." *Energy Economics* Vol. 32, No. 3, 532-537.
- Mankiw, N. G. "The Growth of Nations." *Brookings Papers on Economic Activity* Issue 1 (1995): 373-431.
- Parks, R. (1967). Efficient Estimation of a System of Regression Equations When Disturbances Are Both Serially and Contemporaneously Correlated. *Journal of the American Statistical Association* Vol. 62, 500-509.

- Reed, W.R. and Ye, H. (2010). Which Panel Data Estimator Should I Use? *Applied Economics*, forthcoming.
- Reed, W.R. and Webb, R. (2010). "The PCSE Estimator is Good – Just Not as Good As You Think." *Journal of Time Series Econometrics* Vol. 2, No. 1, Article 8.
- Roll R., Schwartz, E., and Subrahmanyam, A. (2009). "Options Trading Activity And Firm Valuation." *Journal of Financial Economics* Vol. 94, No. 3, 345-360.
- Schucany, W. & Sheather, S.J. (1989). Jackknifing R-estimators. *Biometrika* Vol. 76, No. 2, 393-398.
- Stallmann, J.I., and Deller, S. (2010). "Impacts Of Local And State Tax And Expenditure Limits On Economic Growth." *Applied Economics Letters* Vol. 17, No. 7, 645-648.
- Sunil. S. (2002). A Jackknife Maximum Likelihood Estimator for the Probit Model. *Applied Economics Letters* Vol. 9, No. 2, 73-74.
- Wasylenko, Michael. "Taxation and Economic Development: The State of the Economic Literature." *New England Economic Review* (March/April 1997): 37-52.

TABLE 1
Difference in Coverage Rates for FGLS (Parks) and Jackknife Estimators

<i>Spec.^a</i>	<i>Experimental Data Patterned After...^a</i>	<i>N=5</i>					<i>N=10</i>				<i>N=20</i>	
		<i>T=5</i>	<i>T=10</i>	<i>T=15</i>	<i>T=20</i>	<i>T=25</i>	<i>T=10</i>	<i>T=15</i>	<i>T=20</i>	<i>T=25</i>	<i>T=20</i>	<i>T=25</i>
<i>1</i>	<i>International GDP Data (Level)</i>	39.1	-8.6	-34.7	-45.6	-54.1	60.2	25.6	-3.4	-29.6	57	48.7
<i>1</i>	<i>International GDP Data (Growth)</i>	32.6	-18.5	-30.5	-42.3	-44.6	50.6	-9.1	-35.3	-44.2	70.6	42.6
<i>1</i>	<i>U.S. State PCPI Data (Level)</i>	42.1	22.7	3.7	5	-2.1	52.9	37.6	32.5	33.6	44.5	62.3
<i>1</i>	<i>U.S. State PCPI Data (Growth)</i>	39.7	9.4	-2.7	-14.8	-18.4	51.1	27.4	10.7	-9.5	53.2	57.7
<i>2</i>	<i>International GDP Data (Level)</i>	45.9	4.6	-23.4	7.1	-16.3	64.3	34.5	70.7	36.6	61.9	81.5
<i>2</i>	<i>International GDP Data (Growth)</i>	45.6	-13.2	-38.4	1.8	-67.4	58.9	20.6	61	49.5	69.5	84
<i>2</i>	<i>U.S. State PCPI Data (Level)</i>	39.1	-9.4	34.8	0.5	-25.5	69.5	45.1	82.4	13.6	64.1	88.5
<i>2</i>	<i>U.S. State PCPI Data (Growth)</i>	35.9	-21.3	5.5	-32.3	-21.9	65.6	24.9	54.4	17.8	67.9	89.7
<i>AVERAGE</i>		40	-4.3	-10.7	-15.1	-31.2	59.1	25.8	34.1	8.5	61.1	69.4

^a See text for an explanation of the two specifications of the “residual generating function” and the methodology used to produce simulated data sets patterned after the respective data.

TABLE 2
Coverage Rates for FGLS (Parks) with Jackknifed Standard Errors

<i>RGF</i>	<i>Model Data</i>	<i>N=5</i>					<i>N=10</i>				<i>N=20</i>	
		<i>T=5</i>	<i>T=10</i>	<i>T=15</i>	<i>T=20</i>	<i>T=25</i>	<i>T=10</i>	<i>T=15</i>	<i>T=20</i>	<i>T=25</i>	<i>T=20</i>	<i>T=25</i>
<i>1</i>	<i>International GDP Data (Level)</i>	84.5	57.6	35.7	28.7	21.4	81.3	71.4	50.9	33.4	66	73
<i>1</i>	<i>International GDP Data (Growth)</i>	81.7	59.7	52.1	44.4	43.3	83.7	53.8	37.9	34.5	85.5	79
<i>1</i>	<i>U.S. State PCPI Data (Level)</i>	89.7	87	74.3	76.5	69.6	82.3	91	89.5	87.5	53.1	74.5
<i>1</i>	<i>U.S. State PCPI Data (Growth)</i>	89.5	83.5	79	70.5	67.4	86.3	90.4	82.9	67.5	70.2	93.2
<i>2</i>	<i>International GDP Data (Level)</i>	52.4	42.5	38.4	80	62.4	64.4	40.6	81.7	53.2	61.9	81.5
<i>2</i>	<i>International GDP Data (Growth)</i>	54.1	34.7	28.1	84.5	21	59.2	27.4	79.4	79.3	69.5	84
<i>2</i>	<i>U.S. State PCPI Data (Level)</i>	45.8	24.9	87.6	64.9	47.1	69.7	49.7	89.9	34.2	64.1	88.5
<i>2</i>	<i>U.S. State PCPI Data (Growth)</i>	44.1	18.8	70.7	45	66.6	65.9	31.4	70.6	47.6	67.9	89.7