

BORROWING INFORMATION OVER TIME IN BINOMIAL/LOGIT NORMAL MODELS FOR SMALL AREA ESTIMATION

Carolina Franco¹, William R. Bell²

ABSTRACT

Linear area level models for small area estimation, such as the Fay-Herriot model, face challenges when applied to discrete survey data. Such data commonly arise as direct survey estimates of the number of persons possessing some characteristic, such as the number of persons in poverty. For such applications, we examine a binomial/logit normal (BLN) model that assumes a binomial distribution for rescaled survey estimates and a normal distribution with a linear regression mean function for logits of the true proportions. Effective sample sizes are defined so variances given the true proportions equal corresponding sampling variances of the direct survey estimates. We extend the BLN model to bivariate and time series (first order autoregressive) versions to permit borrowing information from past survey estimates, then apply these models to data used by the U.S. Census Bureau's Small Area Income and Poverty Estimates (SAIPE) program to predict county poverty for school-age children. We compare prediction results from the alternative models to see how much the bivariate and time series models reduce prediction error variances from those of the univariate BLN model. Standard conditional variance calculations for corresponding linear Gaussian models that suggest how much variance reduction will be achieved from borrowing information over time with linear models agree generally with the BLN empirical results.

Key words: area level model, complex surveys, American Community Survey, bivariate model, SAIPE.

1. Introduction

Small area estimation by area level models often uses linear Gaussian mixed models, specifically the model of Fay and Herriot (1979). When such models are applied to data from a repeated survey the question arises as to whether better results may be obtained by borrowing information from past data. Time series extensions to the Fay-Herriot (FH) model have thus been explored. See, e.g., Ghosh et al. (1996), Datta, Lahiri, Maiti, and Lu (1999), Saei and Chambers (2003), Rao

¹U.S. Census Bureau. E-mail: Carolina.Franco@census.gov.

²U.S. Census Bureau. E-mail: William.R.Bell@census.gov.

and Molina (2015, Sections 4.4.3, 8.3, and 10.9), Esteban et al. (2012), and Pratesi et al. (2010, Chapter 3). Huang and Bell (2012) investigated the use of a bivariate FH model that, for each area, borrowed information from an estimate obtained by pooling recent past survey samples, which is similar to borrowing information from an average of past survey estimates.

Area level modeling has also been extended through the use of Generalized Linear Mixed Models (GLMM), which have been discussed in the context of small area estimation by Ghosh, et al. (1998) and Rao and Molina (2015, Section 10.13). GLMMs have potential advantages for modeling inherently discrete data arising from direct survey estimates of the number of persons that possess a certain characteristic (e.g., the number of persons in poverty). This can also be thought of as modeling survey estimates of the corresponding proportions (e.g., poverty rates). Directly applying a linear Gaussian model to such data may risk producing nonsensical results such as negative predictions or, more likely, prediction intervals that include negative values. Taking logarithms can eliminate these problems but creates the problem of dealing with direct estimates of zero that arise when no one in an area's sample possesses the characteristic whose prevalence is being estimated. Analogous problems arise if predicted proportions or their interval limits exceed one, or if direct estimates of proportions equal one. GLMMs avoid such problems and may also help account for the skewness typically inherent in such data when the underlying proportion is near zero or one.

This paper focuses on small area models that combine both extensions just mentioned. To address the challenges posed by discrete survey data, we use a binomial/logit normal (BLN) model. This particular GLMM assumes a binomial distribution for discrete observations, and a normal distribution with a linear regression mean function for logits of the binomial proportions. We determine effective sample sizes for the binomial distributions to preserve sampling variances estimated via a generalized variance function. To borrow information from past data we extend the BLN model to a bivariate version and then to a time series version. The latter uses a first order autoregressive model (AR(1)), although other time series structures could be used. The normality assumption for the random effects in the logits of the proportions facilitates these extensions for modeling dependence. One qualification to note is that the extensions assume independence of the sampling errors of the survey estimates for all years covered by the time series model, as well as for the two equations of the bivariate model.

Our motivating application comes from the U.S. Census Bureau's Small Area Income and Poverty Estimates (SAIPE) program. SAIPE provides annual poverty estimates for various age groups for states, counties, and school districts of the U.S.

An important SAIPE product is school district age 5–17 poverty estimates used by the U.S. Department of Education in allocating federal funds (over \$14 billion in 2013) to school districts. For more information on the SAIPE program, see Bell et al. (2015) or the SAIPE web page at <http://www.census.gov/did/www/saipe/>.

The survey data source used by SAIPE, which we also use here for illustration, is poverty estimates from the American Community Survey (ACS). The ACS is the largest household sample survey in the United States, sampling approximately 3.5 million addresses per year. It collects data on a broad range of population characteristics such as income, health insurance coverage, and education, and publishes estimates annually. For areas with populations of 65,000 or more, ACS publishes estimates based on a single year of data collection. For the smallest places, published estimates use data pooled from five years of ACS samples. The ACS, with its 5-year estimates, has effectively replaced the decennial census long form sample, which was last carried out in Census 2000. SAIPE poverty models use ACS 1-year estimates, which are not publicly released for counties with populations less than 65,000.

We focus here on modeling county poverty for school-aged (5–17) children, a key component of developing the SAIPE poverty estimates for school districts. The SAIPE production model has as a covariate the log of the Census 2000 long form county estimates of age 5–17 children in poverty. This covariate is going further and further out of date, motivating consideration of replacing it with past, but more recent, ACS data. Huang and Bell (2012) thus explored bivariate FH models for current ACS 1-year and past ACS 5-year poverty estimates. This issue also motivates the bivariate and time series extensions to the BLN model that we study here.

Our interest in studying the BLN model applied to SAIPE data stems from its potential advantages for modeling discrete data discussed above, a relevant consideration for the ACS 1-year estimates for small counties. Slud (2000, 2004) did several analyses comparing results from GLMM models to results from models similar to the SAIPE county production model. Slud (2000) showed advantages to the use of a unit level BLN model of sampled counts compared to a linear Fay–Herriot model for logged counts when the data were simulated from the GLMM model.

The rest of the paper proceeds as follows. Section 2 presents the BLN model and its extensions to bivariate and time series (AR(1)) versions. Section 3 presents results from application of these models to ACS county poverty data for 2012. We compare results between the variants of the BLN model to illustrate the potential benefits of the two different ways of borrowing information from past data. Section 4 provides conclusions.

2. Binomial/Logit Normal (BLN) Models

The BLN model may be written as

$$y_i | p_i, n_i \sim \text{Bin}(n_i, p_i) \quad i = 1, \dots, m \quad (1)$$

$$\text{logit}(p_i) = \mathbf{x}_i' \boldsymbol{\beta} + u_i \quad (2)$$

where $\text{logit}(p_i) = \log[p_i/(1 - p_i)]$, $u_i \sim i.i.d. N(0, \sigma_u^2)$, and n_i is the sample size for area i . The model as given by (1)–(2) can be readily applied to unweighted sample counts y_i , but doing this ignores any complex aspects of the survey design. For applications to complex survey data where the y_i are survey weighted estimates, two problems arise. First, the possible values for the y_i will not be the integers $0, 1, \dots, n_i$ for any direct definition of sample size n_i . Instead, y_i will take a value from a finite set of unequally-spaced numbers (not necessarily integers) determined by the survey weights that apply to the sample cases in area i . Second, the sampling variance of y_i implied by the binomial distribution in (1), $n_i p_i (1 - p_i)$, will be incorrect.

To address these problems we start by defining an “effective sample size” $\tilde{n}_i$, and an “effective sample number of successes” $\tilde{y}_i$, determined to maintain: (i) the direct survey weighted estimate $\tilde{p}_i$ of the true proportion, and (ii) a corresponding sampling variance estimate, $\widehat{\text{var}}(\tilde{p}_i)$. For the latter we set

$$\tilde{n}_i = \check{p}_i (1 - \check{p}_i) / \widehat{\text{var}}(\tilde{p}_i) \quad (3)$$

where $\check{p}_i$ is a preliminary model-based prediction of the population proportion p_i (on which $\text{var}(\tilde{p}_i)$ truly depends), and $\widehat{\text{var}}(\tilde{p}_i)$ depends on $\check{p}_i$ through a fitted generalized variance function (GVF). Franco and Bell (2013) give a detailed explanation of the implementation of this GVF for application of the BLN models to the ACS county poverty data used in SAIPE models. Liu, Lahiri, and Kalton (2007) and You (2008) used essentially this type of sampling variance model, but applied it in models of survey estimates of proportions assumed to follow either a normal or a Beta distribution.

Having thus determined $\tilde{n}_i$, we set $\tilde{y}_i = \tilde{n}_i \times \tilde{p}_i$ and, after rounding, substitute $(\tilde{n}_i, \tilde{y}_i)$ for (n_i, y_i) in (1). Note that $\tilde{y}_i = 0$ if $\tilde{p}_i = 0$, but this does not cause problems since the BLN allows for observations of zero. Moreover, $\check{p}_i > 0$ in (3) implies $\tilde{n}_i > 0$ even if $\tilde{p}_i = 0$. Rounding of $\tilde{n}_i$ and $\tilde{y}_i$ may be required by computer software for the fitting of models such as (1)–(2).

We extend the univariate BLN given by (1)–(2) to a bivariate BLN, written as

$$\tilde{y}_{1i}|p_{1i}, \tilde{n}_{1i} \sim \text{Bin}(\tilde{n}_{1i}, p_{1i}) \quad \tilde{y}_{2i}|p_{2i}, \tilde{n}_{2i} \sim \text{Bin}(\tilde{n}_{2i}, p_{2i}) \quad (4)$$

$$\text{logit}(p_{1i}) = \mathbf{x}'_{1i}\beta_1 + u_{1i} \quad \text{logit}(p_{2i}) = \mathbf{x}'_{2i}\beta_2 + u_{2i} \quad (5)$$

$$\begin{bmatrix} u_{1i} \\ u_{2i} \end{bmatrix} \sim i.i.d. N(0, \Sigma), \quad \Sigma = \begin{bmatrix} \sigma_{11} & \sigma_{12} \\ \sigma_{12} & \sigma_{22} \end{bmatrix}$$

for $i = 1, \dots, m$. In (4), for each area i $\tilde{n}_{1i}$ and $\tilde{y}_{1i}$ are the effective sample size and effective number of successes derived as discussed above from a direct survey estimate y_{1i} and a corresponding sampling variance estimate. Similarly, $\tilde{n}_{2i}$ and $\tilde{y}_{2i}$ are derived from another direct survey estimate y_{2i} and corresponding sampling variance estimate. The bivariate BLN model can be applied to estimates y_{1i} and y_{2i} from two different surveys or for two different time points from the same repeated survey. Notice, though, that $\tilde{y}_{1i}$ and $\tilde{y}_{2i}$ are assumed conditionally independent (given $p_{1i}, \tilde{n}_{1i}$ and $p_{2i}, \tilde{n}_{2i}$), as will be the case if the samples on which they are based are drawn independently. This is true for our application of the bivariate BLN in Section 3, where y_{1i} and y_{2i} are ACS 1-year and previous 5-year poverty estimates, respectively, since ACS samples are drawn approximately independently each year. (The ACS housing unit samples are drawn independently each year from one of five population subframes to which U.S. residential addresses are randomly assigned, with rotation of the subframes on a five-year cycle. Sampling fractions for most areas are 5% or less. See U.S. Census Bureau (2014, pp. 32–46).)

Instead of summarizing the information in five prior years of ACS data through the resulting 5-year estimates, a logical alternative to consider is to use the corresponding five individual 1-year estimates. Putting this together with the current 1-year estimates, implies modeling six years of ACS 1-year estimates. We do this by extending the BLN to assume the model errors u_{it} have an AR(1) correlation structure:

$$\tilde{y}_{it}|p_{it}, \tilde{n}_{it} \sim \text{Bin}(\tilde{n}_{it}, p_{it}) \quad i = 1, \dots, m, \quad t = 1, \dots, T \quad (6)$$

$$\text{logit}(p_{it}) = \mathbf{x}'_{it}\beta_t + u_{it} = \mathbf{x}'_{it}\beta_t + \sigma_t \tilde{u}_{it} \quad (7)$$

$$\tilde{u}_{it} = \phi \tilde{u}_{i,t-1} + \varepsilon_{it} \quad (8)$$

where $-1 < \phi < 1$. The ε_{it} are assumed distributed as *i.i.d.* $N(0, 1 - \phi^2)$ so that $\text{var}(\tilde{u}_{it}) = 1$ (Box and Jenkins 1970, p. 58) and $\text{var}(u_{it}) = \sigma_t^2$. Note that this version of the BLN-AR(1) model has different regression coefficients (β_t) and different model variances (σ_t^2) each year. We have three reasons for making this assumption. First, the true regression coefficients and model variances may actually differ year-

to-year. Second, this assumption is implicitly made in current SAIPE production by fitting the univariate production models separately for each year. Third, and most importantly here, the assumption facilitates comparisons of results, especially the comparisons of posterior variances and standard deviations that we make in Section 3, to corresponding results obtained from the univariate and bivariate BLN models. Both the univariate and bivariate BLN models use regression coefficients and a model variance specific to the prediction year.

A more conventional version of the BLN-AR(1) model would set $\beta_t = \beta$ and $\sigma_t^2 = \sigma_u^2$ for all years t in the model. With this assumption, the covariance matrix of $\mathbf{u}_i = (u_{i1}, \dots, u_{iT})'$ has the general form (Box and Jenkins 1970, pp. 56-58)

$$\text{var}(\mathbf{u}_i) = \sigma_u^2 \begin{bmatrix} 1 & \phi & \dots & \phi^{T-1} \\ \phi & 1 & \dots & \phi^{T-2} \\ \vdots & \vdots & \ddots & \vdots \\ \phi^{T-1} & \phi^{T-2} & \dots & 1 \end{bmatrix}. \quad (9)$$

For the heteroscedastic version given by (7)–(8), we drop σ_u^2 in (9) and pre- and post-multiply the matrix there by a diagonal matrix with diagonal elements σ_t . Linear Gaussian models with AR(1) model errors for ACS poverty data were investigated by Taciak and Basel (2012) for application to logs of ACS county 5-17 poverty estimates, and by Hawala and Lahiri (2012) for application to ACS estimates of county 5-17 poverty rates. Esteban et al. (2012) applied such models to data from the Spanish Living Conditions Survey to improve direct survey estimates of the male and female poverty rates for Spanish provinces.

3. Application: Borrowing Information from Past Data in Small Area Estimation of Poverty for U.S. Counties

To illustrate the potential for variance reductions from the bivariate and AR(1) extensions to the BLN models, we apply these models to estimating poverty rates for school aged children in U.S. counties in 2012. The univariate BLN (1)–(2) models the 2012 ACS 1-year county poverty estimates, the bivariate BLN (4)–(5) models these estimates together with the 2007–2011 ACS 5-year county poverty estimates, and the BLN-AR(1) (6)–(8) models the ACS 1-year county poverty estimates from 2007–2012. We shall compare prediction results from these models for 2012 for 3,136 counties, omitting 6 counties from the SAIPE universe which were not consistently defined across all 6 years of data. We did the same analysis with data corresponding to prediction years 2010 and 2011 and obtained very similar results.

The regression variables used in each of the models included 1 for an intercept

term, and logistic transformations of the following:

- the proportion of child tax exemptions “in poverty” for the county, i.e., the ratio of the number of child exemptions claimed on tax returns whose adjusted gross income falls below the poverty threshold divided by the total number of child exemptions for the county. (Notes: (i) In general terms, a “child tax exemption” is a child listed on an income tax return who is economically dependent on the person filing the return. (ii) The poverty threshold used is that applicable to a family of the size implied by the number of exemptions (persons) listed on the tax return.)
- an adjusted version of the county “child tax filer rate,” which is defined as the number of child exemptions in the county claimed on tax returns divided by the county population age 0–17.
- the “SNAP participation ratio,” defined as the ratio of county recipients of benefits from the Supplemental Nutrition Assistance Program (SNAP), a program that subsidizes food expenses of low income persons, in July of the previous year to the county population of the previous year.

Huang and Bell (2012) used the above ratio variables in bivariate models for ACS poverty rates, while Bell et al. (2007) used their logarithms in models for logs of ACS poverty rates. For $\mathbf{x}_{2i}$ in equation (5) of the bivariate BLN, we used the above variables defined for the middle year (2009) of the 5-year interval.

An issue arises for the child tax filer rate in that it often exceeds 1 due to the number of child tax exemptions in a county exceeding the county’s age 0–17 population. This occurs because the upper age limit for a child tax exemption can exceed 17, ranging as high as 23 for university students, and with no age limit for disabled children. The issue was addressed by multiplying all child tax filer rates by a constant factor to bring the maximum rescaled filer rate just below 1, permitting the logistic transformation. This adjustment is discussed further in Franco and Bell (2013).

We used the JAGS software (Plummer 2010) to implement the three models via a Bayesian approach with noninformative priors. Regression parameters were given normal priors with large variances, while the random effect variances in our models were given flat priors on intervals $[0, \kappa]$ chosen wide enough to contain essentially all the posterior probability as judged from examination of their posterior densities for a univariate model. The parameters ρ of the bivariate BLN and ϕ of the BLN-AR(1) models were given flat priors on $(-1, 1)$. We determined the effective sample sizes $\tilde{n}_i$ and effective numbers of successes $\tilde{y}_i$ for the BLN models

as discussed in Section 2 for both the ACS 1-year and ACS 5-year poverty estimates. We used separately fitted GVFs for the sampling variances for each year of the 1-year estimates, as well as a separately fitted GVF for the variances of the 5-year estimates.

3.1. Variance reductions from the extensions of the BLN model

Figure 1 compares the posterior means and standard deviations obtained from JAGS for the rates of school-aged children in poverty for U.S. counties in 2012 from the univariate, bivariate, and AR(1) BLN models. Parts (a) and (b) show that the posterior means are similar regardless of which of the models we choose. Figure 1(c) shows the posterior standard deviations tend to be lower for the bivariate BLN model than for the univariate BLN model, suggesting some value to incorporating the ACS 5-year estimates into the model. The gains are modest, however. The average percentage reduction in posterior standard deviations from using the bivariate versus the univariate model is approximately 5%, with about an 11% corresponding average reduction in posterior variances. The AR(1) model, on the other hand, yields only a 2.3% average decrease in standard deviations and a 4.6% decrease in variances compared to the univariate model. On average, it has larger posterior standard deviations than does the bivariate model, as reflected in Figure 1(d).

As the returns from using the bivariate or AR(1) BLN models to borrow information from past data are so modest, the question arises as to whether the data provide much evidence of dependence over time in the model errors u_{it} . In fact, the posterior mean of ρ from the bivariate BLN is .51 with a 95% posterior (credible) interval of (.43, .60), while the posterior mean of ϕ from the AR(1) model is 0.44, with a 95% interval of (.39, .50). So the data provide clear evidence of dependence over time in the u_{it} , but modeling this dependence does not produce much reduction in prediction uncertainty for the county 5–17 poverty rates.

3.2. How much improvement should we expect from borrowing information from past data?

As a rough guide to how much improvement might be expected from the bivariate or AR(1) models over the univariate model, we consider the linear FH model case when the true dependence structure is a stationary AR(1) model and all model parameters are known. We also assume for simplicity that the model error variance σ_u^2 and the sampling variances v_i remain constant over time. For this case it is straightforward to compute and compare the posterior variances (prediction MSEs) for the univariate, bivariate, and AR(1) versions of the FH model using standard results on

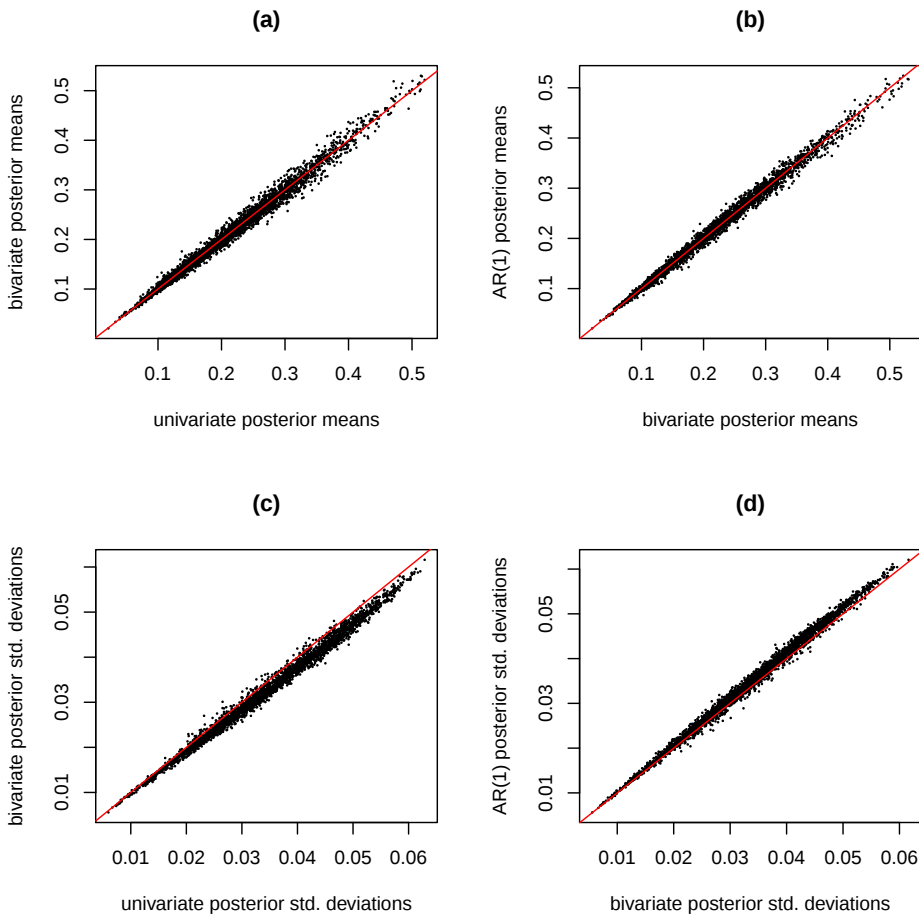


Figure 1: Comparison of posterior means and standard deviations for 2012 U.S. county poverty rates of school-aged children for univariate, bivariate, and AR(1) BLN Models.

conditional variances in a multivariate normal distribution – see the Appendix. Note that, since model parameters are assumed known, the predictions for each model are optimal conditional on the data used, but the data conditioned on differs across the three models.

Percent reductions in posterior variances for the bivariate and AR(1) models compared to the univariate model depend only on the parameter ϕ and variance ratio v_i/σ_u^2 . Figure 2(a) shows contour plots of the percent variance reductions achieved by the AR(1) model as functions of ϕ and v_i/σ_u^2 . (The plot assumes $\phi \geq .4$; a mirror image results for $\phi \leq -.4$, and percent reductions are small for $|\phi| < .4$.) It shows

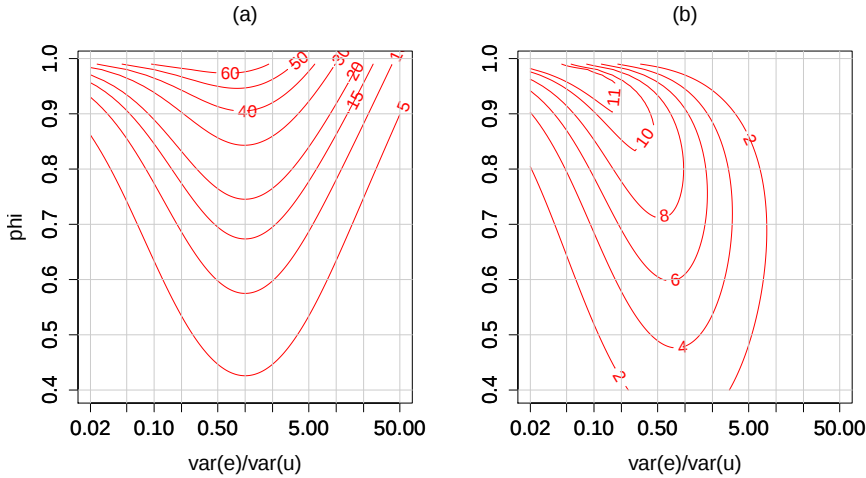


Figure 2: Contour plots of posterior variance percent reductions for small area estimates achieved by the FH-AR(1) model using 6 years of data compared to the univariate and bivariate FH models, when the true population characteristics actually follow an AR(1) model. Contours are shown as functions of the AR(1) parameter ϕ and $\text{var}(e_i)/\text{var}(u_i)$, the ratio of the sampling error variance to the model error variance. (a) Reductions from the AR(1) versus the univariate model. (b) Reductions from the AR(1) versus the bivariate model.

that the variance reductions increase with increasing values of ϕ , and decrease as the value of v_i/σ_u^2 deviates from 1.0. (Note that the x-axis in Figure 2(a) is on a log scale.) For values of ϕ such as .50 or less, the variance reductions are small, no more than about 7% when $v_i/\sigma_u^2 = 1$, and less as v_i/σ_u^2 moves away from 1.0. Large variance reductions require larger values of ϕ . For example, to achieve a 20% or greater reduction in variance requires $\phi \geq .75$.

Esteban et al. (2012) reported results related to those of Figure 2(a) obtained from a simulation study of the FH-AR(1) model, though augmented with a time-invariant area level random effect. This feature, and some other differences (most notably that their simulations provide estimates of the full prediction MSEs, not just a first order approximation) make their specific numerical results not directly comparable to ours. However, their results obtained with the alternative values of $\phi = 0, .25, .5, .75$ (denoted as ρ in their paper) are consistent with the general conclusions we draw from Figure 2(a). First, they found that borrowing from past data yielded little if any benefit for $\phi \leq .5$. Then, for $\phi = .75$, their augmented FH-AR(1) model appears (judging from their Figure 4.1) to reduce prediction MSEs by about

10%, or in some cases slightly more, relative to those obtained by applying this model with ϕ fixed at 0. Their simulation model assumed different values of the sampling variances across areas and time points, resulting in values of their ratio of sampling to model variance ranging from roughly .325 to .8 overall. The values within this range that were covered by a given simulation experiment varied as this depended on the value used for ϕ in the experiment. Esteban et al. also remarked that other simulations they did without the time-invariant random effect led to the same basic conclusions.

Figure 2(b) shows contours of percentage variance reductions from using the AR(1) model versus the bivariate model when the latter is applied to current year survey estimates and the average of survey estimates over the previous five years. We take this use of the five-year average as an approximation to the use of ACS 5-year estimates. Since the calculations assume the AR(1) is the true model, the bivariate model must have higher posterior variances. However, the reduction in variance from using the AR(1) model is generally small – less than 10% except for a small region in the upper left corner of the plot for high values of ϕ and values of $v_i/\sigma_u^2 < .5$.

One might wonder whether larger variance improvements from the AR(1) or bivariate models might result if more years of data were used compared to the six years assumed for the plots of Figure 2. Doing the same contour plots for the cases of 10 years of data and 20 years of data produced little change in the plots, except for very large values of ϕ and within a limited range of large v_i/σ_u^2 values, where more substantial advantages to the AR(1) over the bivariate model were observed. Over almost all of the range of ϕ and v_i/σ_u^2 , using more years of past data appeared to make little difference.

The values of the variance ratios, v_i/σ_u^2 , across the areas $i = 1, \dots, m$ in the model will clearly affect how much variance improvement is achieved in specific areas. To gauge this effect for our application, we fitted a linear FH model to the ACS estimated county poverty rates, for which we had the sampling variances v_i from the GVF, and, using the posterior mean of σ_u^2 , we calculated the ratios v_i/σ_u^2 . Figure 3 shows a histogram of these variance ratios with the x-axis on a log scale. Most of the values lie between .1 and 10, though some extend beyond this. The variance ratios across the U.S. counties thus reflect much of the x-axis range of the contour plots in Figure 2.

Figures 2 and 3, the simulation results of Esteban et al. (2012), and the estimates of ϕ for the AR(1)-BLN model, suggest that for our application only small improvements in posterior prediction variances would be realized from the AR(1) or bivariate models compared to the univariate BLN. This is consistent with the

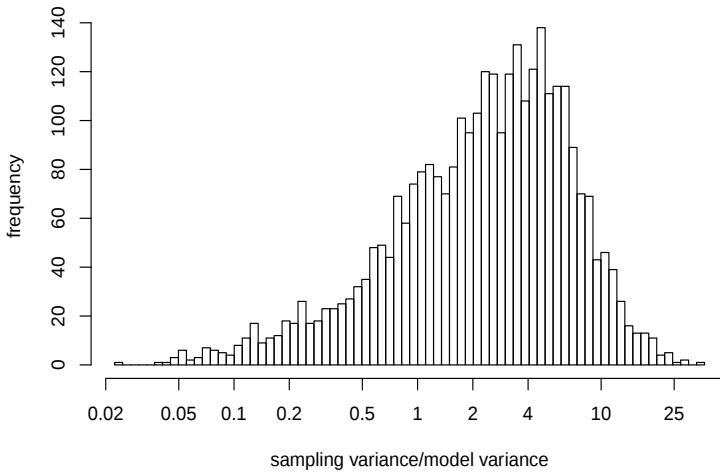


Figure 3: Histogram of the ratios of the sampling variances to the model variance in the FH model for the 2012 U.S. county poverty rates of school-aged children

posterior variance comparisons discussed in Section 3.1. Two other results from these comparisons may still seem surprising. First, the improvements for the bivariate BLN model are somewhat larger than the theoretical calculations for the linear model would suggest. Second, the improvements for the bivariate BLN model are larger than are those for the AR(1) model. While one would expect some limitations on how well calculations for linear FH models with parameters assumed known apply to fitted BLN models, that does not seem to explain these results since we obtained very similar results when we made the same comparisons using the bivariate and AR(1) extensions to the FH model applied to county poverty rates. In this case the bivariate FH model reduced prediction error standard deviations and variances compared to the univariate FH model by, on average, 5% and 9% (compared to 5% and 11% for the bivariate BLN). Corresponding figures for the AR(1) FH model were 2.7% and 5.2% (compared to 2.3% and 4.6% for the AR(1) BLN). In any case, differences between the comparisons for the BLN models and those for the FH model (both empirical and theoretical results) are not large, and all lead to the main conclusion that, given the value of ϕ for the AR(1) model, modest variance reductions would be achieved by the bivariate or AR(1) models relative to the univariate model.

3.3. Impact of removing model covariates

To illustrate a case where greater improvement would be expected from borrowing information over time, we repeated our empirical analyses after removing the regression covariates from the BLN models, leaving only the intercept terms. Without the regressors, the posterior means of ρ and ϕ skyrocketed to .92 and .94, respectively. We are now in the region of the parameter space where, by Figure 2(a), we would expect to see very substantial reductions in posterior variances from using a bivariate or AR(1) model rather than a univariate model. For this case, Figure 4 shows substantial differences between both the posterior means and posterior standard deviations of county poverty rates from the univariate and bivariate BLN models. In fact, we now see an average 25% reduction in posterior standard deviations and a 43% reduction in posterior variances from using the bivariate versus the univariate model. The AR(1) and bivariate BLN models performed similarly (results not shown on the plots), with the AR(1) yielding, on average, 1.3% higher posterior standard deviations compared to the bivariate BLN. The average reductions in standard deviations and variances for both the bivariate and AR(1) FH models for poverty rates were 26% and 45%.

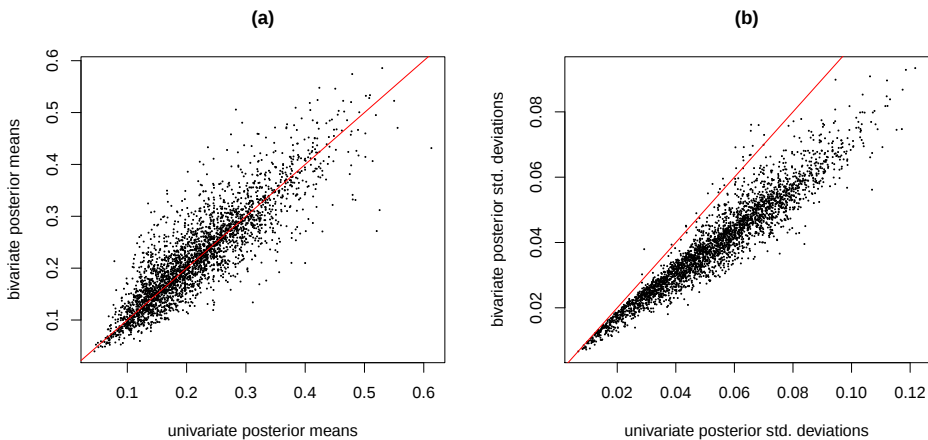


Figure 4: Comparisons of the posterior means and standard deviations for the 2012 U.S. county poverty rates of school-aged children for the univariate and bivariate BLN models with no regressors

3.4. Some model checks

For the linear (FH) model, where $y_i = (x_i'\beta + u_i) + e_i$, examination of standardized residuals defined as $(y_i - x_i'\hat{\beta})/[\widehat{\text{var}}(y_i - x_i'\hat{\beta})]^{1/2}$ provides a standard model check. We seek an analog for the BLN model (1)–(2). Since the inverse to (2) is $p_i = (1 + e^{-(x_i'\beta + u_i)})^{-1}$ and $E(u_i) = 0$, it may seem natural to use residuals defined as $y_i/n_i - \hat{p}_i$ where $\hat{p}_i = (1 + e^{-x_i'\hat{\beta}})^{-1}$ and $\hat{\beta}$ is an estimate of β . However, even with β known, $(1 + e^{-x_i'\beta})^{-1}$ is not an unbiased estimator of $E(y_i/n_i)$ due to the nonlinearity of the logistic transformation and the presence of the random effects u_i . Instead, we define residuals as $y_i/n_i - E(y_i)/n_i$ and compute

$$E(y_i/n_i) = (1/n_i)E_{p_i}[E(y_i|p_i)] = E_{p_i}(p_i) = \int_{-\infty}^{\infty} (1 + e^{-z_i})^{-1} f(z_i) dz_i \quad (10)$$

where $z_i = \text{logit}(p_i)$, $f(z_i)$ is the $N(x_i'\beta, \sigma_u^2)$ density, and $E_{p_i}(\bullet)$ denotes unconditional expectation over the distribution of p_i .

To standardize the residuals we need the unconditional variance

$$\begin{aligned} \text{var}(y_i) &= E_{p_i}[\text{var}(y_i|p_i)] + \text{var}_{p_i}[E(y_i|p_i)] \\ &= E_{p_i}[n_i p_i (1 - p_i)] + \text{var}_{p_i}[n_i p_i] \\ &= n_i E_{p_i}[p_i] - n_i E_{p_i}[p_i^2] + n_i^2 \text{var}_{p_i}[p_i]. \end{aligned} \quad (11)$$

To compute (11) requires computing $E_{p_i}[p_i^2]$ which, analogous to (10), is

$$E_{p_i}[p_i^2] = \int_{-\infty}^{\infty} (1 + e^{-z_i})^{-2} f(z_i) dz_i. \quad (12)$$

Substituting the posterior means of β and σ_u^2 into $f(z_i)$, both (10) and (12) can readily be computed by numerical integration. We used the “integrate” function in R (R Core Team 2013) for this purpose. We then computed standardized residuals as $[y_i/n_i - E(p_i)]/[\text{var}(y_i)]^{1/2}/n_i$.

Figure 5 plots such standardized residuals for 2012 from the equation for $\tilde{y}_{1i}$ of the bivariate BLN given by (4)–(5) against county effective sample sizes $\tilde{n}_{1i}$. (We could equally well do this for residuals from the equation for $\tilde{y}_{2i}$, but focus here on checking the model for $\tilde{y}_{1i}$ since our interest lies in predictions of p_{1i} .) For $\tilde{n}_{1i}$ “sufficiently large”, standard normal distribution inferences (e.g., ± 2.57 for a 99% confidence interval, as denoted by the blue dashed lines on the plot) may be appropriate given the approximate normal distribution of the binomial, although precisely how large $\tilde{n}_{1i}$ must be for this approximation to hold is unclear (Brown, Cai, and DasGupta 2001). In any case, in the plot the bulk of the residuals look reasonably symmetrical, with no systematic biases related to sample size (which is

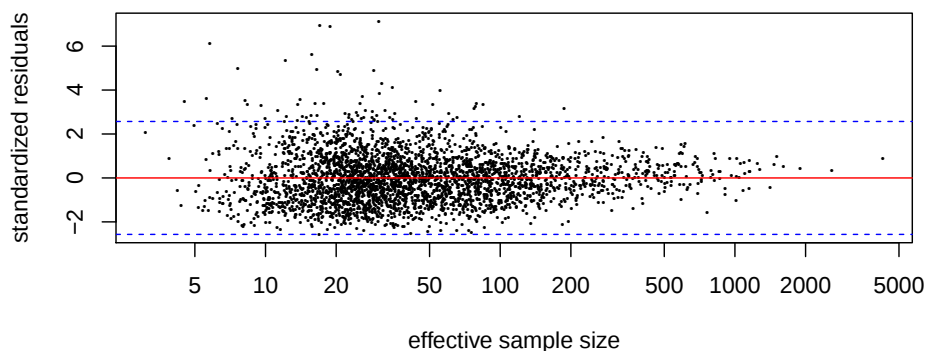


Figure 5: Standardized residuals from the 2012 bivariate BLN model's equation for y_{1i} plotted against county ACS effective sample sizes.

strongly related to population size). There are a number of large positive residuals, though mostly these occur at the smaller effective sample sizes, especially for $\tilde{n}_i$ of about 30 or less, where the direct estimates are erratic. It may seem odd that there is not a corresponding set of large in magnitude negative residuals. This is due to the fact that $\tilde{y}_i/\tilde{n}_i \in [0, 1]$ while all the predicted p_i values are less than 0.54. Extreme negative residuals are thus unlikely, while extreme positive residuals occur when $\tilde{y}_i/\tilde{n}_i$ is large, even 1.0, as happens sometimes with small samples.

We also examined a plot (not shown) of the standardized residuals against the predicted p_i values, which mimics a standard regression diagnostic (plot residuals against fitted values). This plot did not suggest any systematic biases related to the predicted county poverty rates.

Brown et al. (2001) suggest as a “calibration diagnostic” comparing model predictions aggregated to larger areas against corresponding direct survey estimates. In SAIPE production the county model predictions of the number of age 5–17 children in poverty are raked (rescaled) to force agreement with corresponding state estimates obtained from an FH model applied to direct ACS estimates of state poverty rates. For large states substantial weight is given to the direct ACS estimate in the model predictions, and this raking is then similar to raking to the direct estimates. In any case, there is practical interest in how much raking of the county model predictions is required. We examine this here for the bivariate BLN and (unraked) SAIPE production county model predictions derived from the 2012 ACS data.

To explain this in more detail, for the bivariate BLN model we expand our notation slightly to let $\hat{p}_{ji}$ be the bivariate BLN county model prediction of the age 5–17 poverty rate for county i in state j (treating the District of Columbia (DC)

as both a county and a state in this analysis), and N_{ji} be the 5–17 population estimate for county i obtained from the Census Bureau’s population estimates program. (Actually, slight modifications are made of the N_{ji} to estimate the county “poverty universes”, which exclude a relatively small set of persons for whom poverty status cannot be determined (Bell et al. 2015).) The predicted number of age 5–17 children in poverty for state j implied by its county model predictions is then $\sum_{i \in j} \hat{p}_{ji} N_{ji}$. The SAIPE production county model is an FH model for logarithms of the number of children age 5–17 in poverty (Bell et al. 2015). Predictions from this model are transformed to the original (unlogged) scale using a bias adjustment based on properties of the lognormal distribution and accounting for uncertainty due to estimating regression parameters of the model. These predictions are then simply summed across counties to yield state level predictions of the number in poverty.

Figure 6 plots percent differences of the state total estimates of the number of age 5–17 children in poverty from the two county models – bivariate BLN and SAIPE production – compared to the corresponding estimates derived from the SAIPE state model. The percent differences are defined as $100 \times (1 - \text{SAIPE state model estimate} / \text{aggregated county model predictions})$ so positive values indicate aggregated county model predictions exceeding the state model predictions and negative values indicate aggregated county model predictions lower than the state model predictions. The percent differences are plotted for 50 states, with states sorted by their ACS sample sizes (number of addresses). We dropped Alaska because it contained 5 of the 6 counties omitted from the modeling due to their not being consistently defined for all years of our data, which prevented us from getting an implied state poverty prediction for Alaska from the bivariate BLN model. The other omitted county was in Texas, but it had inconsequential effects on the state total.

Somewhat greater percent differences are to be expected at the left of Figure 6 for the small states where the estimation uncertainty is highest. This tendency is apparent in the plot. Apart from this, if we examine the blue solid dots in the plot, we see that the percent differences for the bivariate BLN model appear to be usually no more than a few percent. The corresponding percent differences for the SAIPE production estimates (red circles) appear to usually exceed those from the bivariate BLN, as well as being generally larger in magnitude. These impressions are reflected by Table 1, which summarizes the distributions of the percent differences.

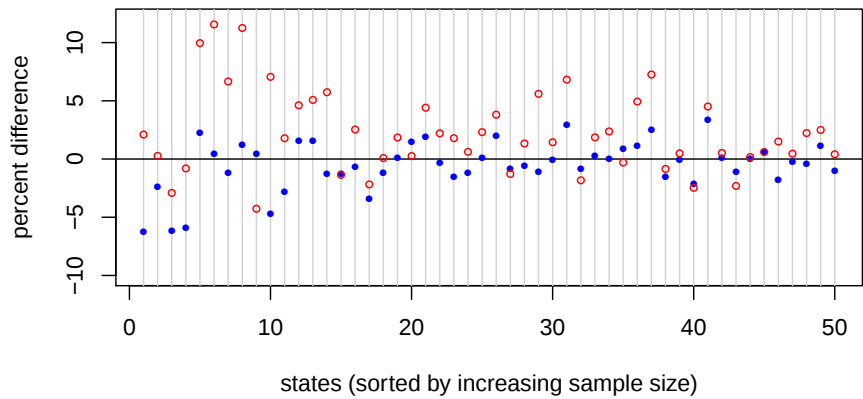


Figure 6: Percent differences between aggregated county model predictions of 2012 state total numbers of age 5–17 children in poverty and corresponding SAIPE state estimates. Red circles = SAIPE production model (unraked predictions); blue solid dots = bivariate BLN model.

Ideally we should take account of statistical uncertainty in the state level percent differences, but this is complicated, particularly for the bivariate BLN, by the dependence between the state and county model predictions due to both coming from models fitted to ACS data. As a conservative indication, 90% prediction intervals for the SAIPE state model predictions, expressed in multiplicative percentae terms, range from lows of around $\pm 1.7\%$ for the largest states (California and Texas) to highs of about $\pm 11\%$ to $\pm 13\%$ for some of the smallest (Wyoming, New Hampshire, and DC). These figures should overstate the uncertainty in the percent differences since we would expect positive dependence between the state and county model predictions.

county model	min	1 st quartile	median	mean	3 rd quartile	max
Bivariate BLN	−6.3	−1.2	−.5	−.3	.8	3.4
SAIPE production	−4.3	.2	2.2	1.8	4.5	11.6

Table 1: Distribution (omitting Alaska – see text) of percent differences between the state aggregates of county model predictions of 5–17 in poverty from the bivariate BLN and SAIPE production models compared to the SAIPE state model estimates for 2012.

4. Conclusions

Several conclusions stand out from the empirical and theoretical results presented in this paper. A general conclusion is that to achieve substantial variance reductions by jointly modeling current and past data requires fairly high levels of dependence over time in the random effects (model errors) of small area models. With modest levels of dependence, variance reductions from including past data are likely to be limited. A conclusion specific to the empirical example on modeling ACS poverty estimates is that the regression covariates used in the models do a good job explaining variation in poverty across counties and over time, leaving residuals with modest levels of dependence. Without these covariates in the models, the dependence over time in the model errors is strong, and borrowing information from past data then substantially reduces posterior (prediction error) variances.

A second general conclusion is that a bivariate model for the current year's estimate and the average of the estimates for some number of immediately preceding years may do about as well as an AR(1) model in borrowing information from past data for small area predictions. In fact, in the example bivariate models did slightly better than the corresponding AR(1) models. Additional comparisons could be made to models with more general dependence structures, such as a higher order AR model or a general 6×6 covariance matrix. While we intend to pursue this, we are confident that this will not alter the main conclusions expressed in the preceding paragraph. We also conjecture that bivariate models may do reasonably well in comparisons to other time series models with stationary autocorrelations, such as higher order AR models. It seems less clear whether this will be the case for models with nonstationary dependence, such as random walks. Consideration of the bivariate model is natural for the SAIPE application given that the ACS annually produces 5-year estimates for all U.S. counties and other small areas, and these 5-year pooled sample estimates can be thought of as similar to 5-year averages of 1-year estimates. While the bivariate model may seem less natural in other applications, it could be considered as a somewhat simpler alternative to using a time series model.

Appendix: Calculating Prediction MSEs for the Bivariate and FH-AR(1) Models

For extending the linear FH model to bivariate and AR(1) versions, let y_{it} be the direct survey estimate for area i and time $t = 1, \dots, T$ of population characteristic Y_{it} , so $y_{it} = Y_{it} + e_{it}$ where e_{it} is the sampling error. For simplicity we assume the model parameters are known (first order approximation) and also assume normality, so that the best linear predictor (BLP) is the conditional expectation and the pre-

diction MSE is the conditional variance. With parameters assumed known we need not explicitly consider the regression mean for $E(Y_{it})$, as this does not affect the conditional variances, which are our focus here. Also, since the FH model assumes independence over areas i , the BLP for area i then uses data for only that area, so we simplify the notation by dropping the subscript i . We further simplify by assuming that $\text{var}(Y_t) = \text{var}(u_t) = \sigma_u^2$ and $\text{var}(e_t) = v$ are constant over time. Within this simplified setup, we seek MSEs for the bivariate and FH-AR(1) predictors of Y_T , the most recent true population quantity, given data $\mathbf{y} = [y_1, \dots, y_T]'$.

Let $\mathbf{z} = \mathbf{A}\mathbf{y} = \mathbf{A}\mathbf{Y} + \mathbf{A}\mathbf{e}$ where $\mathbf{A}$ is a $k \times T$ matrix of the form

$$\mathbf{A} = \begin{bmatrix} A_{11} & \mathbf{0} \\ \mathbf{0}' & 1 \end{bmatrix} \Rightarrow \mathbf{z} = \begin{bmatrix} A_{11}\mathbf{y}_1 \\ y_T \end{bmatrix} \quad (13)$$

where $\mathbf{y}_1 = [y_1, \dots, y_{T-1}]'$ and $\mathbf{0}$ is a $(k-1) \times 1$ vector of zeroes. For the FH-AR(1) model $k = T$ and $A_{11} = I_{T-1}$, while for the bivariate model $k = 2$ and $A_{11} = (T-1)^{-1}[1, \dots, 1]$. Letting $\Sigma_e \equiv \text{var}(\mathbf{e})$, and similarly defining Σ_u , Σ_y , and Σ_z , we have $\Sigma_y = \Sigma_u + \Sigma_e$, $\Sigma_z = \mathbf{A}\Sigma_y\mathbf{A}'$, and $\text{cov}(\mathbf{e}, \mathbf{z}) = \Sigma_e\mathbf{A}'$. From standard results on conditional variance in a multivariate normal distribution, and since predicting $\mathbf{e}$ is equivalent to predicting $\mathbf{Y}$, then

$$\text{var}(\mathbf{Y}|\mathbf{z}) \equiv \text{var}(\mathbf{e}|\mathbf{z}) = \Sigma_e - \Sigma_e\mathbf{A}'(\mathbf{A}\Sigma_y\mathbf{A}')^{-1}\mathbf{A}\Sigma_e.$$

We are assuming $\Sigma_e = vI$, and we write $\Sigma_u = \sigma_u^2 R$, where R is the $T \times T$ correlation matrix of $\mathbf{Y}$. Then $\Sigma_y = \sigma_u^2(\lambda I + R)$ where $\lambda = v/\sigma_u^2$ is the noise-to-signal ratio. Thus,

$$\begin{aligned} \text{var}(\mathbf{e}|\mathbf{z}) &= vI - v\mathbf{A}'[\sigma_u^2\mathbf{A}(\lambda I + R)\mathbf{A}']^{-1}\mathbf{A}v \\ &= v\{I - \lambda\mathbf{A}'[\mathbf{A}(\lambda I + R)\mathbf{A}']^{-1}\mathbf{A}\}. \end{aligned}$$

Let $\Omega \equiv [\omega_{j\ell}] = [\mathbf{A}(\lambda I + R)\mathbf{A}']^{-1}$. We are interested in the (T, T) th element of $\text{var}(\mathbf{e}|\mathbf{z})$, which is

$$\text{var}(Y_T|\mathbf{z}) \equiv \text{var}(e_T|\mathbf{z}) = v \left\{ 1 - \lambda [\mathbf{0}', 1]\mathbf{A}'\Omega\mathbf{A} \begin{bmatrix} \mathbf{0} \\ 1 \end{bmatrix} \right\}.$$

From the definition of $\mathbf{A}$ in equation (13), $[\mathbf{0}', 1]\mathbf{A}' = [\mathbf{0}', 1]$, so that this reduces to

$$\text{var}(Y_T|\mathbf{z}) = v(1 - \lambda\omega_{TT}) \quad (14)$$

where ω_{TT} is the (T, T) th element of Ω . The expression (14) is easily computed

given v , σ_u^2 , and R . For our comparisons, R is the AR(1) correlation matrix given in equation (9), which is determined solely by ϕ . Hence, Ω is determined by λ and ϕ . Note that for the bivariate model, A is $2 \times T$ and Ω is then a 2×2 matrix.

The prediction MSE of the univariate FH model is $\text{var}(Y_T|y_T) = \sigma_u^2 v / (\sigma_u^2 + v)$ (Rao and Molina 2015, eq. (6.1.8)). The percent reduction in prediction MSE from the FH-AR(1) or bivariate models relative to the univariate FH model is thus 100 times

$$\begin{aligned} 1 - \frac{\text{var}(Y_T|\mathbf{z})}{\text{var}(Y_T|y_T)} &= 1 - \frac{\sigma_u^2 + v}{\sigma_u^2 v} v(1 - \lambda \omega_{TT}) \\ &= 1 - (1 + \lambda)(1 - \lambda \omega_{TT}). \end{aligned}$$

This expression depends on only λ and ϕ .

Acknowledgements

We thank Roderick Little, Hal Stern, and Alan Zaslavsky who, in a 2007 review of SAIPE models, suggested investigation of binomial models using effective sample sizes for modeling ACS poverty estimates.

Disclaimer

This report is released to inform interested parties of research and to encourage discussion. The views expressed on statistical, methodological, technical, or operational issues are those of the authors and not necessarily those of the U.S. Census Bureau.

REFERENCES

- BELL, W. R., BASEL, W. W., CRUSE, C., DALZELL, L., MAPLES, J. J., O'HARA, B., POWERS, D. (2007). Use of ACS Data to Produce SAIPE Model-Based Estimates of Poverty for Counties. SAIPE technical report, U.S. Census Bureau, URL <http://www.census.gov/did/www/saipe/publications/files/report.pdf>.
- BELL, W. R., BASEL, W. W., MAPLES, J. J., (2015). An Overview of the U.S. Census Bureau's Small Area Income and Poverty Estimates (SAIPE) Program. In *Analysis of Poverty Data by Small Area Methods*, Monica Pratesi (ed.), London: Wiley, to appear.
- BOX, G. E. P., JENKINS, G. M. (1970). *Time Series Analysis: Forecasting and Control*, San Francisco: Holden Day.

- BROWN, L. D., CAI, T. T., DASGUPTA, A., (2001). Interval Estimation for a Binomial Proportion, *Statistical Science*, 16, 101-133.
- BROWN, G., CHAMBERS, R., HEADY, P., HEASMAN, D. (2001). Evaluation of Small Area Estimation Methods – An Application to Unemployment Estimates from the UK LFS. In *Proceedings of the Statistics Canada Symposium 2001, Achieving Data Quality in a Statistical Agency: A Methodological Perspective*.
- DATTA, G. S., LAHIRI, P., MAITI, T., LU, K. L. (1999). Hierarchical Bayes Estimation of Unemployment Rates for the States of the U.S. *Journal of the American Statistical Association*, 94, 1074-1082.
- ESTEBAN, M. D., MORALES, D., PEREZ, A., SANTAMARIA, L., (2012). Small Area Estimation of Poverty Proportions Under Area-Level Time Models. *Computational Statistics and Data Analysis*, 56, 2840-2855.
- FAY, R. E., HERRIOT, R. A. (1979). Estimation of Income for Small Places: An Application of James-Stein Procedures to Census Data. *Journal of the American Statistical Association*, 74, 269–77.
- FRANCO, C., BELL, W. R. (2013). Applying Bivariate Binomial/Logit Normal Models to Small Area Estimation. *Proceedings of the American Statistical Association, Section on Survey Research Methods*, 690–702, URL <http://www.amstat.org/sections/srms/Proceedings/>.
- GHOSH, M., NATARAJAN, K., STROUD, T.W.F., CARLIN, B.P. (1998). Generalized Linear Models for Small Area Estimation. *Journal of the American Statistical Association*, 93, 273–282.
- GHOSH, M., NANGIA, N., KIM, D. H. (1996). Estimation of Median Income of Four-Person Families: A Bayesian Time Series Approach. *Journal of the American Statistical Association*, 91, 1423-1431.
- HAWALA, S., LAHIRI, P. (2012). Hierarchical Bayes Estimation of Poverty Rates. *Proceedings of the American Statistical Association, Survey Research Methods Section*, 3410–3424, URL <http://www.census.gov/did/www/saibe/publications/files/hawalalahirishpl2012.pdf>.
- HUANG, E. T., BELL, W. R. (2012). An Empirical Study on Using Previous American Community Survey Data Versus Census 2000 Data in SAIPE Models for Poverty Estimates. Research Report Number RRS2012–4, Center for Statistical Research and Methodology, U.S. Census Bureau, URL <http://www.census.gov/srd/papers/pdf/rrs2012-04.pdf>.
- LIU, B. M., LAHIRI, P., KALTON, G. (2007). Hierarchical Bayes Modeling of Survey Weighted Small Area Proportions. *Proceedings of the American Statistical Association, Survey Research Section*, 3181–3186.
- PLUMMER, M. (2010). JAGS - Just Another Gibbs Sampler (JAGS 2.1.0.,

- May 12, 2010). URL <https://sourceforge.net/projects/mcmc-jags>.
- PRATESI, M., GIUSTI, C., MARCHETTI, S., SALVATI, N., TZAVIDIS, N., MOLINA, I., DURBÁN, M., GRANÉ, A., MARÍN, J. M., VEIGA, M. H., MORALES, D., ESTEBAN, M.D., SAÑCHEZ, A., SANTAMARÍA, L., MARHUENDA, Y., PÉREZ, A., PAGLIARELLA, M. C., RAO, J. N. K., FERRETTI, C. (2010). Small Area Estimation of Poverty and Inequality Indicators: Pilot Applications. Work Package 2 (D17) of Small Area Methods for Poverty and Living Conditions Estimates Project. URL <http://www.sample-project.eu/SAMPLEwp2d17.pdf>.
- R CORE TEAM (2013). R: A Language and Environment for Statistical Computing. Vienna, Austria: R Foundation for Statistical Computing. URL <http://www.R-project.org/>.
- RAO, J. N. K., MOLINA, I. (2015). *Small Area Estimation* (2nd ed.), Hoboken, New Jersey: Wiley.
- SAEI, A., CHAMBERS, R. (2003). Small Area Estimation Under Linear and Generalized Linear Mixed Models with Time and Area Effects. Southampton, UK, Southampton Statistical Sciences Research Institute. (S3RI Methodology Working Papers, (M03/15)). URL <http://eprints.soton.ac.uk/8165/>.
- SLUD, E. V. (2000). Models for Simulation and Comparison of SAIPE Analyses. SAIPE Technical Report, U.S. Census Bureau, URL <http://www.census.gov/did/www/saipe/publications/files/saipemod.pdf>.
- SLUD, E. V. (2004). Small Area Estimation Errors in SAIPE Using GLM versus FH Models. *Proceedings of the American Statistical Association, Section on Survey Research Methods*, 4402–4409, URL <http://www.amstat.org/sections/srms/Proceedings/>.
- TACIAK, J., BASEL, W. W. (2012). Time Series Cross Sectional Approach for Small Area Poverty Models. *Proceedings of the American Statistical Association, Section on Government Statistics*, URL <http://www.census.gov/did/www/saipe/publications/files/JTaciakBaseljsm2012.pdf>.
- U.S. CENSUS BUREAU (2014). *American Community Survey Design and Methodology (version 2.0, January 2014)*, URL http://www.census.gov/acs/www/methodology/methodology_main/.
- YOU, Y. (2008). An Integrated Modeling Approach to Unemployment Rate Estimation for Sub-Provincial Areas of Canada. *Survey Methodology*, 34, 19–27.