

BAYESIAN SAE WITH AN EMPHASIS ON LOW- AND MIDDLE-INCOME COUNTRIES: LECTURE 3: AREA-LEVEL SAE MODELS

Jon Wakefield

Departments of Statistics and Biostatistics
University of Washington

MOTIVATION

THE NON-SPATIAL AREA-LEVEL MODEL

THE SPATIAL AREA-LEVEL MODEL

AREA-LEVEL PREVALENCE MODELS

SUBNATIONAL NMR ESTIMATION IN RWANDA:
AREA-LEVEL MODELS

MOTIVATION

BEYOND WEIGHTED ESTIMATES

- ▶ **Aim of Small Area Estimation (SAE):** Provide inferential summaries across geographical areas, using survey data – in a LMICs context, the household surveys use **stratified two-stage unequal probability cluster sampling**.
- ▶ Recall the **weighted estimator**, which provides a simple approach and has many desirable properties in data-rich situations.
- ▶ As data become sparser, weighted estimates can be highly unstable (high variance) or we can't calculate a measure of uncertainty.
- ▶ One then needs to consider SAE models, of which there are two distinct types:

1: Area-Level Models

2: Unit-Level Models

THE AREA-LEVEL FAY-HERRIOT MODEL

- ▶ An amazingly advanced model was introduced by Fay and Herriot (1979).
- ▶ The basic idea behind the area-level Fay-Herriot model, is to **simultaneously model** the collection of weighted estimates

$$\hat{\theta}_i^w = \frac{\hat{T}_i}{\hat{N}_i} = \frac{\sum_{k \in S_i} w_{ik} Y_{ik}}{\sum_{k \in S_i} w_{ik}}. \quad (1)$$

from **all areas**.

- ▶ These estimates, $\hat{\theta}_i^w$, along with their standard errors, $\hat{V}_i^{1/2}$, constitute **the data** for $i = 1, \dots, m$, areas.
- ▶ We model between-area variation in $\hat{\theta}_i^w$ using covariates and by assuming the residuals arise from a common distribution.
- ▶ We are effectively **increasing the sample size** in each area, by leveraging similarity in the means.

THE NON-SPATIAL AREA-LEVEL MODEL

THE MOST BASIC FORM OF AREA-LEVEL MODEL

- ▶ We let $Y_i = \hat{\theta}_i^w$ correspond to the “data” and $E[Y_i] = \theta_i$ – here we are performing **design-based inference** and the expectation is over samples that could have been drawn, and θ_i is the true **finite population mean** in area i .
- ▶ We describe a two-stage model in which we have a sampling model and a linking model.
- ▶ If there are sufficient samples within the area, a design-based asymptotic arguments (Breidt and Opsomer, 2017) gives $Y_i \mid \theta_i \sim N(\theta_i, \hat{V}_i)$ where $\hat{V}_i$ is the design-based variance.
- ▶ When we fit this model, it is an example of **indirect estimation** because the data from all areas is being used to enhance the estimates in each area.

THE MOST BASIC FORM OF AREA-LEVEL MODEL

The simplest version of the area-level model is:

$$Y_i = \theta_i + \epsilon_i, \quad \epsilon_i \sim_{ind} N(0, \hat{V}_i) \quad (2)$$

$$\theta_i = \alpha + u_i, \quad u_i \sim_{iid} N(0, \sigma_u^2). \quad (3)$$

- ▶ In the **first stage of the model**, in line (2) we have the **sampling model** with variances $\hat{V}_i$ treated as known.
- ▶ The notation **ind** is short for **independent**.
- ▶ In the **second stage of the model**, in line (3) we have the **linking model** which ties together $\theta_1, \dots, \theta_m$.
- ▶ The notation **iid** is short for **independent and identically distributed**.
- ▶ **Parameters:** $\alpha = E[\theta_i]$ is the intercept and is the average of the population means, $\sigma_u^2 = \text{var}(\theta_i)$ is the variance of the second-stage residuals u_i .

ESTIMATION FOR THE AREA-LEVEL MODEL

- ▶ Suppose for simplicity that α and σ_u^2 are known.
- ▶ In this case the posterior mean estimate of θ_i is the so-called **Best Linear Unbiased Predictor (BLUP)**, which is the posterior mean,

$$\begin{aligned}\hat{\theta}_i^{\text{BLUP}} = \text{E}[\theta_i \mid y_i] &= Y_i - \frac{\hat{V}_i}{\hat{V}_i + \sigma_u^2} (Y_i - \alpha) \\ &= F_i \times Y_i + (1 - F_i) \times \alpha\end{aligned}$$

where

$$F_i = \frac{\sigma_u^2}{\hat{V}_i + \sigma_u^2} \leq 1$$

is the **shrinkage factor** – estimates are shrunk towards α .

- ▶ The intuition is that the estimate is a compromise between the observed data in the area, Y_i , and the average of the data over the complete set of areas, α .

ESTIMATION FOR THE AREA-LEVEL MODEL

Recall,

$$\hat{\theta}_i^{\text{BLUP}} = F_i \times Y_i + (1 - F_i) \times \alpha, \quad F_i = \frac{\sigma_u^2}{\hat{V}_i + \sigma_u^2}$$

- A measure of uncertainty is the **posterior variance**,

$$\text{var}(\theta_i | y_i) = F_i \times \hat{V}_i,$$

so that the uncertainty is reduced, compared to the direct (weighted) estimate variance ($\hat{V}_i$).

- Notice that if σ_u^2 is big and/or $\hat{V}_i$ is small, we have $F_i \approx 1$ so that

$$\hat{\theta}_i^{\text{BLUP}} \approx Y_i, \quad \text{var}(\theta_i | y_i) \approx \hat{V}_i.$$

If we imagine the sample size increasing in an area, then we are in this situation as $\hat{V}_i \rightarrow 0$.

- Informally, this explains why Fay-Herriot estimates are **design consistent**.

FAY-HERRIOT MODEL WITH COVARIATES

To model between-area variation in Y_i we may introduce area-level covariates $\mathbf{x}_i$, a $p \times 1$ vector.

With covariates:

$$Y_i = \theta_i + \epsilon_i, \quad \epsilon_i \sim_{ind} \mathcal{N}(0, \hat{V}_i) \quad (4)$$

$$\theta_i = \alpha + \mathbf{x}_i^T \boldsymbol{\beta} + u_i, \quad u_i \sim_{iid} \mathcal{N}(0, \sigma_u^2). \quad (5)$$

► The BLUP becomes:

$$\begin{aligned} \hat{\theta}_i^{BLUP} = E[\theta_i | y_i] &= Y_i - \frac{\hat{V}_i}{\hat{V}_i + \sigma_u^2} (Y_i - \alpha - \mathbf{x}_i^T \boldsymbol{\beta}) \\ &= F_i \times Y_i + (1 - F_i) \times (\alpha + \mathbf{x}_i^T \boldsymbol{\beta}) \end{aligned}$$

with $F_i = \sigma_u^2 / (\hat{V}_i + \sigma_u^2)$, as before.

INCLUDING COVARIATES

- ▶ Covariates derived from a **census** are desirable but in LMICs censuses are infrequently carried out, and the data may not be easily accessible at the level required.
- ▶ Instead, **satellite data** are often used.
- ▶ It is appealing to use **survey covariates**, but there are two issues that require consideration:
 1. The covariates are not population-level summaries, but rather estimates – this **errors-in-variables** aspect requires careful thought see, for example, Ybarra and Lohr (2008).
 2. **Correlated errors** at the first stage of the model, if response and covariate are taken from the same survey – we discuss this a little further.

SURVEY-BASED COVARIATES

- ▶ For simplicity assume a univariate x_i – if the same response data Y_i and covariate data X_i (the version with error) are from the **same survey** then we violate a key assumption of the model.
- ▶ We have

$$\begin{aligned} Y_i &= \theta_i + \epsilon_i, & \epsilon_i &\sim_{ind} N(0, \hat{V}_{yi}) \\ X_i &= x_i + \eta_i, & \eta_i &\sim_{ind} N(0, \hat{V}_{xi}) \end{aligned}$$

- ▶ Because both the area-level response Y_i and the auxiliary covariate X_i are direct (weighted) estimates derived from the same survey sample, the standard Fay-Herriot assumption of mutually independent sampling errors $\text{cov}(\epsilon_i, X_i) = 0$ is violated.
- ▶ The shared sampling variance induces an **endogeneity** problem via area-level measurement error.

THE LINEAR MIXED EFFECTS MODEL

We can write the two-stage model (4) and (5) as

$$Y_i = \alpha + \mathbf{x}_i^\top \boldsymbol{\beta} + u_i + \epsilon_i,$$

with

$$u_i \mid \sigma_u^2 \sim_{iid} \mathcal{N}(0, \sigma_u^2), \quad \epsilon_i \sim_{ind} \mathcal{N}(0, \hat{V}_i), \quad i = 1, \dots, m.$$

In this model:

- ▶ α and $\boldsymbol{\beta}$ are **fixed effects**.
- ▶ $u_1, \dots, u_m$ are **random effects** – they have a distribution!
- ▶ Hence, we have a **mixed effects model**, and it is linear, and so, more specifically, it is a **linear mixed effects model (LMEM)**.
- ▶ Note: This is a **frequentist** distinction since in a **Bayesian** framework, all parameters are random.

THE SPATIAL AREA-LEVEL MODEL

AREA-LEVEL (FAY-HERRIOT) SPATIAL MODEL

Recall the [area-level Fay-Herriot model](#):

$$\begin{aligned}Y_i &= \theta_i + \epsilon_i, & \epsilon_i &\sim_{ind} \mathbf{N}(0, \widehat{V}_i) \\ \theta_i &= \alpha + \mathbf{x}_i^\top \boldsymbol{\beta} + u_i,\end{aligned}$$

where previously we assumed $u_i \mid \sigma_u^2 \sim_{iid} \mathbf{N}(0, \sigma_u^2)$, $i = 1, \dots, m$.

- In the spatial version of the model, the “residuals” (random effects)

$$\mathbf{u} = (u_1, \dots, u_m)$$

are modeled in such a way that if areas i and j are “close” they are more likely to be similar.

- The most common way in which closeness is defined, is to say that areas i and j are [neighbors](#) if they share a [common boundary](#).

NEIGHBORHOOD GRAPH IN RWANDA

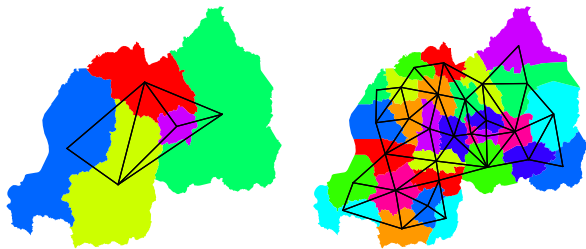


FIGURE 1: Neighborhood structure for 5 Admin1 areas and 30 Admin2 areas in Rwanda – this is used in the [spatial smoothing models](#). Black lines link [neighbors](#).

AREA-LEVEL (FAY-HERRIOT) SPATIAL MODEL

There are many spatial modeling possibilities for $\mathbf{u}$ (Banerjee *et al.*, 2015) including:

- ▶ Conditional Autoregressive (CAR) model.
- ▶ Simultaneous Autoregressive (SAR) model.
- ▶ Intrinsic CAR (ICAR) model.
- ▶ Combining an independent and identically distributed (IID) model with an ICAR model, to give the famous Besag, York, Mollié (BYM) model (Besag *et al.*, 1991).
- ▶ We describe a reformulation of the BYM model which has certain advantages – known as the **BYM2 model** and due to Riebler *et al.* (2016).

AREA-LEVEL (FAY-HERRIOT) SPATIAL MODEL

We write

$$\mathbf{u} \mid \sigma_u^2, \phi \sim \text{BYM2}(\sigma_u^2, \phi).$$

The spatial residual is decomposed as:

$$\mathbf{u} = \begin{bmatrix} u_1 \\ \vdots \\ u_m \end{bmatrix} = \lambda \left(\sqrt{\phi} \begin{bmatrix} s_1 \\ \vdots \\ s_m \end{bmatrix} + \sqrt{1 - \phi} \begin{bmatrix} e_1 \\ \vdots \\ e_m \end{bmatrix} \right)$$

where

- ▶ The s_i are the **spatial terms** and the $e_i \sim N(0, 1)$ are the **non-spatial (IID) terms**, $i = 1, \dots, m$.
- ▶ Two parameters in the model: ϕ is the **proportion of the variation that is spatial** and λ is the **residual standard deviation** – tells us how bumpy the surface is and corresponds to the **smoothing parameter**.

THE ICAR SPATIAL MODEL

- ▶ The ICAR model is very simple – the “best guess” of the level is taken as the mean of the neighbors, with variation allowed around that level.
- ▶ The explicit form for the ICAR spatial terms is,

$$s_i \mid s_j, j \in \text{ne}(i), \sigma_s^2 \sim_{\text{ind}} \text{N} \left(\bar{s}_i, \frac{\sigma_s^2}{m_i} \right),$$

which is a **local smoothing model** in which:

- ▶ $\text{ne}(i)$ is the set of neighbors of i , of which there are m_i ,
- ▶ $\bar{s}_i = \frac{1}{m_i} \sum_{j \in \text{ne}(i)} s_j$ is the mean of the neighbors,
- ▶ σ_s^2 is the variance that describes the variation around the conditional mean.

THE ICAR SPATIAL MODEL

- We can write the spatial model as:

$$\begin{aligned}\pi(\mathbf{s} \mid \sigma_s^2) &\propto (\sigma_s^2)^{-(m-1)/2} \exp \left(-\frac{1}{2\sigma_s^2} \sum_{i \sim j} (s_i - s_j)^2 \right) \\ &= (\sigma_s^2)^{-(m-1)/2} \exp \left(-\frac{\mathbf{s}^\top \mathbf{Q} \mathbf{s}}{2\sigma_s^2} \right)\end{aligned}$$

- The “precision” matrix $\mathbf{Q}$ associated has elements

$$Q_{ij} = \begin{cases} m_i & i = j \\ -1 & i \sim j \\ 0 & \text{otherwise} \end{cases}$$

- This model is extremely computationally efficient to work with, because $\mathbf{Q}$ is sparse.

THE ICAR SPATIAL MODEL

- ▶ In the weeds, in the implementation we use, the ICAR precision is scaled so that the variance σ_s^2 has a marginal interpretation, see Riebler *et al.* (2016) for details.
- ▶ This leads to a precision matrix $\mathbf{Q}_\star$ and scaled random effects $\mathbf{s}_\star$.

$$\mathbf{u} = \begin{bmatrix} u_1 \\ \vdots \\ u_m \end{bmatrix} = \sigma_u^2 \left(\sqrt{\phi} \begin{bmatrix} \mathbf{s}_{\star,1} \\ \vdots \\ \mathbf{s}_{\star,m} \end{bmatrix} + \sqrt{1-\phi} \begin{bmatrix} \mathbf{e}_1 \\ \vdots \\ \mathbf{e}_m \end{bmatrix} \right)$$

where $e_i \sim_{iid} \mathcal{N}(0, 1)$, $i = 1, \dots, m$, as before.

- ▶ We have

$$\text{var}(\mathbf{u} \mid \sigma_u^2, \phi) = \sigma_u^2 [\phi \mathbf{Q}_\star^- + (1 - \phi) \mathbf{I}]$$

where $\mathbf{Q}_\star^-$ is the generalized inverse.

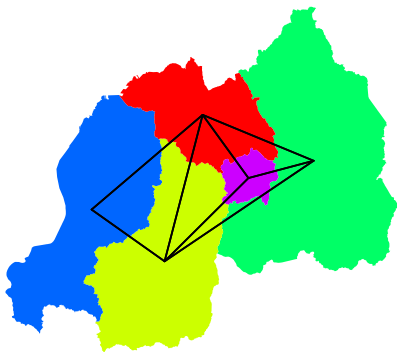


FIGURE 2: Neighborhood structure for 5 Admin1 areas. Black lines link neighbors.

- For example, the level for area 1 (note: numbering is arbitrary) is being pulled towards the average of the levels in areas 2, 3, 4 (its neighbors):

$$E[S_1 \mid S_2, S_3, S_4] = \frac{1}{3}(S_2 + S_3 + S_4)$$

AREA-LEVEL PREVALENCE MODELS

AREA-LEVEL PREVALENCE MODELS

- ▶ When the target variable is not on the whole of the real line, we may choose to take a transformation.
- ▶ Fay and Herriot (1979) carried out SAE for per capita income for local geographic areas like townships, cities, and boroughs level — the response variable is positive and the log transformation was taken.
- ▶ Prevalences are by far and away the most common response types in the DHS.
- ▶ Prevalances lie between 0 and 1, which makes modeling a little more tricky – various transformations are possible and we choose the **logit transform**; the inverse arcsin is an alternative.

AREA-LEVEL PREVALENCE MODELS

Let,

$$\theta_i = \log \left(\frac{p_i}{1 - p_i} \right), \quad i = 1, \dots, m,$$

be the log odds of the observed prevalence, which is on the real line.

Modeling proceeds as before, but on the logit scale.

The response variable is taken as,

$$Y_i = \log \left(\frac{\hat{p}_i^w}{1 - \hat{p}_i^w} \right), \quad \text{for } i = 1, \dots, m.$$

Input Data:

$$\underbrace{(Y_1, \hat{V}_1), (Y_2, \hat{V}_2), \quad \dots, \quad (Y_{m-1}, \hat{V}_{m-1}), (Y_m, \hat{V}_m)}_{\text{AREAS ARE LINKED THROUGH COMMON MODEL}}$$

AREA-LEVEL PREVALENCE MODELS

In large samples,

$$Y_i = \underbrace{\theta_i}_{\text{TRUTH}} + \underbrace{\epsilon_i}_{\text{SAMPLING ERROR}},$$

with

- ▶ ϵ_i following a $N(0, \hat{V}_i)$ distribution,
- ▶ $\hat{V}_i$ estimated using formulas relevant to the survey design.
- ▶ We might expect the logit transform to improve the normal approximation to the sampling distribution, but note that $E[Y_i] \neq \theta_i$ though Y_i is a consistent estimator of θ_i – later we will discuss this aspect further in the context of **unmatched linking models**.

AREA-LEVEL PREVALENCE MODELS

- ▶ As with the area-level model we described previously, we can have IID or spatial random effects and can include area-level covariates.
- ▶ In the **Fay-Herriot spatial model with covariates** we have:

$$Y_i = \theta_i + \epsilon_i, \quad \epsilon_i \sim_{ind} N(0, \hat{V}_i)$$
$$\theta_i = \underbrace{\alpha}_{\text{OVERALL LEVEL}} + \underbrace{\mathbf{x}_i \boldsymbol{\beta}}_{\text{COVARIATE MODEL}} + \underbrace{u_i}_{\text{SPATIAL RESIDUAL}}, \quad i = 1, \dots, m$$

The spatial terms $u_1, \dots, u_m$ are assigned a BYM2 model with parameters σ_u^2, ϕ .

AREA-LEVEL (FAY-HERRIOT) MODEL

Advantages:

- ▶ Weighted estimates and uncertainty measures that feed into the hierarchical model account for the design, which avoids potential bias due to (say) stratified and PPS sampling, and accounts for the clustering.
- ▶ As the sample size in each area increases, the estimates tend to the true prevalences, which is known as **design consistency**.
- ▶ By modeling all the data in unison, **uncertainty in each area is reduced** (on average).
- ▶ The benefits arise from both the **random effects** and the **covariates**.

AREA-LEVEL (FAY-HERRIOT) MODEL

Disadvantages:

- ▶ The modeling is based on the weighted estimates and standard errors in each area, so when the data are sparse and estimates are not available or unreliable, the method cannot be used.
- ▶ When the data are sparse, the estimates may be **overshrunk** since the data are not sufficiently informative to discriminate from the “flat” map case.
- ▶ The model produces **shrinkage** which can lead to interpretation being more tricky (since bias is introduced in each estimate).

SUBNATIONAL NMR ESTIMATION IN RWANDA: AREA-LEVEL MODELS

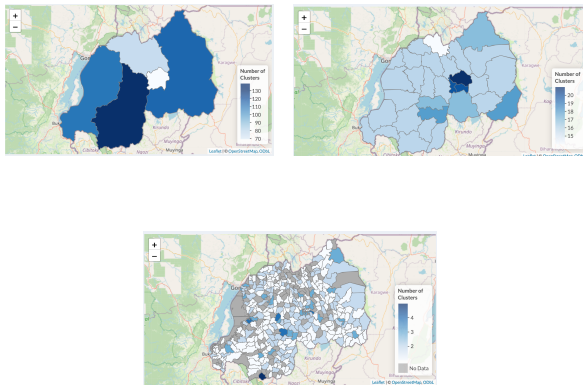


FIGURE 3: Number of sampled clusters, for Admin1/Admin2/Admin3 areas, from Rwanda 2019 DHS.

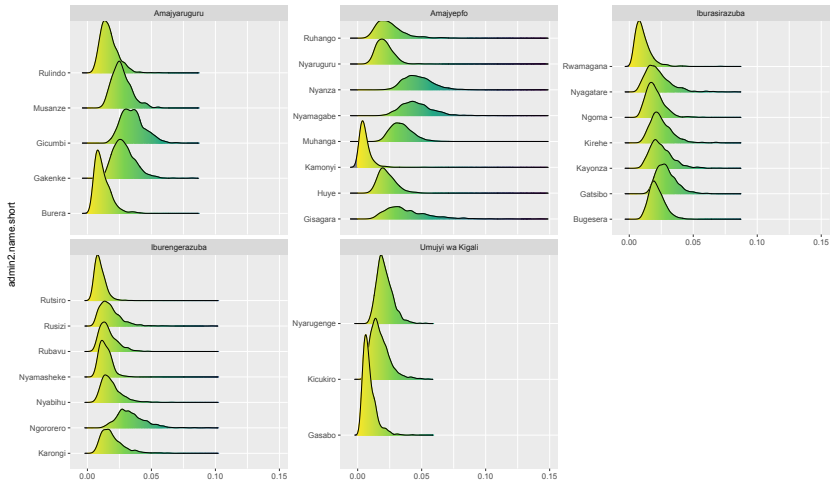


FIGURE 4: Ridgeplots for Admin2 areas *within* Admin1, direct estimates.

- For some Admin1, the constituent Admin2 estimates have similar distributions, while for others they are quite different.

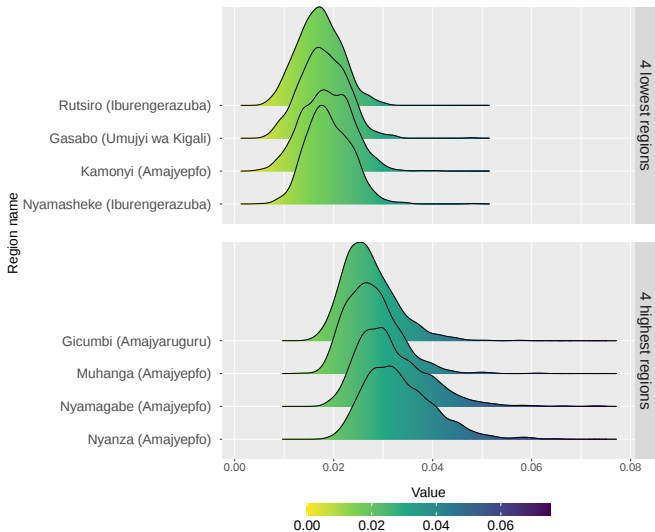


FIGURE 5: Ridgeplots for extreme (low and high) Admin2 areas.

- Lack of information at Admin2, which explains the relatively wide distributions and large overlap across areas.

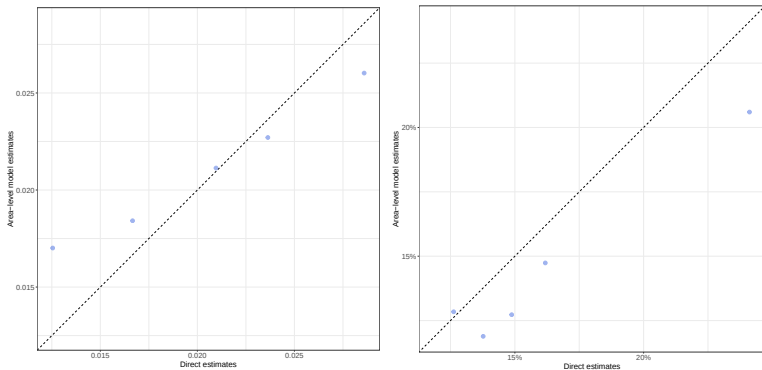


FIGURE 6: Area-level estimates versus direct estimates (left) and area-level CVs versus direct CVs (right), at **Admin1**.

- ▶ Some shrinkage (narrowing of range of estimates) of estimates at Admin1 (left).
- ▶ Gains in precision in all but one Admin1 area (right).

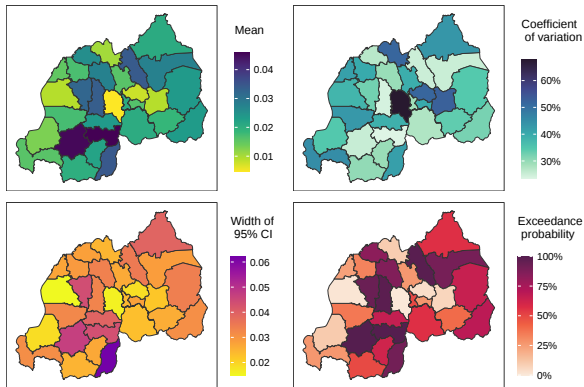


FIGURE 7: Summary measures for direct estimates at Admin2.

- Relatively large range in point estimates.
- Uncertainty measure show large variation also, and for some areas, estimates are quite imprecise.

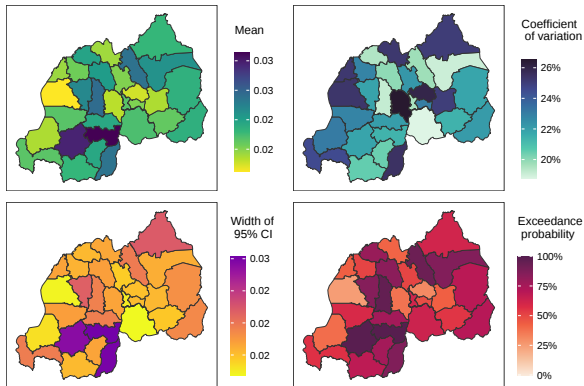


FIGURE 8: Summary measures for area-level (Fay-Herriot) estimates at Admin2.

- ▶ Compared to direct estimates (Figure 7) a narrower range in the point estimates.
- ▶ Uncertainty measure show less variation also, due to sharing of information.

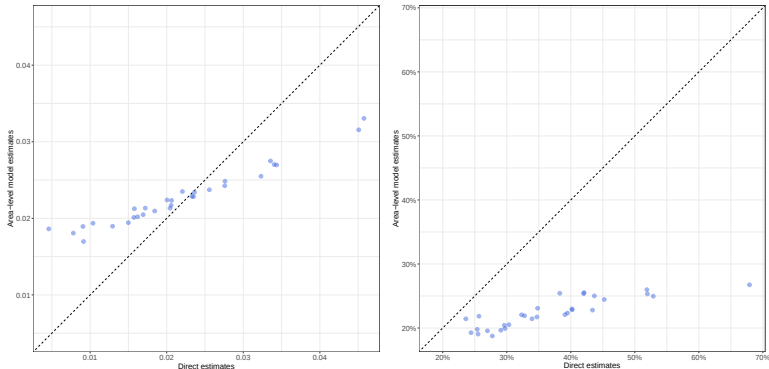


FIGURE 9: Area-level estimates versus direct estimates (left) and area-level CVs versus direct CVs (right), at **Admin2**.

- ▶ As compared to Figure 6, much greater shrinkage at Admin2 (left).
- ▶ As compared to Figure 6, much greater gains in precision in all **Admin2** areas (right).

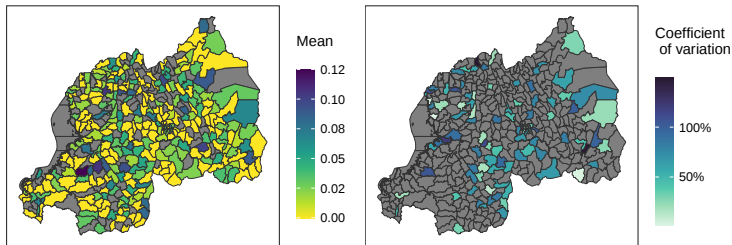


FIGURE 10: Direct prevalence estimates and CV maps at Admin3.

- ▶ Direct estimates are not available in quite a few areas (since no data).
- ▶ CVs are available in very few, because even if data are available, some configurations do not allow calculation of a variance, e.g., all responses zero.
- ▶ So unable to fit area-level models at Admin3 – unit-level models are the only possibility.
- ▶ The shinyApp gives warnings.

Data Sparsity Check				
	National	Admin-1	Admin-2	Admin-3
Direct Estimates	✓Model is ready to be implemented.	✓Model is ready to be implemented.	✓Model is ready to be implemented.	⚠ Data is not present/sparse in 77% of areas.
Area-level Model		✓Model is ready to be implemented.	✓Model is ready to be implemented.	✗ Data is not present/sparse in 77% of areas, so model will not be fitted.
Unit-level Model		✓Model is ready to be implemented.	✓Model is ready to be implemented.	✓Model is ready to be implemented.

Model Fitting				
Run models that passed check Run all selected models				
	National	Admin-1	Admin-2	Admin-3
Direct Estimates	✓ Successful	✓ Successful	✓ Successful	⚠ Model fitted, but interpret with caution due to data sparsity.
Area-level Model		✓ Successful	✓ Successful	❌ Model will not be fitted.
Unit-level Model		✓ Successful	✓ Successful	✓ Successful

FIGURE 11: The shinyApp gives a warning that direct estimates not available in many Admin3 areas, and so this does not allow the fitting of an Admin3 area-level Fay-Herriot model.

- ▶ **Fay-Herriot area-level models** are by far the most popular SAE models.
- ▶ **Overshrinkage** a worry: **nested spatial models** (which we discuss later) are a worthwhile alternative.
- ▶ **Model checking** still in infancy: we use cross-validation to predict weighted estimates.
- ▶ **Machine learning** approaches are very intoxicating, but very difficult to quantify uncertainty.

- ▶ Benchmarking to known totals is less popular in LMICs because the population information may be inaccurate or out of date, see Okonek and Wakefield (2022) for references to the benchmarking literature.
- ▶ Variance modeling (which we discuss later) can be carried out to increase the utility of this method (Gao and Wakefield, 2023) – a method for variance modification is available in the `surveyPrev` package.

References

- Banerjee, S., Carlin, B., and Gelfand, A. (2015). *Hierarchical Modeling and Analysis for Spatial Data, Second Edition*. CRC Press.
- Besag, J., York, J., and Mollié, A. (1991). Bayesian image restoration with two applications in spatial statistics. *Annals of the Institute of Statistics and Mathematics*, **43**, 1–59.
- Breidt, J. and Opsomer, J. (2017). Model-assisted survey estimation with modern prediction techniques. *Statistical Science*.
- Fay, R. and Herriot, R. (1979). Estimates of income for small places: an application of James–Stein procedure to census data. *Journal of the American Statistical Association*, **74**, 269–277.
- Gao, P. A. and Wakefield, J. (2023). A spatial variance-smoothing area level model for small area estimation of demographic rates. *International Statistical Review*, **91**, 493–510.
- Okonek, T. and Wakefield, J. (2022). A computationally efficient approach to fully Bayesian benchmarking. *Journal of Official Statistics*. Published online: May 27, 2024.
- Riebler, A., Sørbye, S., Simpson, D., and Rue, H. (2016). An intuitive Bayesian spatial model for disease mapping that accounts for scaling. *Statistical Methods in Medical Research*, **25**, 1145–1165.

Ybarra, L. M. and Lohr, S. L. (2008). Small area estimation when auxiliary information is measured with error. *Biometrika*, **95**, 919–931.