

1

Introduction on Measuring Poverty at Local Level Using Small Area Estimation Methods

Monica Pratesi and Nicola Salvati

Department of Economics and Management, University of Pisa, Pisa, Italy

1.1 Introduction

All over the world, fighting against poverty is assuming a more and more central role and recent radical economic and social transformations have caused a renewed interest in this field. Poverty is a complex concept. As a consequence, the focus should not be only on monetary poverty, but also on the larger concept of well-being, which preliminarily includes the definition and measure of the following aspects: capability of income production, being involved in a satisfying job, being in good health, living in an adequate house, achieving a proper level of education, having good social relations, and so on. These characteristics require poverty to be defined in a multidimensional setting.

Given that, the reduction of the risk of becoming poor can be achieved only through a very wide range of policy actions and tools: from the mere monetary transfer to a varied supply of social services.

Local governments play a fundamental role in implementing actions to provide help to vulnerable people. By means of providing social services and transfers in kind, Local Governmental Agencies (LGAs) are able to adapt their service supply to multiple and different needs. The governance of local areas must be concerted and shared creating a virtuous pool of governmental and not governmental actors and agencies.

So the policy makers need to know the situation as it is and the impact of their actions at this local level and also stakeholders and citizens are interested in better understanding the effect of policies on their own territory.

However the main sources of statistical data on monetary and non-monetary poverty are from sample surveys on income and living conditions. These rarely give credible estimates at sub-regional and local level. From this comes the importance of the Small Area Estimation (SAE) methods for measuring poverty at local level. This is confirmed also by the large amount of literature on these local estimates resulting from many projects, conferences and books in the last decade.

This chapter has a twofold scope. It serves as necessary background to introduce the book as it constitutes also a useful preparation to the specific methodologies described in each chapter, and a common reference for the notation to use. We start from the definition of poverty indicators and the problem of their estimation (Section 1.2), to present then the main issues related to the data as data integration and data quality that are cross-cutting the methodologies presented in the book (Section 1.3). Section 1.4 reviews the model-assisted and model-based methods used in the book and also gives advice and recommendations on the previous issues.

1.2 Target Parameters

1.2.1 Definition of the Main Poverty Indicators

In order to monitor the process of social inclusion, a list of 18 indicators monitoring poverty and social exclusion was proposed in 2001 (Atkinson *et al.* 2002). The list is constantly modified and complemented. It contains both indicators based on household incomes (monetary indicators) and indicators based on non-monetary symptoms of poverty (non-monetary indicators). Among poverty indicators, the so-called Laeken indicators are very often used to target poverty and inequalities. They are a core set of statistical indicators on poverty and social exclusion agreed by the European Council in December 2001, in the Brussels suburb of Laeken, Belgium.

Referring to the monetary poverty and starting from the Income distribution the most frequently used indicators are the average mean of the equalized income, the Head Count Ratio (HCR) and the Poverty Gap (PG). The HCR measures the incidence of poverty and it is the percentage of individuals of households under a poverty line, that can be defined at national or regional level. For example, the European Commission fix it as 60% of the median value of the equalized income distribution. The PG index measures the intensity of poverty, that are the depth of poverty by considering how far, on average, the poor are from that poverty line.

Formally, the incidence of poverty or HCR and the PG can be obtained by the generalized measures of poverty introduced by 1984. Denoting the poverty line by t , the Foster-Greer-Thorbecke (FGT) poverty measures are defined as:

$$F(\alpha, t) = \frac{1}{N} \sum_{j=1}^N \left(\frac{t - y_j}{t} \right)^\alpha \mathbf{I}(y_j \leq t). \quad (1.1)$$

Here y is a measure of income for individual/household j , N is the number of individuals/households and α is a “sensitivity” parameter. Setting $\alpha = 0$ defines the HCR, $F(0, t)$, whereas setting $\alpha = 1$ defines the PG, $F(1, t)$.

The HCR indicator is a widely used measure of poverty. The popularity of this indicator is due to its ease of construction and interpretation, even if it has some limitations. As it assumes that all poor individuals/households are in the same situation, the easiest way of reducing its value is by implementing actions to target benefits to people who are just below the poverty

line. In fact, they are the ones who are the cheapest to move across the line. Hence, policies based on the headcount index might be not completely effective, as they are not based on the exam of the whole income distribution. For this reason, estimates of the PG indicator are important. The PG can be interpreted as the average shortfall of poor people. It shows how much would have to be transferred to all the poor to bring their expenditure up to the poverty line.

Together with the above indicators, the average value of the distribution of the household income is also important. This is especially true when the level of income is modest and the distribution of income has a long tail. In this case the median value on which the poverty line is computed is expected to be low and the HCR tends to be low as well. Also the PG can lose its relevance, giving a misleading indication of the deprivation of the population under study.

In many cases these measures are considered as a starting point for more in depth studies of poverty and living conditions. In fact, analyses are done using also non-monetary indicators in order to give a more complete picture of poverty and deprivation (Cheli and Lemmi, 1995). In addition, as poverty is a question of graduation, the set of indicators is generally enlarged with other indicators belonging to vulnerable groups, from which it can be likely to move towards the status of poverty (see Chapter 2 of this book). The spatial distribution of these poverty indicators is a feature of high interest. It can be illustrated and represented by building poverty maps. Poverty maps can be constructed using censuses, surveys, administrative data and other data. Here we refer to poverty mapping to visualize the spatial distribution of poverty indicators. This is particularly useful, as it is shown in Chapter 2, to monitor the localization of poverty and the individuation of the most vulnerable areas.

1.2.2 Direct and Indirect Estimate of Poverty Indicators at Small Area Level

The estimates of the different poverty indicators at area level can be done under the design-based (Hansen *et al.* 1953; Kish 1965; Cochran 1977), model-assisted (Särndal *et al.* 1992) and model based approach (Gosh and Meeden, 1997, Valliant *et al.* 2000; Rao 2003), as direct or indirect small area estimates. The direct estimates are produced under the design-based approach using only data coming from one survey, the indirect estimates use auxiliary information (variables) to improve the quality and accuracy of survey estimates or to break down the known values referred to larger areas by using regression-type models. All these estimates belong to the broad class of Small Area Estimation (SAE) methods.

Let us start introducing the notation we use in this chapter and in particular in the review of the small areas model-assisted and model-based methods. Consider that a population U of size N is divided into D non-overlapping subsets U_d (domains of study or areas) of size N_d , $d = 1, \dots, D$. We index the population units by j and the small areas by d , the variable of interest is y_{jd} , $\mathbf{x}_{jd}$ is a vector of p auxiliary variables. We assume that $\mathbf{x}_{ij}$ contains 1 as its first component. Suppose that a sample s is drawn according to some, possibly complex, sampling design such that the inclusion probability of unit j within area d is given by π_{jd} , and that area-specific samples $s_d \subset U_d$ of size $n_d \geq 0$ are available for each area. Note that non-sample areas have $n_d = 0$, in which case s_d is the empty set. The set $r_d \subseteq U_d$ contains the $N_d - n_d$ indices of the non-sampled units in small area d .

Values of y_{jd} are known only for sampled values while for the p -vector of auxiliary variables it is assumed that area level totals $\mathbf{X}_d$ or means $\bar{\mathbf{X}}_d$ or individual values $\mathbf{x}_{jd}$ are accurately known from external sources.

The straightforward approach to calculate FGT poverty indicators referring to the areas of interest is to compute direct estimates. For each area, direct estimators use only the data referring to the sampled households, since for these households the information on the household income is available.

The direct estimators of the FGT poverty indicators are of the form:

$$F_d^{dir}(\alpha, t) = \frac{1}{\sum_{i \in s_d} w_{jd}} \sum_{j \in s_d} w_{jd} \left(\frac{t - y_{jd}}{t} \right)^\alpha I(y_{jd} < t), \quad d = 1, \dots, D, \quad (1.2)$$

where w_{jd} is the sampling weight (inverse of the probability of inclusion) of household j belonging to area d and $\sum_{i \in s_d} w_{jd} = N_d$. In the same way, the mean of the household equivalized income in each small area can be computed as:

$$m_d^{dir} = \frac{1}{\sum_{i \in s_d} w_{jd}} \sum_{j \in s_d} w_{jd} y_{jd}, \quad d = 1, \dots, D. \quad (1.3)$$

When the sample size in the areas of interest is limited, estimators such as (1.2) and (1.3) cannot be used. In fact the size is too small to obtain acceptable statistical significance of the direct estimates obtained under the sample design. Then the purely design-based solution and the usage of direct estimates often implies the increase of the sample size, oversampling of the studied domains. If oversampling is done, credible estimates can be obtained with appropriate direct estimators and the SAE problem is solved. Nevertheless, in many practical situations oversampling is far from being an option as cost-benefit analysis excludes it as a time-consuming and unaffordable solution.

In these cases, model-assisted and model-based SAE techniques need to be employed. Therefore, the estimation of poverty indicators (target parameters) at local level is computed with indirect methods by using auxiliary variables, usually coming from administrative data available also at local area level. The relationship between the target parameters and the auxiliary variables is described by a suitable model. Considering Särndal *et al.* (1992) we clarify that in this context a model consists of “some assumptions of relationship, unverifiable but not entirely out of place, to save survey resources or to bypass other practical difficulties”.

Under these approaches it is useful to express the mean and the FGT indicators for the small area d as shown in the following.

The population small area mean can be written as:

$$m_d = N_d^{-1} \left(\sum_{j \in s_d} y_{jd} + \sum_{j \in r_d} y_{jd} \right). \quad (1.4)$$

Since the y values for the r_d non-sampled units are unknown, they need to be predicted.

The FGT poverty indicators in small area d can be written as:

$$F_d(\alpha, t) = N_d^{-1} \left(\sum_{j \in s_d} z_{jd}(\alpha, t) + \sum_{j \in r_d} z_{jd}(\alpha, t) \right), \quad (1.5)$$

where

$$z_{jd}(\alpha, t) = \left(\frac{t - y_{jd}}{t} \right)^\alpha I(y_{jd} < t). \quad (1.6)$$

Also the z values for the r_d non-sampled units are unknown, and they need to be predicted on the basis of the predicted y values.

The prediction of the y is generally based on a set of auxiliary variables following a regression model. In this perspective, the model-based methodologies allow for the construction of efficient estimators and their confidence intervals by borrowing the strength through use of a suitable model.

The prediction process can encounter inadequacies, difficulties, and problems due both to the characteristics of the available data and the specification and fitting of the SAE model. These issues depend on the amount and the extent of the information on the study variable and on the auxiliary information, and on the typology of the study variable we are interested in. Other problems are linked to the specification of the model as the under/over shrinkage effect of the variability of the estimates between the areas, the modeling of the spatial relationships among the areas and/or the units and the treatment of out-of-sample areas (see Section 1.3).

1.3 Data-related and Estimation-related Problems for the Estimation of Poverty Indicators

The data-related problems are faced when preparing the data information available to set up the estimation phase.

There are various sample surveys, both at EU and country level, on household income, consumption, labor force and living conditions that can be used to compute direct estimates of poverty and related indicators. However, these surveys have at least two limitations: (i) problems of incoherent definitions may rise, because no single data source is able to cover all the aspects; and (ii) the estimates are accurate only at the level of large areas, because the sample is sized at regional level (e.g., in Italy not at province and municipality level).

To overcome the first limitation, it is necessary to check the coherence among the different definitions of the target variables and to improve their comparability, as well as to integrate the micro data coming from different surveys and other data sources to increase the accuracy of the direct estimations.

The second limitation means that the survey data do not support reliable estimation at the level of a local area because sample sizes are often too small to provide direct estimates with acceptable variability (as measured by the coefficient of variation). Sometimes, these estimates could be obtained with larger samples, oversampling the areas of interest, but increasing also the survey costs, and this is not a generally feasible solution to the problem.

When the administrative register data are used as covariate in the SAE model, it is frequently necessary to integrate data coming from different administrative sources in order to derive more adequate auxiliary variables and more accurate and complete final statistics. This is not a straightforward procedure, as it is shown in Chapters 3 and 4 of this book. The keyword is the harmonization of the registers in such a way that information from different sources and observed data should be consistent and coherent.

Other data-related problems arise when indirect methods based on sample surveys are used:

- (i) *The out-of-sample areas.* The estimation of target parameters at local area use both the data collected by the related survey and the auxiliary variables data available at that area level. Frequently, for some or many areas the values of the study variable are not available,

and obviously the SAE have to face with this situation, that is known as the problem of out-of-sample areas or domains.

- (ii) *The benchmarking.* Often the target parameters to be estimated at area level are to be related with known values referred to larger areas we want to break down with the estimation models. Once obtained, the small area estimates should be consistent with already known values for larger areas. Benchmarking is the consistency of a collection of small area estimates with a reliable estimate obtained according to ordinary design-based methods for the union of the areas. The population counts or the values of the target parameters in larger areas serve as a benchmark accounting for under coverage or over coverage and underreporting of the small area target values. Realignment of the small area estimates with the known values is an automatic result of the application of some small area methods. This is also particularly important for National Statistical Institutes to ensure coherence between small area estimates and direct estimates produced at higher level planned domains. In Section 1.4 we examine the methods from this perspective giving advice and warnings about their features and impact on the estimates, guiding the reader to other chapters of the book.
- (iii) *The excess of zero values.* The excess of significant zero values in the data requires a preliminary investigation to formulate a model of behavior for the study variable in the population. There are many practical situations where the study variable can be conceptualized as skewed and strictly positive: in a population of individuals income and consumption follow those models. The problem of the zero excess emerges in situations where the target variable is not only skewed and strictly positive, but defined over the whole positive axis, zero included. Also, when analyzing significant variables to build up poverty indicators it is likely to be in the presence of survey data where there are many zero values of that variable for many sampled households. We refer here to the case of negative income values that are substituted by zero values. A high frequency of zeros can occur also when the study variable is a characteristic of the households, such as presence of households not able to keep their home adequately warm or with arrears on utility bills in a local area where living conditions are acceptable. In this case the problem is different and should be treated under the umbrella of SAE for a rare population.
- (iv) *The outlier.* Outlier detection in the study variable have always been an interesting challenge when examining data to prepare the estimation of small area target parameters. If they are significant and not to be eliminated cleaning up the data set, they require methods that are robust against their effect on the validity of the small area model.

There are solutions described in recent literature to deal with the problem of excess of zeros and with the estimation in the presence of outliers which we will mention in Section 1.4 and they also are presented in the following chapters.

Part III of this book contains chapters devoted to the design-based estimation of poverty indicators and on related themes. Particularly Chapter 5 provides evidence on the effect of the sample design on SAE methods. Chapter 6 shows applications of the design-based framework to SAE and Chapter 7 illustrates the cumulation of panel data to estimate the sampling variance.

The estimation-related problems are inherent to the selected SAE model and its specification and fitting procedure. They produce an effect on the set of small area estimates affecting their heterogeneity and the meaning of their relation with other variables:

- (v) *The shrinkage effect.* The SAE estimates can often be motivated from both a Bayesian and a frequentist point of view, can be obtained using the theory of best linear unbiased prediction (BLUP) or empirical best linear unbiased prediction (EBLUP) or under non-parametric and semi-parametric approaches using also M-quantile models. The chapters of Part III and Part V of this book show many of these models and present simulation studies and application to real poverty data. Nevertheless, there are situations where the models have the tendency for under/over-shrinkage of small area estimators. In fact, it is often the case that, if we consider a collection of small area estimates, they misrepresent the variability of the underlying “ensemble” of population parameters. In other words, the expected sampling variance of the set of predictions is less than the expected sampling variance of the ensemble of the true Small Area parameters (see Rao, 2003, section 9.6 for a discussion of this problem and also of adjusted predictors).
- (vi) *The spatial modeling.* In recent years there have been significant developments in model-based small area methods that incorporate spatial information in an attempt to improve the efficiency of small area estimates by borrowing strength over space. The possible gains from modeling the correlations among small area random effects used to represent the unexplained variation of the small area target quantities are examined and compared with other parametric and non parametric approaches. The reader can find a review of spatio-temporal models in the chapters of Part IV. In Chapters 11, 12 and 13 there are examples of how these spatial models perform when estimation is for out-of-sample areas that is areas with zero sample, and issues related to estimation of mean squared error (MSE) of the resulting small area estimators are discussed. The emphasis is on point prediction of the target area quantities, and mean square error assessments. However, these alternative small area models using data with geographical information have to be studied also with reference to their performance whenever the Modifiable Area Unit Problem (MAUP) occurs.
- (vii) *The Modifiable Area Unit Problem.* The MAUP appears when analyzing the relation (spatial or not) between variables. It is a potential source of error that can affect spatial studies, which utilize aggregate data sources and also the SAE results. The result can be diverse when the same relation is measured on different areal units. This can give misleading results in the specification of SAE models and affect the quality of the small area estimates. A simple strategy to deal with the problem of MAUP in SAE is to undertake analysis at multiple scales or zones. In Section 1.4 we will indicate some preliminary results on the scale effect of MAUP when obtaining small area estimates.

1.4 Model-assisted and Model-based Methods Used for the Estimation of Poverty Indicators: a Short Review

1.4.1 Model-assisted Methods

In the last 30 years mixture modes of making inference have become common in survey sampling; in many cases design-based inference is model assisted. Also in the SAE context the model-assisted approach has become popular and in this section we briefly review the most common estimators under this approach.

Among design-based methods assisted by the specification of a model for the study variable there are three families of methods that have been recently applied in poverty mapping: Generalized Regression (GREG) estimators; pseudo-EBLUP estimators; and M-quantile weighted estimators.

The GREG approach can be used to estimate several poverty indicators. With reference to the estimation of the small area mean, the estimators under this approach share the following structure:

$$\hat{m}_d^{GREG} = \sum_{j \in U_d} \hat{y}_{jd} + \sum_{j \in s_d} w_{jd}(y_{jd} - \hat{y}_{jd}), \quad (1.7)$$

where w_{jd} is the sampling weight of unit j within area d that is the reciprocal of the respective inclusion probability π_{jd} . Different GREG estimators are obtained in association with different models specified for assisting estimation, that is for calculating predicted values $\hat{y}_{jd}$, $j \in U_d$. In the simplest case a fixed effects regression model is assumed: $E(y_{jd}) = \mathbf{x}_{jd}^T \boldsymbol{\beta}$, $\forall j \in U_d$, $\forall d$ where the expectation is taken with respect to the assisting model. Lehtonen and Veijanen (1999) introduce an assisting two-level model where $E(y_{jd}) = \mathbf{x}_{jd}^T (\boldsymbol{\beta} + \mathbf{u}_d)$, which is a model with area-specific regression coefficients. In practice, not all coefficients need to be random and models with area-specific intercepts mimicking linear mixed models may be used (Lehtonen *et al.* 2003). In this case the GREG estimator takes the form of (1.7) with $\hat{y}_{jd} = \mathbf{x}_{jd}^T (\hat{\boldsymbol{\beta}} + \hat{\mathbf{u}}_d)$. Estimators $\hat{\boldsymbol{\beta}}$ and $\hat{\mathbf{u}}$ are obtained using generalized least squares and restricted maximum likelihood methods (Lehtonen and Pahkinen, 2004). See Chapter 6 of this book.

Under the pseudo-EBLUP approach the estimators are derived taking into account the sampling design both via the sampling weights and the auxiliary variables in the models. The estimators of the area mean proposed by Prasad and Rao (1999) and You and Rao (2002) are based on the assumption of a population nested error regression model and it is also assumed that the sampling design is ignorable given the auxiliary variables included in the model. As for the error terms it is assumed that $u_d \stackrel{i.i.d.}{\sim} N(0, \sigma_u^2)$ and $e_{ij} \stackrel{i.i.d.}{\sim} N(0, \sigma_e^2)$.

By combining a Hájek type direct estimator of $\bar{m}_d$ defined as $\bar{y}_{dw} = \sum_{j \in s_d} w_{jd} y_{jd}$ where $w_{jd} = w_{jd} \left(\sum_{j \in s_d} w_{jd} \right)^{-1}$, and the nested error regression model, Prasad and Rao (1999) obtain the following aggregated area level model:

$$\bar{y}_{dw} = \bar{\mathbf{x}}_{dw}^T \boldsymbol{\beta} + v_d + \bar{e}_{dw}, \quad (1.8)$$

with $\bar{e}_{dw} = \sum_{j \in s_d} w_{jd} e_{jd}$ and $\bar{\mathbf{X}}_{dw} = \sum_{j \in s_d} w_{jd} \mathbf{x}_{jd}$.

The design consistent pseudo-EBLUP estimator $\hat{\eta}_{dw}$ of the d th area mean is then given by:

$$\hat{\eta}_{dw} = \hat{\gamma}_{dw} \bar{y}_{dw} + (\bar{\mathbf{X}}_d - \hat{\gamma}_{dw} \bar{\mathbf{x}}_{dw})^T \hat{\boldsymbol{\beta}}_w, \quad (1.9)$$

where $\hat{\gamma}_{dw} = \hat{\sigma}_u^2 (\hat{\sigma}_u^2 + \hat{\sigma}_e^2 \delta_d)^{-1}$, $\delta_d = \sum_{j \in s_d} w_{jd}^2$ and

$$\hat{\boldsymbol{\beta}}_w (\hat{\sigma}_u^2, \hat{\sigma}_e^2) = \left(\sum_{d=1}^D \sum_{j \in s_d} w_{jd} \mathbf{x}_{jd} (\mathbf{x}_{jd} - \hat{\gamma}_{dw} \bar{\mathbf{x}}_{dw}^T) \right)^{-1} \left(\sum_{d=1}^D \sum_{j \in s_d} w_{jd} (\mathbf{x}_{jd} - \hat{\gamma}_{dw} \bar{\mathbf{x}}_{dw}^T) y_{jd} \right). \quad (1.10)$$

The variance components (σ_u^2, σ_e^2) can be estimated using for example, Restricted Maximum Likelihood (REML) or the fitting-of-constants method. Both Prasad and Rao (1999) and You and Rao (2002) provided formulae for the model-based MSE associated with the

pseudo-EBLUP estimators of the area mean. Jiang and Lahiri (2006) noted that these estimators are not second-order correct. Torabi and Rao (2010) derived a second order unbiased predictor for the pseudo-EBLUP estimator (1.9).

An alternative family of model-assisted small area estimators is based on the M-quantile methodology (Chambers and Tzavidis, 2006), see Chapter 9 of this book. Recently, under this model, Fabrizi *et al.* (2014a) proposed a design consistent estimator of area-specific poverty indicators using the Rao–Kovar–Mantel estimator of the distribution function of income F_i (Rao *et al.* 1990) defined as:

$$\hat{F}_d^{WMQ/RKM} = N_d^{-1} \left[\sum_{j \in s_d} w_{jd} I(y_{jd} \leq t) + \sum_{j \in U_d} I(\mathbf{x}_{jd}^T \hat{\beta}_{w\bar{\theta}_d} \leq t) - \sum_{j \in s_d} w_{jd} I(\mathbf{x}_{jd}^T \hat{\beta}_{w\bar{\theta}_d} \leq t) \right], \quad (1.11)$$

where $\hat{\beta}_{wq}$ is a design consistent estimator of β_q . In the application of M-quantile regression to SAE, Chambers and Tzavidis (2006) characterize the variability across the population, beyond what is accounted for by the model covariates, by using the so-called M-quantile coefficients of the population units. For unit j in area d , this coefficient is the value θ_{jd} such that $Q_{\theta_{jd}}(y_{jd} | \mathbf{x}_{jd}) = y_{jd}$, where $Q_q(y_{jd} | \mathbf{x}_{jd})$ is the conditional M-quantile that is assumed to be a linear function of the auxiliary information. The authors observe that if a hierarchical structure does explain part of the variability in the population data, units within areas defined by this hierarchy are expected to have similar M-quantile coefficients. Average area coefficients $\bar{\theta}_d$ may be calculated and this represents an alternative approach to estimating area random effects without the need for using parametric assumptions.

More specifically, the weighted M-quantile-based small area estimator of the mean from (1.11) is:

$$\hat{m}_d^{WMQ} = \int t d\hat{F}_d^{WMQ/RKM}(t) = \frac{1}{N_d} \sum_{j \in s_d} w_{jd} y_{jd} + \left(\frac{1}{N_d} \sum_{j \in U_d} \mathbf{x}_{jd}^T - \frac{1}{N_d} \sum_{j \in s_d} w_{jd} \mathbf{x}_{jd}^T \right) \hat{\beta}_{w\bar{\theta}_d}. \quad (1.12)$$

The M-quantile method can be also used for estimating the HCR and the PG. Using t to denote the poverty line, different poverty indicators are defined by the area-specific mean of the variable derived:

$$f_{jd}(\alpha, t) = \left(\frac{t - y_{jd}}{t} \right)^\alpha I(y_{jd} \leq t), d = 1, \dots, D; \quad j = 1, \dots, N_d. \quad (1.13)$$

The population-level small area-specific poverty indicator can be decomposed as:

$$F_d(\alpha, t) = N_d^{-1} \left[\sum_{j \in s_d} f_{jd}(\alpha, t) + \sum_{j \in r_d} f_{jd}(\alpha, t) \right]. \quad (1.14)$$

The first component in (1.14) is observed in the sample, whereas the second component has to be predicted by using the M-quantile model. Tzavidis *et al.* (2014) propose a non-parametric approach by using a smearing-type estimator. More specifically:

$$F_d(\alpha, t) = N_d^{-1} \left[\sum_{j \in s_d} f_{jd}(\alpha, t) + \sum_{j \in r_d} E(f_{jd}(\alpha, t)) \right]. \quad (1.15)$$

For simplicity let us focus on the simplest case when $\alpha = 0$. An estimator of $F_d(0, t)$ is obtained by substituting an estimator of $E(f_{jd}(\alpha, t))$ in (1.15) leading to

$$\hat{F}_d(0, t) = N_d^{-1} \left[\sum_{j \in s_d} w_{jd} f_{jd}(0, t) + \frac{1}{\sum_{j \in s_d} w_{jd}} \sum_{k \in r_d} \sum_{j \in s_d} w_{jd} I(\mathbf{x}_{kd}^T \hat{\boldsymbol{\beta}}_{w\bar{\theta}_d} + \hat{e}_{jd} \leq t) \right], \quad (1.16)$$

where $\hat{e}_{jd}$ s are the estimated residuals from the M-quantile fit. The same approach can be followed to estimate $\hat{F}_d(1, t)$ or any other of the FGT poverty measures.

For the estimation of the variance of the M-quantile (MQ) predictors see Fabrizi *et al.* (2014a) where two alternative estimators of the variance of the MQ predictors are proposed.

Even if the use of design consistent estimators in SAE is somewhat questionable because of the small sample sizes in some or all of the areas, as Pfeffermann noted (Pfeffermann, 2013), the families of methods we have described above offer generally design consistent estimators.

The three approaches previously described give partial solutions to the problems listed in Section 1.3: they give practical solutions to benchmarking, they deal with the presence of outliers, the estimates that they provide are differently affected by the shrinkage effect, and they all offer out-of-sample predictions.

Also to protect against possible model failures, *benchmarking* procedures make the total of small area estimates match a design consistent estimate for a larger area. With respect to benchmarking, all the families of methods offer a solution.

There are two kinds of benchmarked estimators: estimators that are internally benchmarked (or self-benchmarked) and those that are externally benchmarked. Self-benchmarked predictors are the GREG estimator and the pseudo-EBLUP introduced by You and Rao (2002). The externally benchmarked ones are more common under the model-based approach. For a recent review see Wang *et al.* (2008).

The GREG procedure uses the higher level totals as auxiliary data in calculating survey weights, thereby adjusting the lower level weights so that the total and subtotal estimates are consistent (see also Smith and Hidioglou, 2005). In addition, the weights that are used for direct estimation using survey data in GREG expression are often constructed using calibration methods. Often benchmarking to auxiliary totals is used together with weight equalization. Benchmarking (forcing certain estimates to match known totals) has been shown to reduce variances for statistics correlated with the auxiliary characteristics, and weight equalization (forcing the weights within higher-level units to be equal) has been shown to further reduce variances for statistics measured on the higher-level units (Lehtonen and Veijanen, 1999). The pseudo-EBLUP estimators satisfy the benchmarking property without any adjustment in the sense that they add up to the direct survey regression estimator when aggregated over the areas. A drawback of this type of self-benchmarked estimators is that they force the use of the same auxiliary information used for the direct usually GREG-type estimator also for the model-based small area predictors, whereas it could be very profitable to allow for different auxiliary variables at the small area level. Coming to the M-quantile approach note that expression (1.12) has a GREG-type form. This is the basis to see that the MQ predictors do not satisfy the benchmarking property as it is shown in Fabrizi *et al.* (2014b). Here the authors propose a method of constraining M-quantile regression. It can be applied to obtain benchmarking MQ small area estimates.

The treatment of the *outliers* is not the focus of the estimators of GREG type nor of those under the pseudo-EBLUP approach, while the weighted M-quantile approach this issue.

There are studies under the AMELI (2008) project that illustrate the behavior of the GREG-like estimators in the presence of different models of outlier-contamination of the observed data. The results show that even if a robust method of fitting the logistic mixed model was not available, the poverty rate estimators are fairly robust: this happens both under a simple random sampling design and under a complex sampling design (AMELI, 2008, Deliverable 2.2). To deal with outliers, Beaumont and Alavi (2004) use the weighted generalized M-estimation technique to reduce the influence of units with large weighted population residuals. With respect to the empirical pseudo best approach recalled before there is no contribution addressing the robustification of the estimates against the presence of outliers. Jiang *et al.* (2011) relaxed some of the classical EBLUP model to obtain robust-model based predictors. These relaxations may work also under the pseudo-EBLUP approach but until now no evidence of it has been produced. The AMELI project provides evidence also on the behavior of the Empirical Best Predictor type estimator based on a logistic mixed model. This estimator is least affected by contaminations when the data come from a simple random sample but it is not based on the pseudo-EBLUP approach. As it concerns the M-quantile estimator with respect to GREG-S popular in small area literature (see Rao, 2003, section 2.5), note that: (i) the use of an area-specific coefficient ($\bar{\theta}_d$) in M-quantile regression accounts for area characteristics not explained by the auxiliary variables; and (ii) the use of M-estimation offers outlier robust estimation. Specifically, the recourse to M-quantile regression reduces the impact that outlier observations have on the estimated regression coefficients and thereby on the small area means.

The models which are assisting the estimation under the design-based approach can have the tendency for *under/over-shrinkage* of small area estimators.

The desirable property of neutral shrinkage is not achieved under the pseudo-EBLUP approach. In this case it is reasonable that the over-shrinking behavior of the Empirical Best predictors is confirmed. The understatement of extreme values, referred to as over-shrinkage in this context, is problematic when the goal is the description of the overall distribution among areas. However this tendency can be adjusted (see EURAREA, 2001, section B.3) and it is likely that the adjustment can work even under the pseudo-EBLUP approach, but up to now no evidence of it has been produced.

The tendency of GREG estimators is similar to that of direct estimators and in contrast to that of the over-shrinking empirical Bayes (EB) predictors, as the results of the EURAREA project have shown. The behavior of M-quantile-based predictors is then more similar to that of direct estimators and GREG. Fabrizi *et al.* (2014b) propose an adjustment of the benchmarked MQ predictors in order to obtain estimators with approximately neutral shrinkage. This adjustment parallels the one used to adjust EB predictors (Rao, 2003, see Section 9.6). They extend the methodology of Fabrizi *et al.* (2014b) to obtain estimates that enjoy “ensemble” properties, that is properties related to the estimation of a functional of an ensemble of parameters (Frey and Cressie, 2003). An ensemble of estimators is said to be neutral with respect to shrinkage if the variance of the ensemble of the parameters can be unbiasedly estimated by the variance of the ensemble of the estimators. This guarantees a correct representation of the geographical variation of the variable in question. Otherwise, this geographical variation may be over- or underestimated. Neutral shrinkage is important when small area estimators are used to create “maps”.

For the set $E = \{d | n_d = 0\}$ of the *out-of-sample areas*, that is areas where $n_d = 0$, the GREG-like estimators cannot be computed. The pseudo-EBLUP approach provides predictors under the specified models which are likely to underestimate the variability of the estimates among areas. Consistently with Chambers and Tzavidis (2006), the small area estimator $\hat{m}_d^{WMQ}$ can be defined as $N_d^{-1} \sum_{j \in U_d} \mathbf{x}_{jd}^T \hat{\boldsymbol{\beta}}_{w0.5}$, that is a synthetic estimator based on the weighted M-regression.

1.4.2 Model-based Methods

The most popular method used for model-based SAE employs linear mixed models. In the general case such a model has the form:

$$y_{jd} = \mathbf{x}_{jd}^T \boldsymbol{\beta} + u_d + e_{jd}, \quad (1.17)$$

where u_d is the area-specific random effect and e_{jd} is an individual random effect. The empirical best linear unbiased predictor (EBLUP) of m_d (Henderson, 1975, Rao, 2003, chapter 7) is then

$$\hat{m}_d^{LM} = N_d^{-1} \left[\sum_{j \in s_d} y_{jd} + \sum_{j \in r_d} \{\mathbf{x}_{jd}^T \hat{\boldsymbol{\beta}} + \hat{u}_d\} \right], \quad (1.18)$$

where $\hat{\boldsymbol{\beta}}$, $\hat{u}_d$ are defined by substituting an optimal estimator for the covariance matrix of the random effects in (1.17) in the best linear unbiased estimator of $\boldsymbol{\beta}$ and the BLUP of u_d , respectively. A widely used estimator of the MSE of the EBLUP is based on the approach of Prasad and Rao (1990). This estimator accounts for the variability due to the estimation of the random effects, regression parameters, and variance components.

Models presented in Parts IV and V of this book rely on and often enlarge the assumptions of this popular approach: Chapter 8 introduces the issue of measurement error in the covariates; Chapter 10 extends it to a non-parametric regression environment; and Chapters 11, 12 and 13 extend it to take into account spatial and temporal correlations and the characteristics of geographical patterns.

Assuming model (1.17) on the logarithmically transformed values of income y_{jd} , the most widely used method for small area poverty mapping is the so-called World Bank (WB) or Elbers, Lanjouw and Lanjouw (ELL) method (Elbers *et al.* 2003). Chapter 18 describes links, alternatives and models used under this approach. The model is fitted to clustered survey data from the population of interest, with the random effects in the model corresponding to the cluster used in the survey design. Once the model has been estimated using the survey data, the ELL method uses the following bootstrap population model to generate L synthetic censuses:

$$y_{jd}^* = \mathbf{x}_{jd}^T \hat{\boldsymbol{\beta}} + u_d^* + e_{jd}^*, u_d^* \sim N(0, \hat{\sigma}_u^2), e_{jd}^* \sim N(0, \hat{\sigma}_e^2) \quad (1.19)$$

For each draw, using the synthetic values of the welfare variable y_{jd}^* , values of the poverty indicators of interest for the different small areas are calculated. These are averaged over the L Monte Carlo simulations to produce the final estimates of the poverty quantities, with the simulation variability of these estimates used as an estimate of their uncertainty.

Molina and Rao (2010) point out that when small areas and clusters coincide, in the simplest case of estimating a small area mean, the ELL method leads to a synthetic regression

estimator that, in many cases, could be less efficient than the alternative model-based estimators. Molina and Rao (2010) propose a modification of the ELL method (the empirical best predictor (EBP) method) introducing random area effects (rather than random cluster effects) into the linear regression model for the welfare variable, and also simulated out-of-sample data by independent drawings the conditional distribution of the out-of-sample data, given the sample data. Deeper insights on this and recent enhancements of the method are described in Chapter 17.

An alternative approach to EBLUP has been discussed in Chandra and Chambers (2005) and it is based on the use of model-based direct estimation (MBDE) within the small areas. In this case an estimate for a small area of interest corresponds to a weighted linear combination of the sample data for that area, with weights based on a population level version of the linear mixed model. These weights “borrow strength” via this model, which includes random area effects. Provided the assumed small area model is true, the EBLUP is asymptotically the most efficient estimator for a particular small area. In practice however the “true” model for the data is unknown and the EBLUP can be inefficient under misspecification. In such circumstances, Chandra and Chambers (2005) note that MBDE offers an alternative to potentially unstable EBLUP. In particular, MBDE is easy to implement, produces sensible estimates when the sample data exhibit patterns of variability that are inconsistent with the assumed model (e.g., contain too many zeros) and generates robust MSE estimates. The MBDE is presented in Chapter 14 of this book for the estimation of the Cumulative Distribution Function.

A different approach has been proposed in the literature for further robustification of the inference by relaxing some of the model assumptions. This approach is based on M-quantile regression (Breckling and Chambers, 1988). It provides a “quantile-like” generalization of regression based on influence functions (Breckling and Chambers, 1988). A linear M-quantile regression model is one where the q th M-quantile $Q_q(y_{jd}|\mathbf{x}_{jd})$ of the conditional distribution of y given x satisfies:

$$Q_q(y_{jd}|\mathbf{x}_{jd}) = \mathbf{x}_{jd}^T \boldsymbol{\beta}_q. \quad (1.20)$$

That is, it allows a different set of regression parameters for each value of q . For specified q and continuous influence function ψ , an estimate $\hat{\boldsymbol{\beta}}_q$ of $\boldsymbol{\beta}_q$ can be obtained via an iterative weighted least squares algorithm.

As stated in the previous section, extending this line of thinking to SAE, Chambers and Tzavidis (2006) observed that if variability between the small areas is a significant part of the overall variability of the population data, then units from the same small area are expected to have similar M-quantile coefficients. In particular, when (1.20) holds, and $\boldsymbol{\beta}_q$ is a sufficiently smooth function of q , these authors suggest a predictor of m_j of the form:

$$\hat{m}_d^{MQ} = N_d^{-1} \left[\sum_{j \in s_d} y_{jd} + \sum_{j \in r_d} \hat{Q}_{\bar{\theta}_d}(y_{jd}|\mathbf{x}_{jd}) \right], \quad (1.21)$$

where $\hat{Q}_{\bar{\theta}_d}(y_{jd}|\mathbf{x}_{jd}) = \mathbf{x}_{jd}^T \hat{\boldsymbol{\beta}}_{\bar{\theta}_d}$ and $\bar{\theta}_d$ is an estimate of the average value of the M-quantile coefficients of the units in area d . Typically this is the average of estimates of these coefficients for sample units in the area. When there is no sample in the area, we can form a “synthetic” M-quantile predictor by setting $\bar{\theta}_d = 0.5$. Tzavidis *et al.* (2010) refer to (1.21) as the “naïve” M-quantile predictor and note that this can be biased and they propose a bias adjusted M-quantile predictor of m_d .

The M-quantile small models are used also for estimating the poverty indicators such as HCR and PG (Tzavidis *et al.* 2014) by using a smearing-type estimator (Duan, 1983). A small area estimator of the HCR is obtained as:

$$\hat{F}_d(0, t) = N_d^{-1} \left[\sum_{j \in s_d} f_{jd}(0, t) + \hat{E}[f_{jd}(0, t)] \right] \quad (1.22)$$

where

$$\hat{E}[f_{jd}(0, t)] = \int I(\mathbf{x}_{jd}^T \hat{\boldsymbol{\beta}}_{\hat{\theta}_d} + \hat{e}_{jd} \leq t) d\hat{F}(\hat{e}) = n^{-1} \sum_{k \in r_d} \sum_{j \in s_d} I(\mathbf{x}_{kd}^T \hat{\boldsymbol{\beta}}_{\hat{\theta}_d} + \hat{e}_{jd} \leq t)$$

with the distribution function estimated as $\hat{F}(\hat{e}) = n^{-1} \sum_{j=1}^n I(\hat{e}_j \leq e)$. The same approach can be used to estimate the PG indicator or any other of the FGT poverty measures.

Under the model-based approach many of the problems listed in Section 1.3 have a solution, for example all of them offer out-of-sample predictors. Among the other issues we focus here on the excess of zero values in the data and in the treatment of geographic information and spatial data.

Model-based estimators usually do not have the *benchmarking* property under a complex sampling design. Given a small area estimator, that does not show the benchmarking property, a first simple way of achieving benchmarking is by a ratio type adjustment. Externally benchmarked predictors are obtained through an a-posteriori adjustment of model-based predictors. Among the others, Pfeiffermann and Barnard (1991) propose an externally restricted benchmarked estimator of small area means. This is constructed under an area linear mixed model for a continuous response variable.

Many variables of interest in economics surveys on poverty and living conditions are semicontinuous in nature, that is they either take a single fixed value (typically 0, zero) or they have a continuous, often skewed, distribution on the positive real line. They present an *excess of zero values*. A semicontinuous variable is quite different from one that has been left censored or truncated, because the zeros are valid self-representing data values, not proxies for negative or missing responses. A two-part random effects model (Olsen and Schafer, 2001) is widely used for SAE with zero-inflated variables, see for example, Pfeiffermann *et al.* (2008) and Chandra and Sud (2012). Chandra and Chambers (2014) propose a SAE method for semicontinuous variables under a two part random effects model. The issues which arise when the data are lognormal are discussed in Chapter 15.

In poverty studies observations that are spatially close may be more alike than observations that are further apart. One approach for incorporating spatial information in *spatial modeling* and in a small area regression model is to assume that the model coefficients themselves vary spatially across the geography of interest and/or the random effects of the model be correlated. Both EBLUP predictors and MQ predictors can be extended to include the effect of the spatial characteristics of the data. These extensions can be applied to poverty studies (see SAMPLE, 2008, deliverables), but are not reviewed in this book.

When geography is included as auxiliary information in modeling, the spatial correlation and the consequent correlation between the random effect in the EBLUP model require the extension of the EBLUP estimator to the Spatial Empirical Best Linear Unbiased Predictor (SEBLUP) estimator (Petrucchi and Salvati, 2006, Pratesi and Salvati, 2009).

Under the MQ approach the reference to the Geographically Weighted Regression (GWR) (Brundson *et al.* 1996) helps in modeling spatial variation. This uses local rather than global parameters in the regression model. That is, a GWR model assumes spatial non-stationarity of the conditional mean of the variable of interest. Salvati *et al.* (2012) propose an M-quantile GWR model, that is a local model for the M-quantiles of the conditional distribution of the outcome variable given the covariates. This approach is semi-parametric in that it attempts to capture spatial variability by allowing model parameters to change with the location of the units, in effect by using a distance metric to introduce spatial non-stationarity into the mean structure of the model. The model is then used to define a predictor of the small area characteristic of interest. As a consequence, it integrates the concepts of bias-robust SAE and borrowing strength over space within a unified modeling framework. By construction, the model is a local model and so can provide more flexibility in SAE, particularly for out-of-sample small area estimation, that is areas where there are no sampled units. For the estimation of the variance of the predictors see Chambers *et al.* (2011, 2014).

When studying the spatial distribution of local poverty indicators obtained by SAE methods, it can be relevant to consider the possible effect of the MAUP. This is a source of statistical bias that can radically affect the results of statistical analysis. It affects results when point-based measures of spatial phenomena (e.g., population density) are aggregated into larger areas. The resulting summary values (e.g., totals, rates, proportions) are influenced by the choice of the boundaries of the areas. For example, point-based census or survey data may be aggregated into census enumeration districts, or post-code areas, or any other spatial partition (thus, the “areal units” are “modifiable”).

The topic has not yet been treated explicitly in the current literature on SAE. The only empirical study is due to Pratesi and Petrucci (2014) who studied the scale effect on SAE predictors by a simulation experiment. They provide evidence to assess the robustness of SAE methods to different scale of aggregation of the point-based measures inside the pre-defined small areas (domains) of interest. The rationale of this simulation study is to verify to what extent we can aggregate the individual values inside the small areas and still have an acceptable accuracy of the estimate of the small area parameter. Under this simulation experiment, methods that are naturally robust to outliers and not linked to distributional assumption on the study variable as M-quantile methods perform better than the alternative methods for SAE and are found to be resilient to changing scale of analysis. This is likely due to the fact that the changes in geography do not affect the M-quantile coefficients at area level.

References

- AMELI: Advanced Methodology for European Laeken Indicators 2008 Project no. SSH-CT-2008-217322. FP7-SSH-2007-1.
- Atkinson AB, Cantillon B, Marlier E, and Nolan B 2002 *Social Indicators: The EU and Social Inclusion*. Oxford: Oxford University Press.
- Beaumont JF and Alavi A 2004 Robust Generalized Regression Estimation. *Survey Methodology* **30**, 195–208.
- Breckling J and Chambers R 1988 M-quantiles. *Biometrika* **79**(4), 761–771.
- Brundson, C, Fotheringham AS, and Charlton M 1996 Geographically weighted regression: a method for exploring spatial nonstationarity. *Geographical Analysis* **28**, 281–298.
- Chambers R, Chandra H, and Tzavidis N 2011 On bias-robust mean squared error estimation for pseudo-linear small area estimators. *Survey. Methodology*, **37**, 153–170.
- Chambers R and Tzavidis N 2006 M-quantile models for Small Area Estimation. *Biometrika*, **93**, 255–268.

- Chambers R, Chandra H, Salvati N, and Tzavidis N 2014 Outlier robust small area estimation. *Journal of the Royal Statistical Society: Series B* **76**(3), 47–69.
- Chandra H and Chambers R 2014 Comparing EBLUP and C-EBLUP for Small Area Estimation. *Biometrical Journal*, doi: 10.1002/bimj.201300233.
- Chandra H and Chambers R 2005 Small area estimation for semicontinuous data. *Statistics in Transition*, **7**, 637–648.
- Chandra H and Sud UC 2012 Small area estimation for zero-inflated data. *Communications in Statistics — Simulation and Computation* **41**(5), 632–643.
- Cheli B and Lemmi A 1995 A totally fuzzy and relative approach to the multidimensional analysis of poverty. *Economic Notes* **24**, 115–134.
- Cochran WG 1977 *Sampling Techniques*, 3rd ed. New York: John Wiley & Sons, Inc.
- Duan N 1983 Smearing estimate: a non parametric retransformation method. *Journal of the American Statistical Association* **78**, 605–610.
- Elbers C, Lanjouw JO, and Lanjouw P 2003 Micro-level estimation of poverty and inequality. *Econometrica* **71**, 355–364.
- EURAREA 2001 Programme funded by Eurostat under the Fifth Framework (FP5) Programme of the European Union.
- Fabrizi E, Giusti C, Salvati N, and Tzavidis N 2014a Mapping average equivalized income using robust small area methods. *Papers in Regional Science* **93**, 685–701.
- Fabrizi E, Salvati N, Pratesi M, and Tzavidis N 2014b Outlier robust model-assisted small area estimation. *Biometrical Journal* **56**, 157–175.
- Foster J, Greer J, and Thorbecke E 1984 A class of decomposable poverty measures. *Econometrica* **52**, 761–766.
- Frey J and Cressie N 2003 Some results on constrained Bayes estimators. *Statistics and Probability Letters* **65**, 389–399.
- Gosh M and Meeden G 1997 *Bayesian Methods for Finite Population Sampling*. London: Chapman & Hall.
- Hansen MH, Hurwitz WN and Madow WG 1953 *Sample Survey Methods and Theory*. New York: John Wiley & Sons, Inc.
- Henderson C 1975 Best linear unbiased estimation and prediction under a selection model. *Biometrics* **31**, 423–447.
- Jiang J and Lahiri P 2006 Estimation of finite population domain means: a model-assisted empirical best prediction approach. *Journal of the American Statistical Association* **101**, 301–311.
- Jiang J, Nguyen T, and Rao JS 2011 Best predictive small area estimation. *Journal of the American Statistical Association* **106**, 732–745.
- Kish L 1965 *Survey Sampling*, New York: John Wiley & Sons, Inc.
- Lehtonen R and Pahkinen E 2004 *Practical Methods for Design and Analysis of Complex Surveys* Chichester, England: John Wiley & Sons, Ltd.
- Lehtonen R, Särndal CE, and Veijanen A 2003 The effect of model choice in estimation for domains, including small domains. *Survey Methodology* **29**, 33–44.
- Lehtonen R and Veijanen A 1999 Domain estimation with logistic generalized regression and related estimators. IASS Satellite Conference on Small Area Estimation. Riga: Latvian Council of Science, 121–128.
- Molina I and Rao JNK 2010 Small area estimation of poverty indicators. *The Canadian Journal of Statistics* **38**, 369–385.
- Olsen MK and Schafer JL 2001 A two-part random-effects model for semicontinuous longitudinal data. *Journal of the American Statistical Association* **96**, 730–745.
- Petrucci A and Salvati N 2006 Small area estimation for spatial correlation in watershed erosion assessment. *Journal of Agricultural, Biological, and Environmental Statistics* **11**, 169–182.
- Pratesi M and Salvati N 2009 Small area estimation in the presence of correlated random area effects. *Journal of Official Statistics* **25**(1), 37–53.
- Pfeffermann D and Barnard CH 1991 Some new estimators for small-area means with application to the assessment of farmland values. *Journal of Business & Economics Statistics* **9**, 73–84.
- Pfeffermann D, Terry B, and Moura FAS 2008 Small area estimation under a two-part random effects model with application to estimation of literacy in developing countries. *Survey Methodology* **34**, 235–249.
- Pfeffermann D 2013 New important developments in small area estimation. *Statistical Science* **28**(1), 40–68.
- Prasad NGN and Rao JNK 1990 The estimation of the mean squared error of small-area estimators. *Journal of the American Statistical Association* **85**, 163–171.
- Prasad NGN and Rao JNK 1999 On robust small area estimation using a simple random effects model. *Survey Methodology* **25**, 67–72.

- Pratesi M and Petrucci A 2014 Methodological and operational solutions to gaps and issues on methods for producing agricultural and rural statistics at small domains level and on methods for aggregation, disaggregation and integration of different kinds of geo-referenced data for increasing the efficiency of agricultural and rural statistics. In *FAO report*.
- Rao JNK 2003 *Small Area Estimation*. Hoboken, New York: John Wiley & Sons, Inc.
- Rao JNK, Kovar JG, and Mantel HJ 1990 On estimating distribution functions and quantiles from survey data using auxiliary information. *Biometrika* **77**, 365–375.
- Salvati N, Tzavidis N, Pratesi M, and Chambers R 2012 Small area estimation via M-quantile geographically weighted regression. *TEST* **21**, 1–28.
- SAMPLE: Small Area Methods for Poverty and Living Condition Estimates 2008 Project no. SSH-CT-2008-217565. FP7-SSH-2007-1.
- Särndal CE, Swensson B, and Wretman J 1992 *Model Assisted Survey Sampling*. New York: Springer-Verlag.
- Smith P and Hidirolou M 2005. Benchmarking through calibration of weights for microdata. *Working Papers and Studies*. Luxembourg: Office for Official Publications of the European Communities.
- Torabi M and Rao JNK 2010 Mean squared error estimators of small area means using survey weights. *The Canadian Journal of Statistics* **38**, 598–608.
- Tzavidis N, Marchetti S, and Chambers R 2010 Robust prediction of small area means and distributions. *Australian and New Zealand Journal of Statistics* **52**, 167–186.
- Tzavidis N, Marchetti S, and Donbavand S 2014 Outlier robust semi-parametric small area methods for poverty estimation. In *Poverty and Social Exclusion, New Methods and Analysis* edited by Gianni Betti and Achille Lemmi, Chapter 15, 283–300. Routledge, Abingdon, Oxon and simultaneously New York.
- Valliant R, Dorfman AH, and Royall RM 2000 *Finite Population Sampling and Inference: A Prediction Approach*. New York: John Wiley & Sons, Inc.
- Wang J, Fuller WA, and Qu Y 2008 Small area estimation under a restriction. *Survey Methodology* **34**, 29–36.
- You Y and Rao JNK 2002. A pseudo-empirical best linear unbiased prediction approach to small area estimation using survey weights. *The Canadian Journal of Statistics* **30**, 431–439.

