

TI 2026-066/V
Tinbergen Institute Discussion Paper

A Bayesian-Inspired Approach to Improving the Precision of Synthetic Panel Estimates of Poverty Dynamics

Gerton Rongen ¹

Peter Lanjouw ²

Chris Elbers ³

¹ Vrije Universiteit Amsterdam

² Vrije Universiteit Amsterdam and Tinbergen Institute

³ Vrije Universiteit Amsterdam

A Bayesian-inspired approach to improving the precision of synthetic panel estimates of poverty dynamics

Gerton Rongen, Chris Elbers, Peter Lanjouw

Abstract

This paper proposes a Bayesian-inspired refinement to the synthetic panel method introduced by Dang et al. (2014), yielding point estimates of poverty transitions and other income mobility metrics. It uses bootstrap samples of household income (and consumption expenditure) correlation coefficients observed in available panel surveys to proxy for correlations over time in countries where no panel surveys are available. This allows more precise estimation of poverty dynamics based on repeated cross-sectional surveys. We present results in the form of a band of two standard deviations (both plus and minus) around the point estimate. A validation analysis using actual Tanzanian panel data shows that, with few exceptions, these bands overlap with the true estimates' confidence intervals. Subsequent application to Malaysian poverty dynamics over the period 2004-2022 illustrates how the method can be employed. The resulting four standard deviation bands are much narrower than intervals based on the original, conservative, synthetic panel bound estimates, in particular for conditional probabilities, such as the chance of escaping poverty. Hence, the results of this approach yield more relevant policy information.

1. Introduction

The study of poverty dynamics is an integral part of poverty analysis. Researchers and policy makers want to know which population groups are least able to escape poverty, or identify which provinces have most chronic poverty. Based on this information, policy makers can take appropriate action, e.g. aimed at overcoming structural barriers driving chronic deprivation, or to institute policies that cushion shocks if poverty is mostly transient. The aim of this paper is to propose a method that, in the absence of true panel data, estimates such poverty transition rates with more precision than existing synthetic approaches.

Typically, estimating poverty dynamics (and measures of general income mobility) requires panel data, which are absent in many countries. In addition, panel surveys usually have smaller samples sizes than cross-sectional surveys, such that panels may not be representative of smaller population groups. They may further be plagued by attrition. The synthetic panel method provides a solution to this information problem. It is based on repeated cross-sectional surveys and, building on a household income model, enables the estimation of upper and lower bounds on poverty transitions and other income mobility statistics (Dang et al., 2014).

Importantly, the model links household income (or consumption expenditures) only to explanatory variables that do not vary over time. Bounds on poverty transitions are then based on two scenarios for the behavior of the regression residuals in this model: a minimum mobility scenario, in which residuals are assumed to be perfectly correlated over time, and a maximum mobility scenario, in which residuals are uncorrelated over time. These assumptions imply correlation coefficients of 1 and 0, respectively. Unfortunately, the bounds derived in this way can be wide, especially for conditional probabilities.

In an effort to overcome this limitation, proposals for sharpening these bounds have been advanced. These include drawing on correlation coefficients from ancillary panel data, and exploiting cohort-level correlations (Dang et al. (2014); Dang & Lanjouw (2023)). However, relevant ancillary data are not always available, and the cohort approach does not yield robust results (Herault & Jenkins (2019); Garces-Urzainqui (2023)).

This paper proposes a refinement to the synthetic panel method that takes the use of ancillary data one step further. Our approach is based on taking a bootstrap sample of documented income and consumption expenditure correlations from actual panel surveys. These empirical correlations span a range that is much narrower than the 0 – 1 interval that underpins the conventional Dang et al. (2014) approach: we observe a range of 0.43 – 0.82 in our sample. As a result, we are able to present point estimates of transitions, with a related robust

measure of the spread of these estimates. Applications to Tanzania and Malaysia demonstrate the validity and relevance of this bootstrap point estimate approach.

This paper thus contributes to the literature by proposing a method to obtain more precise insights into poverty dynamics and broader income mobility. Moreover, we put together a dataset of empirical correlation coefficients that may be of use to other researchers. Finally, the method carries policy relevance because its improved precision means that it can provide a stronger evidence base for policy.

The paper proceeds as follows: Section 2 provides a detailed discussion of our approach and an analysis of our sample of correlation coefficients. Section 3 describes validation efforts, including a comparison with transition outcomes based on actual Tanzanian panel data, while Section 4 discusses an application of the method to Malaysia. Section 5 concludes.

2. Method and sample of correlations

We propose a Bayesian-inspired innovation to the synthetic panel method that allows us to present point estimates of poverty transition quantities. First, we provide an outline of the original synthetic panel approach as devised by Dang et al. (2014).¹ Thereafter, we elaborate on our proposed refinement that makes use of bootstrap sampling. Lastly, we describe our sample of correlation coefficients.

2.1. Basics of the synthetic panel method

Our aim is to estimate joint and conditional probabilities of poverty dynamics. For example, we wish to know the likelihood that an individual is poor in survey round 1 and non-poor in round 2, a joint probability which we can write as follows, denoting per capita income by y_{it} and the poverty line by z_t :

$$P(y_{i1} \leq z_1 \text{ and } y_{i2} > z_2).$$

¹ This subsection draws heavily on the exposition of the method in Rongen and Lanjouw (2024) and Rongen et al. (2024). We mostly refer to income as the measure of welfare, but the method can equally be applied to consumption expenditures.

Alongside evidence on joint probabilities, we are also interested in conditional probabilities. For instance, we want to estimate the likelihood of moving out of poverty, conditional on being poor initially. Such a conditional probability can be written as

$$P(y_{i2} > z_2 | y_{i1} \leq z_1) = \frac{P(y_{i1} \leq z_1 \text{ and } y_{i2} > z_2)}{P(y_{i1} \leq z_1)}$$

and can be computed straightforwardly from the joint probabilities, as the right-hand side of the equation shows. In the absence of panel data, we do not observe incomes in round 1 and round 2 for the same unit i and would normally not be able to estimate these probabilities. The original synthetic panel method, however, draws on cross-sectional data to yield upper and lower bounds on each of the four possible outcomes of the dynamic process, both in conditional and unconditional terms.

The basic idea is to use time-invariant household characteristics x_i to predict household per capita income (or consumption expenditure) for the year in which the actual observation is lacking. The resulting model is:

$$y_{i1} = \beta'_1 x_i + \varepsilon_{i1}$$

and

$$y_{i2} = \beta'_2 x_i + \varepsilon_{i2},$$

where y_i represents the logarithm of household per capita income and the subscripts 1 and 2 indicate the survey round. By definition, the explanatory variables remain unchanged across the two rounds. Mostly, they will be characteristics of the household head, such as the year of birth, educational attainment, ethnicity, religion, and birth district. In addition, if the survey includes retrospective questions, e.g. if the household owned a car or television at the time of the previous survey, then these could also be included. As we have data for each survey round, we estimate the coefficient vector β on these time-invariant explanatory variables separately for each round.

In practical terms, we predict the round 1 income for households in the round 2 survey by multiplying the coefficient vector resulting from the round 1 regression ($\widehat{\beta}_1$) with the characteristics of households in round 2 (x_{i2}). This prediction looks as follows for the observations i of round 2, with the addition of an error term estimate:

$$\widehat{y}_{i1} = \widehat{\beta}'_1 x_{i2} + \widehat{\varepsilon}_{i1}.$$

We turn first to the assumptions underlying this approach and will then discuss how estimates for the error term are obtained.

The method for estimating bounds on poverty transition quantities is predicated on two important assumptions. First, we need to assume that the underlying population from which the data are sampled is the same in both survey rounds. This implies that the method cannot be applied in cases where there are large shocks to the population, for example, if there were massive migration into or out of the country. More precisely, we want to ensure that the income model coefficients that we estimate in round 1 are good predictors for what the round 1 income of the households surveyed in round 2 would have been. Hence, in an effort to establish stability of the reference population, typically only households whose heads are between 25 and 60 years in survey round 1 are included.²

Second, we need to make assumptions about the relationship between the error terms ε_{i1} and ε_{i2} . Specifically, we assume that these are, on average, positively correlated over time, i.e. that their correlation coefficient ρ_ε is positive. In the presence of household-specific effects and persistence of shocks, this does not seem to be an unreasonable assumption, as is also confirmed empirically by Dang et al. (2014). No further assumptions about the shape of the distribution of the error terms are made for these estimates; in that sense, the resulting bounds on poverty transitions are non-parametric. The upper- and lower-bound estimates for mobility derive from two extreme cases of the assumed relationship between the error terms.³

At one extreme, we consider the case of $\rho_\varepsilon = 0$, which corresponds to the upper-bound mobility estimate. In this scenario, the errors are uncorrelated and mobility, given household characteristics, will be at a maximum. For example, households are more likely to experience a large positive shock in one round, and a large negative one in another, such that both movements out of poverty and into it are more probable. Because we assume that errors in fact are positively correlated, on average the upper bound should be strictly larger than actual mobility. We implement this scenario by randomly drawing a residual from the actual distribution of round 1 residuals. This random residual is then added to the fitted income to yield predicted household income in round 1 and corresponding poverty status. Aggregating these quantities at the population level, we obtain simulated country-level poverty dynamics.⁴ Due to the random nature

² This age range is chosen so that formation of new households and dissolution of existing households should be at a minimum. Moreover, it ensures that in round 1 heads will generally have obtained their maximum level of education. The optimal age range may vary from on context to another.

³ Note that upper bounds on mobility are the mirror image of lower bounds on immobile categories, i.e., remaining poor or remaining non-poor, and vice versa.

⁴ Transition probabilities for each household are population-weighted when we aggregate, so that the country-level figures represent population proportions.

of this process, these steps are repeated R times, with $R = 100$ in our applications, and averaged to obtain a robust estimate of the upper bound of mobility.

At the other extreme, our lower-bound estimates assume a perfect positive correlation between the errors: $\rho_\varepsilon = 1$. In this case, average mobility will be at its minimum. For the estimation, this means we simply take the estimated round 2 residual, rescale it to match the variance of the round 1 residual standard error, and add it to the fitted round 1 income.⁵ Analogous to the upper-bound scenario, but without the repetition, we obtain the lower-bound population-level poverty dynamics estimates. Conditional probabilities are subsequently computed from the population- or subgroup-level joint probability estimates.

2.2. Bootstrap point estimates

We do not know the actual degree of serial household income correlation in our target country, but panel studies for a broad range of countries permit computation of such correlations. This range of empirical serial correlations is much narrower than the 0 - 1 range that the bound estimates are based on.⁶ Hence, we assume that the true values for our target country will be somewhere within this empirically observed range. Our proposition, then, is that a bootstrap sample of these empirical income correlations, coupled with an additional assumption that errors follow a bivariate normal distribution, yields a meaningful point estimate for transition quantities, and plausible standard deviations. Obviously, we do not know what the actual correlations are in each interval, so we prefer to focus on a bandwidth of four standard deviations (4SD band) around the point estimate (two plus and two minus), rather than the actual estimate itself. Below, we will show that this provides considerably more precise estimates compared to the original bound estimates.

This approach is inspired by Bayesian statistics. Our assertion that the correlation coefficient has a distribution over an interval narrower than 0 to 1 can be characterized as a flat, or diffuse, prior. In that case, the posterior distribution equals the empirical distribution function. Hence, taking a bootstrap sample from the empirical distribution is equivalent to sampling from the posterior distribution.

The calculation of our point estimates is based on a method already put forward by Dang et al. (2014). Its precision comes at the cost of introducing additional assumptions, however.

⁵ We use a scaling factor $\gamma = \hat{\sigma}_{\varepsilon 1} / \hat{\sigma}_{\varepsilon 2}$ to adjust for the difference in error variance between the two rounds ($\hat{\sigma}_\varepsilon$ is the standard error of the residuals).

⁶ Note that we make a jump here from model error correlations to income correlations. We outline below how they are connected.

Notably, we will assume that the errors of our income model, ε_{i1} and ε_{i2} , follow a bivariate standard normal distribution characterized by the correlation coefficient ρ_ε .⁷ We impute hypothesized values for the correlation of the error terms (ρ_ε) by drawing on income correlation coefficients (ρ_y) from actual panel data in other countries. Dang and Lanjouw (2023) provide the following formula for linking the serial correlation of income to the serial correlation of model residuals:

$$\rho_\varepsilon = \frac{\rho_y \sqrt{\text{Var}(y_1) \text{Var}(y_2)} - \beta_1' \text{Var}(x) \beta_2}{\sqrt{\text{Var}(\varepsilon_1) \text{Var}(\varepsilon_2)}}$$

In our application, the equation combines expressions of the variation in our cross-sectional data with an income correlation from our bootstrap sample to arrive at an estimate of ρ_ε .⁸

We start implementation by compiling a database of empirical correlation coefficients (see the section below for a description). Correlations are standardized to a two-year reference interval, since the correlations reported cover intervals of different duration. To that end, we assume correlations follow a logarithmic trend over time⁹:

$$\rho_{y,p} = \beta_0 + \beta_1 \log \delta_p$$

where δ_p indicates the interval between survey waves in years for panel p . We expect correlations to weaken over time, yielding a negative β_1 . Then, correlations over shorter or longer intervals can be standardized to a two-year interval as follows:

$$\rho_{y,p}^S = \rho_{y,p} + \widehat{\beta}_1 \log \frac{2}{\delta_p}$$

We implement this standardization separately for income and consumption correlations.

Subsequently, we draw, with replacement, a bootstrap sample of size 100 from these standardized correlations.¹⁰ For each draw, we estimate transition probabilities as follows. First,

⁷ We test this assumption in the validation section. Results are satisfactory.

⁸ Combining data from different, unrelated sources may in some cases lead to a value of ρ_ε that contradicts our own assumptions, as we will describe further down.

⁹ Colgan (2026) concludes that this approach yields the best point estimates, based on a comparison of extrapolation approaches using New Zealand data compiled by Ball (2016). A sensitivity analysis using our sample of correlation coefficients leads to a similar conclusion; see the Appendix for a detailed description of results.

¹⁰ As an alternative approach, we used maximum likelihood estimation to fit beta and normal distributions to the empirical data, from which we were then able to draw a random sample of coefficients. However, the

we fit our household income model and estimate the resulting residual correlation $\widehat{\rho}_\varepsilon$ by applying the formula provided by Dang & Lanjouw (2023). Then, we employ the bivariate standard normal cumulative distribution function $\Phi_2(\cdot)$ to obtain estimates for all household poverty transition probabilities, e.g. as follows for the poor-to-poor transition¹¹:

$$\hat{P}(y_{i1} < z_1 \text{ and } y_{i2} < z_2) = \Phi_2\left(\frac{z_1 - \widehat{\beta}_1'x_{i2}}{\widehat{\sigma}_{\varepsilon_1}}, \frac{z_2 - \widehat{\beta}_2'x_{i2}}{\widehat{\sigma}_{\varepsilon_2}}, \widehat{\rho}_\varepsilon\right)$$

For each draw, household level results are aggregated up to population and subgroup level. Lastly, averaging over all 100 draws yields the bootstrap mean probabilities, i.e. our point estimates, and allows us to compute standard deviations over the 100 draws.¹²

This approach allows for the analysis of a variety of intragenerational mobility questions. This paper focuses mostly on poverty dynamics. However, we do also study a broader concept of economic mobility, where we divide the population into three states: poverty, vulnerability and economic security. In this categorisation, those classified as vulnerable face an elevated chance of falling into poverty in the next period (Dang & Lanjouw, 2017). Upward mobility is then defined as moving up one or two categories, downward mobility as moving down one or two categories. Nonetheless, a multitude of other analyses is feasible by adapting the categorizations, such as analysis of relative income mobility by studying movements between quintiles.

2.3. Empirical correlation coefficients

Panel data offer the opportunity to estimate actual serial correlation coefficients of household per capita income or consumption expenditures. Colgan (2023, 2026) and Dang and Lanjouw (2023) report such estimates for fourteen countries in the EU-SILC database and Bosnia-Herzegovina, Lao PDR, Peru, the United States and Vietnam. In addition, we obtained correlations from the harmonized World Bank LSMS-ISA database for five countries in Sub-Saharan Africa (Bentze & Wollburg, 2024). Table 1 provides summary statistics on this sample of correlations, while Table 3 in the Annex gives the full list.¹³ Our sample consists of 45 correlations,

draws from these theoretical distributions were not sufficiently similar to the empirical distribution, such that we preferred to stick with a bootstrap sample from the latter.

¹¹ We refer to Equations 20-23 in Dang et al. (2014) for all transition formulae.

¹² In the application to Tanzania, 100 draws are sufficient to obtain a stable estimate of the standard deviation.

¹³ The full correlations dataset and a template Stata do-file to estimate bootstrap point estimates are available at the website of the Data and Evidence to End Extreme Poverty (DEEP) program, at <https://doi.org/10.55158/DEEPWP41>.

for 24 different countries. For eleven countries, we have standardized correlations for at least two non-overlapping time periods.¹⁴ Figure 1 visualizes this collection of correlations, indicating source study and differentiating between income and consumption correlations. In addition, Figure 2 plots kernel density estimates of the correlations, by welfare concept and overall, while Figure 3 compares these kernel densities with fitted normal distributions. A first observation is that the sample of income correlations is more concentrated and has a higher mean than the consumption sample. A second observation is that the latter sample seems close to following a normal distribution.

Table 1 Sample of standardized panel correlation coefficients - summary statistics

Unweighted							
	mean	median	sd	min	max	N	Countries
Consumption	0.644	0.654	0.095	0.430	0.810	25	9
Income	0.739	0.755	0.059	0.630	0.822	20	15
Full sample	0.686	0.687	0.094	0.430	0.822	45	24
Weighted							
	mean	median	sd	min	max	N	Countries
Consumption	0.639	0.649	0.100	0.430	0.810	25	9
Income	0.737	0.750	0.058	0.630	0.822	20	15
Full sample	0.700	0.700	0.090	0.430	0.822	45	24
Authors' calculations. Correlation coefficients have been standardized to cover a two-year interval. Weighted estimates are weighted by the number of observed coefficients by country, such that all countries have equal weight.							

¹⁴ The method used to standardize coefficients to a two year interval is discussed above. The estimated log-trend coefficients are -0.10 for the consumption correlation sample, and -0.08 for the income sample.

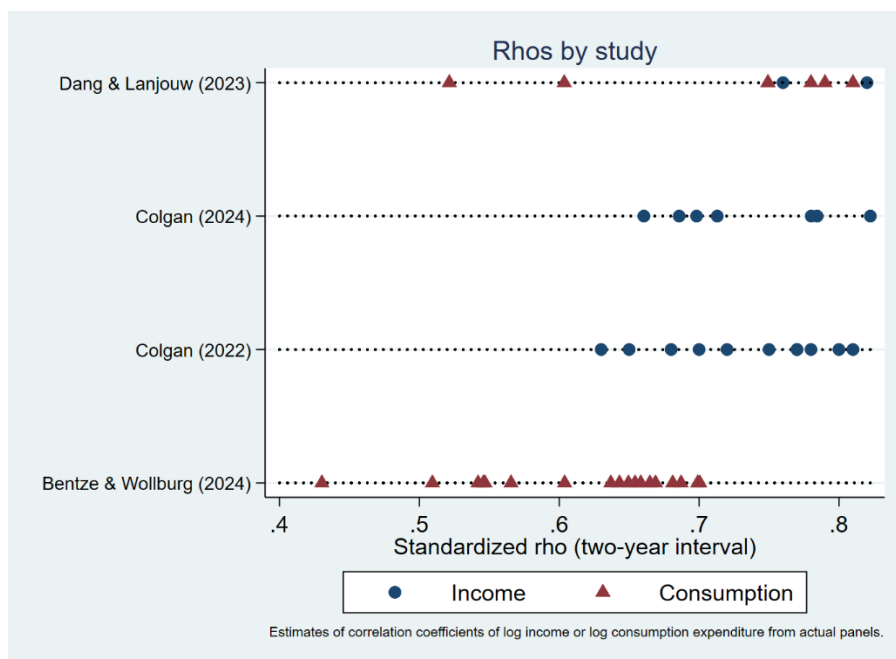


Figure 1 Sample of panel correlation coefficients by study

Importantly, the sample seems to vindicate a crucial assumption for producing synthetic panel bound estimates: that the correlation coefficient is positive, which implies in practice a number between 0 and 1. In fact, the range of standardized coefficients is much narrower than that, minimum and maximum values being 0.43 and 0.82. Weighted sample means are 0.64 for consumption and 0.74 for income, with an overall mean of 0.70. We weight observed coefficients by the number of coefficients by country, such that all countries have equal weight. The aim is to account for the fact that observed correlations for a country are not independent. For example, we have seven observations for Uganda, but only one for Niger; these should not all carry equal weight. Nevertheless, Table 1 shows that weighting does not affect the parameters of the distribution.

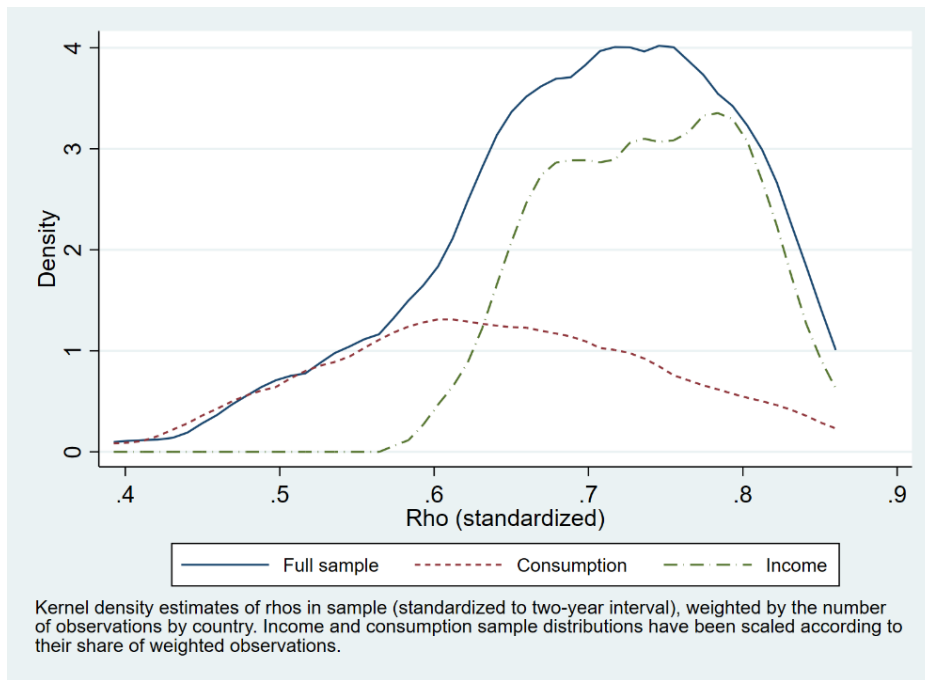


Figure 2 Distribution of correlation coefficients

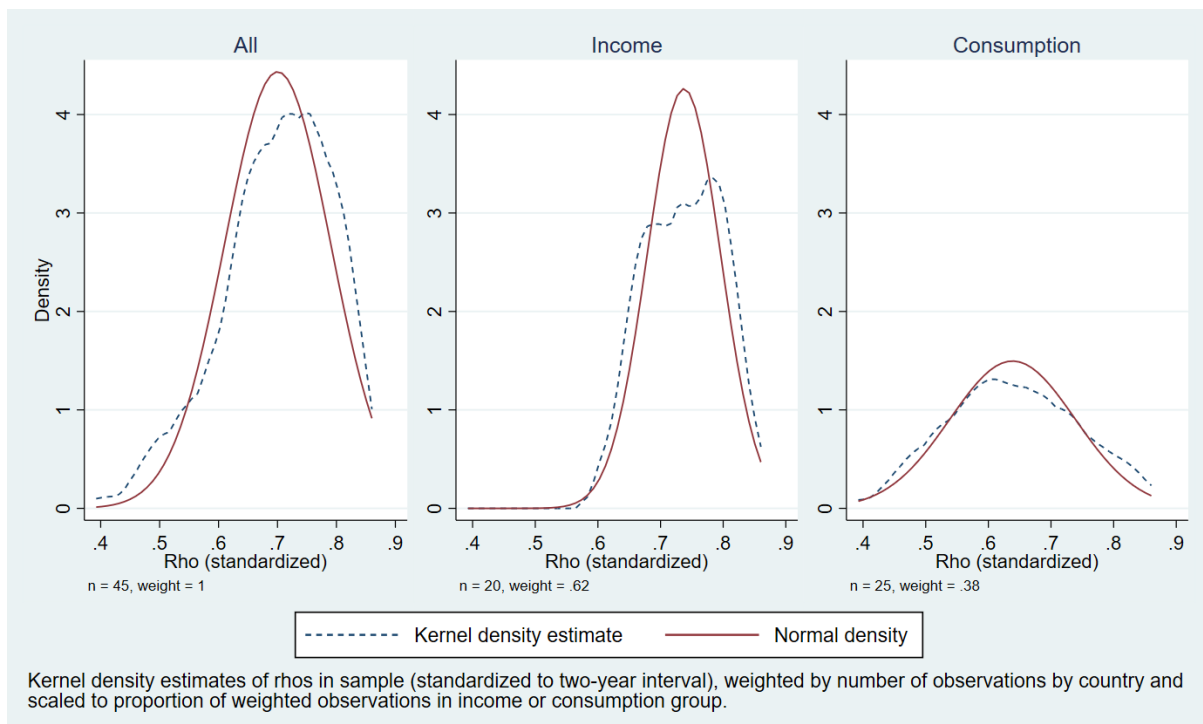


Figure 3 Distributions of correlation coefficients compared against the normal distribution

We investigate the clustering of correlation coefficients by plotting the standardized coefficients against country income levels (Figure 4). It is obvious that countries are sorted into a high income group with income correlations and a low and middle income group with

consumption correlations, reflecting common views about which concept best captures household welfare at different stages of economic development. The income group is more closely clustered, while the consumption one is more spread out, as observed in the kernel density plots too. Taking all coefficients together, a simple OLS regression shows that country log income levels and correlations are positively correlated: the estimated slope parameter is 0.04 with a standard error of 0.01. At a first glance, one could think that the numbers in our sample run counter to the theory of consumption smoothing, since the mean correlation for consumption expenditures is lower than the mean for income. However, we do not have both types of correlations for countries at the same level of income, which is required for a fair frame of comparison. In this sample, countries with a consumption coefficient are also those with lower levels of welfare, i.e., likely to have a larger proportion of informal jobs with variable earnings and to be more vulnerable to variations in agricultural harvests.

Nevertheless, also at similar levels of prosperity, there is considerable variation. This is not surprising given existing differences in economic institutions and social protection systems. At the lower end, correlations vary from 0.43 for Ethiopia (2010-12), to 0.65 for Niger (2011-14), see Table 3 in the Annex. Around an income level of nine log points, we observe a range from 0.52 for Bosnia and Herzegovina (2001-04) to 0.81 for Viet Nam (2004-06). Even within one country, variation can be notable, as in Nigeria, where we have a range from 0.57 (2010-12) to 0.67 (2015-18). Overall, however, we clearly observe a central tendency as shown in Figure 3.

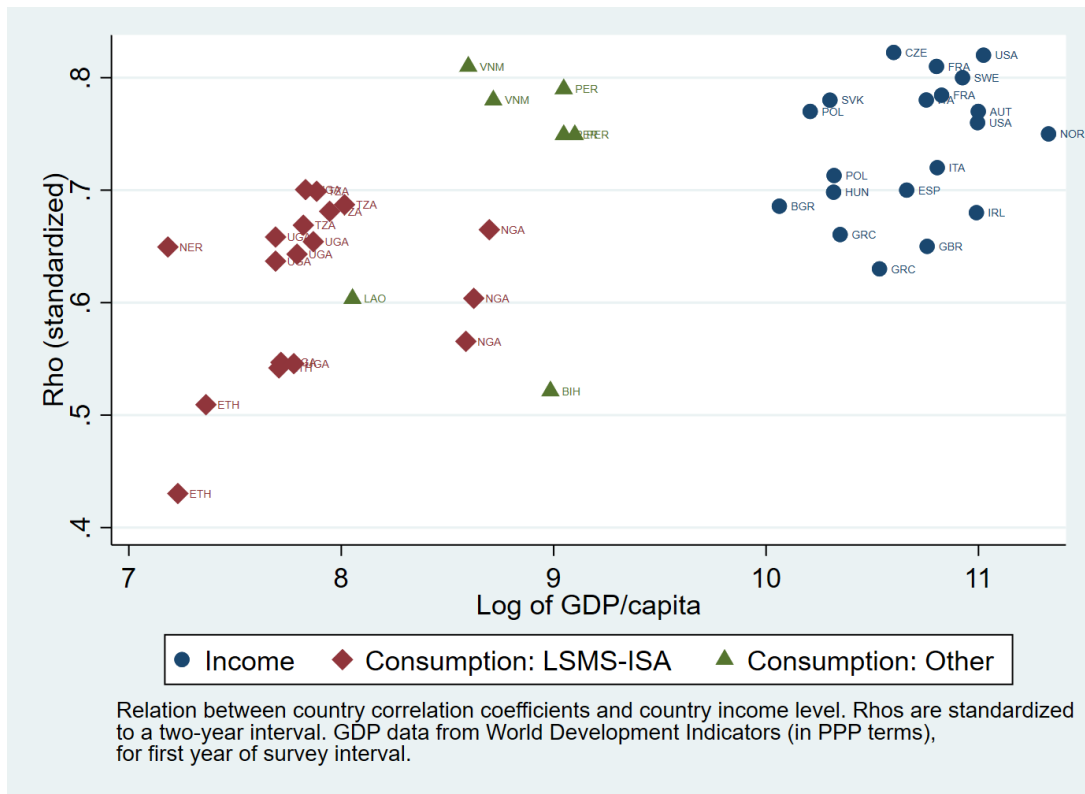


Figure 4 Correlation coefficients by income level

3. Validation

As a first step in validating the method, we compare the correlations from Colgan (2023, 2026) and Dang and Lanjouw (2023) that we had collected for a preliminary application (Rongen & Lanjouw, 2024) to additional correlations from the LSMS-ISA surveys (Bentze & Wollburg, 2024). Figure 1 and Figure 4 show that these two subsamples are generally in line with each other, despite being drawn from countries with differing income levels. This strengthens the evidence for the assumption that there is indeed an underlying distribution of correlation coefficients that is narrower than the 0/1 interval.

A second validation exercise is that we compare the results from our bootstrap approach with estimated true transition outcomes based on an actual panel of Tanzanian data for the interval 2014-2020.¹⁵ In this analysis, we focus on poverty transition outcomes and examine to

¹⁵ This validation analysis is based on a research note for the program Data and Evidence to End Extreme Poverty (DEEP) (Rongen, 2026). That publication is accompanied by a template Stata do-file and our dataset of correlations, and contains usage instructions for the template do-file.

what extent our estimates align with the true estimates, and how the bootstrap estimates improve precision.

3.1. Data and implementation

We use two waves of the Tanzania National Panel Survey, implemented by the Tanzania National Bureau of Statistics: Wave 4 (2014-2015) and Wave 5 (2020-2021).¹⁶ The balanced panel consists of 2,968 households for which total consumption expenditure data is available. Applying the age range restrictions on household heads (between ages 25 and 60 in Wave 4), between 2,250 to 2,400 households remain, depending on the model specification. The poverty line used is the World Bank's extreme poverty line of USD 2.15 per person per day in 2017 PPP terms, equivalent to TZS 1672 per person per day. At this line, static poverty estimates in the balanced panel sample are as documented in Table 2 below.

Table 2 Estimates of static headcount poverty rates in Tanzania (balanced panel sample)

Wave	Urban population %	Headcount rate %		
		Rural	Urban	Total
4	32.1	47.2	11.9	35.9
	(2.6)	(1.9)	(1.7)	(1.6)
5	30.6	46.8	11.7	36.1
	(2.4)	(1.7)	(1.7)	(1.5)
Authors' calculations based on Tanzania National Panel Survey, Refresh Panel balanced sample. Note: standard errors in parentheses. The decrease in urban share of the population could be due to differential attrition rates.				

For the purpose of this validation analysis, the balanced panel was split into two random parts, stratified by rural/urban status. Each part has an equal number of observations, as if they were separate cross-sectional samples. Subsequently, one part was used to estimate the Wave 4 consumption model, the other to fit the Wave 5 one. This halved the sample size compared to the panel. However, in practice, cross sectional surveys often have larger sample sizes than panel surveys, making them more suited to establish profiles of the population that is living in, for example, chronic poverty.

To obtain bootstrap point estimates, we draw a bootstrap sample from our set of empirical correlation coefficients. This set contains a coefficient based on the Tanzanian data for

¹⁶ These surveys are part of the Harmonized LSMS-ISA Panel dataset (Bentze & Wollburg, 2024), which formed the basis for our analysis. Additional variables were merged in from the separate original data files (NBS (2015) and (2021), see <https://microdata.worldbank.org/index.php/catalog/2862> and <https://microdata.worldbank.org/index.php/catalog/5639> for further information). Note that we only use the households from the so-called 'Refresh Panel', not the additional ones from the NPS-SSD extended panel.

the 2014-2020 interval. This particular coefficient was removed prior to drawing the bootstrap sample.

We employ two different specifications of a consumption model for the synthetic panel estimation, in order to assess its performance under differing model quality. Specification 1 is parsimonious, leading to rather low R^2 statistics (0.29 and 0.20 in Waves 4 and 5 respectively). The explanatory regression variables are characteristics of the household head (age, the square of age, a gender dummy and a formal education dummy) and a dummy for the household's location in either an urban or rural location.¹⁷ Specification 2 replaces the formal education dummy with five dummies for education level. Moreover, dummies are added for the region in which the household is located.¹⁸ This last specification emulates the synthetic panel model used by Aikaeli et al. (2021) and increases the R^2 considerably, to 0.46 and 0.37 respectively. In both models, the dependent variable is the logarithm of real annual household per capita consumption expenditures in Tanzanian shilling (TSH). Regression results are given in Table 2 below. The low R^2 of Specification 1 implies that the regression residuals are relatively large, which affects the precision of the synthetic panel estimates, causing, for example, the 0/1 bounds to be wider.

3.2. Results

Poverty transition outcomes are assessed by comparing the synthetic panel estimates (both the 0/1 bounds and the 4SD band around the point estimates) to the true panel estimate and its confidence interval.¹⁹ We do so at the population level and separately for the urban and rural subgroups, for the two specifications of our model. To start, Figure 6 considers joint poverty transition probabilities, such as the probability that someone is chronically poor (i.e. poor in both survey waves). These probabilities reflect the share of the population or subgroup that experienced this transition. In addition, Figure 7 shows conditional transition probabilities.

It is clear that the 0/1 bounds perform well in terms of covering the true estimates, as was also demonstrated for other countries (see Garcés-Urzainqui et al. (2021) for an overview). In fact, all true transition estimates fall within the 0/1 bounds. Our point estimate 4SD bands perform rather well too. A first observation is that there is no major difference in performance between,

¹⁷ The LSMS-ISA Harmonized Panel has only two education variables; of those, inclusion of the primary education indicator is not feasible because its values are not consistent over the two waves.

¹⁸ One drawback of including such regional dummies is that this adds the assumption that households have not moved between regions.

¹⁹ In order to make a more realistic comparison, the true panel estimates are based on the full balanced panel, not on the random subsamples used in the synthetic panel analysis.

on the one hand, estimates based on our full sample of correlation coefficients, which includes both income and consumption correlations, (diamond markers in the figure) and, on the other hand, estimates based on consumption correlations only (triangle markers). They cover the true estimates equally often, so in the following we focus on the full sample estimates.

A second observation, however, is that there is a clear difference in how the two specifications of the income model perform. Rounding to the nearest percentage point, the Specification 2 4SD bands cover true joint estimates 9 out of 12 times, compared to 7/12 times for Specification 1.²⁰ Notably, all full population and rural subgroup joint estimates are within Specification 2's 4SD bands. Only for the urban subgroup does Specification 1 perform marginally better in this respect, but the differences are small. Consider, for example, the share of the urban population that is non-poor in Wave 4, but poor in Wave 5 (middle row, Figure 6), which in the true panel is estimated at 7 percent when rounded (95% CI: 5 – 9 percent). The 4SD band in Specification 1 is 7 – 10 percent, just covering the true estimate, while it is 8 – 11 percent for Specification 2, just missing the mark. Something similar holds for the conditional probabilities, where Specification 2 misses all true estimates for urban areas, while Specification 1 misses only two. Where both miss, however, Specification 2 is much less off than Specification 1. Inspecting the conditional probability of escaping poverty in urban areas (middle row, right side in Figure 7), we have a true estimate of 49 percent (95% CI: 35 – 63 percent).²¹ The Specification 2 band is 51 – 69 percent, while it is 54 – 79 percent for Specification 1, meaning that Specification 2 is closer and narrower. Much the same holds for the conditional probability of remaining in poverty that people in urban areas face.

From the examples above, we can conclude that in those cases where the true estimate is not within the 4SD band, the bands do not lead us to infer conclusions that are dramatically different from ones based on the true estimates. Either there is considerable overlap between the true estimate's confidence interval and the 4SD band, or estimates are so condensed that the size of differences is not meaningful in a practical sense. This conclusion is clear when considering the four Specification 2 conditional probabilities for the urban subgroup, on the right side, middle row, of Figure 7. None of these 4SD bands cover the true estimate, but all come sufficiently close. For example, the conditional probability of remaining non-poor is estimated in the panel at 92 percent (95% CI: 90 - 95), while the 4SD band for Specification 2 is 88 – 90 percent.

²⁰ These twelve cases consist of the four joint transition probabilities for each of the three groups.

²¹²¹ Note that these true estimates are for Specification 2; they are 48 percent (95% CI: 35 – 61) for Specification 1. The difference is due to a slightly smaller sample for Specification 2 because of missing values (see Table 4 in the Appendix).

All these estimates would lead us to conclude that, by and large, urban residents have a 9 out of 10 chance of staying out of poverty.

A third observation that we can garner from these validation results is that the synthetic panel estimates underperform in urban areas in Tanzania. There are a couple of potential explanations for this.²² One factor is that poverty rates are much lower here, at about 12 percent, compared to the full population average of around 36 percent. Moreover, since just under a third of the population lives in urban areas, transition estimates are based on a smaller sample of about 500 households in each round. It is surprising that this proportion decreases from Wave 4 to Wave 5 (see Table 2), whereas in the context of urbanization an increase would have been expected. The decrease could be explained by higher attrition rates in urban areas. Another reason for underperformance in urban areas may be that the value of the urban dummy variable varies between survey waves, which are six years apart. In the balanced panel, 17 percent of survey households change classification between rounds; this could be due to official area reclassification, migration, or imputation errors. True estimates for the subgroups are based on household urban/rural classification in the Wave 4 survey. Note that the synthetic panel results assumes that households do not migrate, because household location is included as a time-invariant explanatory variable.

Another observation is that the 4SD bands, especially under Specification 2, have a tendency to somewhat underestimate the static categories (poor - poor and non-poor - non-poor), while overestimating the mobile ones; the estimates based on the consumption sample do so most. This follows from the fact that the actual panel correlation coefficient over the six year interval between survey rounds is estimated at 0.61, while the means of both our bootstrap samples of correlations are smaller, at 0.54 for the consumption sample and at 0.59 for the full sample mean.²³ The closer the correlation coefficient is to zero, the more mobility is estimated. Thus, the estimates based on the consumption correlations, which are closer to zero, overestimate mobility most. Correspondingly, since the mean of the full sample is closer to the

²² An explanation that does not seem to be playing a role is the fact that correlations differ between groups. The actual consumption correlation coefficient is higher in urban areas (at 0.58) than in rural areas (at 0.49), but the urban coefficient is closer to the mean of our full sample of empirical correlations (0.59 when adjusted to the six-year interval). Thus, this cannot explain underperformance in urban areas. Moreover, in the Specification 2 residuals, the difference in correlations even disappears: the residual coefficient is about 0.38 for both groups.

²³ The actual coefficient equals 0.61 in both model specifications (when rounded to two decimals). The coefficient is estimated across the household head age range of 25–60 in Wave 4, and robust to the age range selected. The means of the samples of correlations are derived by adjusting the means of the bootstrap sample of standardized correlations (0.655 and 0.699 respectively) to a six-year interval based on a log-trend relation.

actual correlation, the estimates derived from this sample perform marginally better: they are somewhat less off when the 4SD band fails to cover the true estimate.

A final observation is that the 4SD bands for Specification 2 are somewhat narrower. This relates to the higher R^2 : compared to Specification 1, the variation in the residuals is smaller.²⁴ The spread in our bootstrap estimates derives from the distribution of these residuals, which are assumed to follow a bivariate normal distribution. Smaller standard deviations mean a more concentrated distribution. This implies that there is less variation in poverty transition outcomes across the draws of our bootstrap sample (the predicted transition probabilities for a given household will be less variable), and thus that the 4SD band will be narrower. The average width over the four population-level joint transitions is 4.8 percentage points for Specification 2, against 6.3 percentage points for Specification 1.

As noted, the 0/1 bounds have a perfect score in terms of covering the true estimate in this validation analysis. However, the point estimate 4SD bands perform well in the sense that they cover the true point estimate most of the time and that they are much narrower than the 0/1 bounds. This holds in particular for the more elaborate Specification 2. Moreover, both sets of samples of correlation coefficients yield good results. In those cases where the true estimate does not fall within the 4SD band, this does not lead to conclusions straying off the mark: there is significant overlap between the 4SD band and the true estimate's confidence interval, or differences are small in a practical sense.

The improved precision of these 4SD band estimates carries practical significance. Imagine that the panel data are not available. Then, looking at just the static poverty numbers, one might think that between 2014 and 2020, poverty did not change much. The synthetic panel provides evidence that much more is going on. Based on Specification 2, the 0/1 bounds indicate that some 4 to 18 percent of the population fell into poverty. The 4SD band based on the full sample of correlations narrows this down to 13 – 18 percent. The true estimate is 13.8 percent (95% CI: 12 – 16 percent). Precision improvements are even larger for certain conditional probabilities, for example when it comes to the chances of escaping poverty that poor rural Tanzanians face. The 0/1 bounds are rather wide, estimating a conditional probability of 9 – 45 percent, making them of little practical use. The full sample 4SD band of 33 – 46 percent chance represents a considerable improvement. The true estimate of 33 percent is just inside this band (95% CI: 29 – 37 percent).

²⁴ The regression residuals in Specification 2 have a smaller standard deviation than the Specification 1 residuals (0.54 log points against 0.62 in Wave 4, 0.66 points against 0.77 in Wave 5).

Another feature of the panel data is that it enables us to test an important assumption of the synthetic panel, i.e., that the regression residuals follow a bivariate normal distribution. Figure 5 provides a visual analysis of the marginal and joint distributions for both model specifications, based on the full balanced sample. The marginal distributions of the residuals are approximately normal – deviations seem minimal. The standard deviation of the residuals is smaller under Specification 2. The panels on the right side plot the residuals in each wave against each other. We can observe that both panels have more observations in the lower-left and upper-right quadrants, consistent with a correlation coefficient that is positive. Indeed, the coefficients are estimated at 0.475 for Specification 1 and at 0.379 for Specification 2. The lower residual correlation in the latter specification can also be observed directly from the lower right panel, as the points are scattered more randomly than in the upper right panel. The two scatter plots do not indicate that the assumption of bivariate normality is untenable, although formal tests reject univariate and bivariate normality.

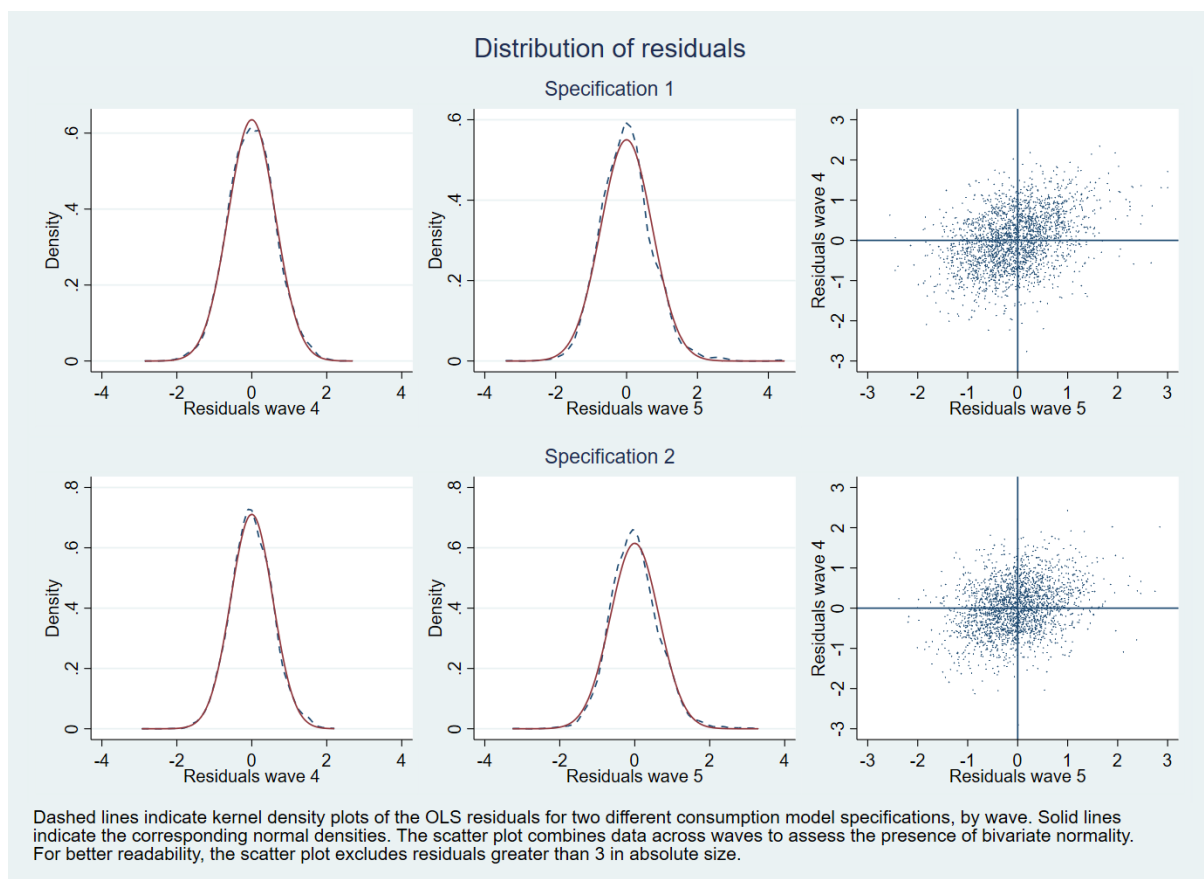


Figure 5 Distribution of regression residuals Tanzania

In sum, the method seems to perform well. In the trade-off between covering the actual estimate and delivering a relevant narrow band of estimates, the bootstrap point estimates strike a good balance. They thus serve as a relevant tool to inform policy makers about poverty dynamics.

Based on this validation analysis, we cannot make a recommendation as to whether it is preferable to use the full sample of correlation coefficients, or the sample of only those coefficients that align with the welfare measure surveyed. Our conclusion is that, in this case, estimates based on the full sample perform somewhat better because the mean correlation coefficient in that sample is closer to the actual consumption serial correlation mean for this survey. However, the researcher cannot know this a priori absent a true panel, as in the application presented in the next section.

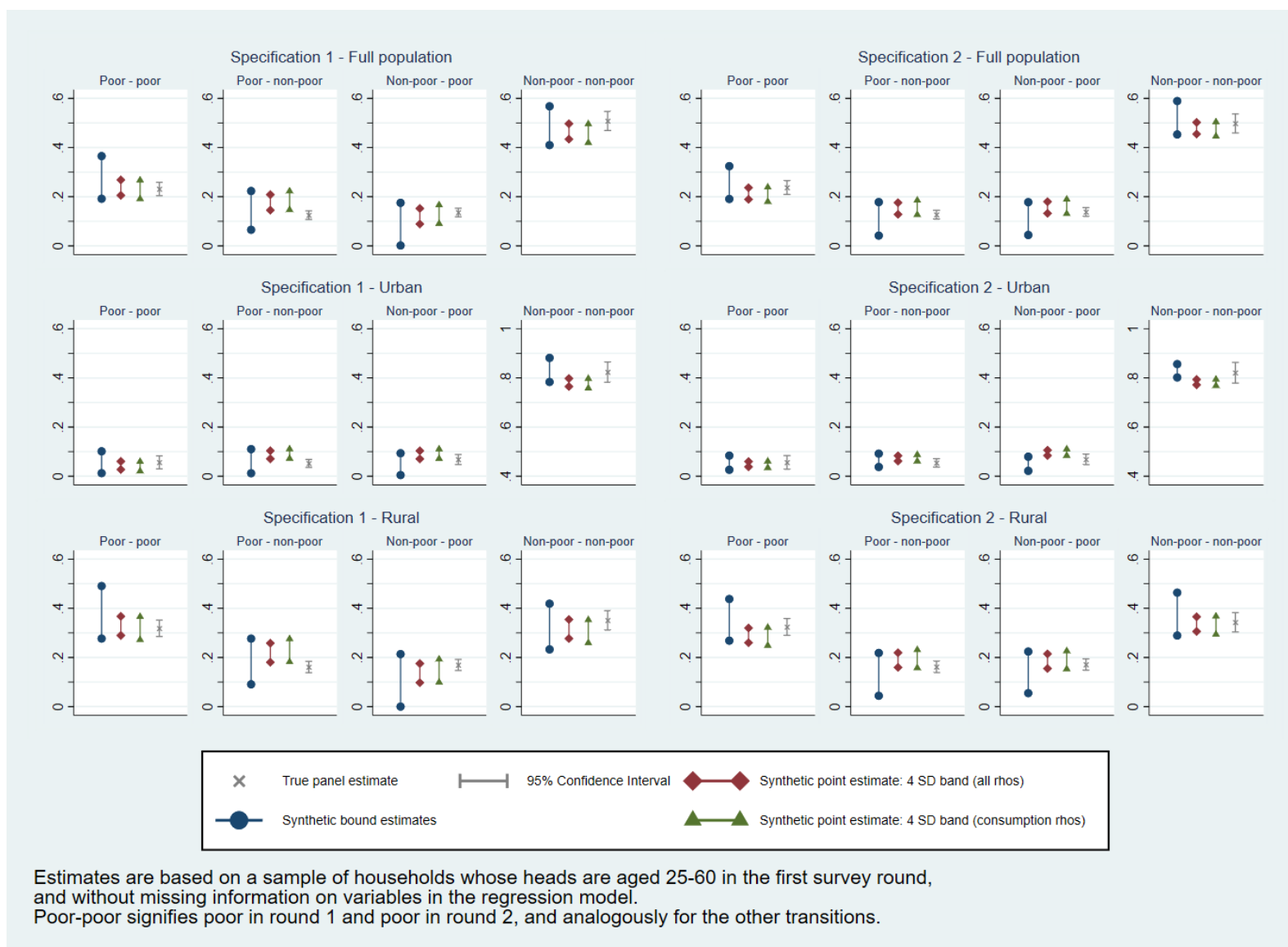


Figure 6 Population shares poverty transitions

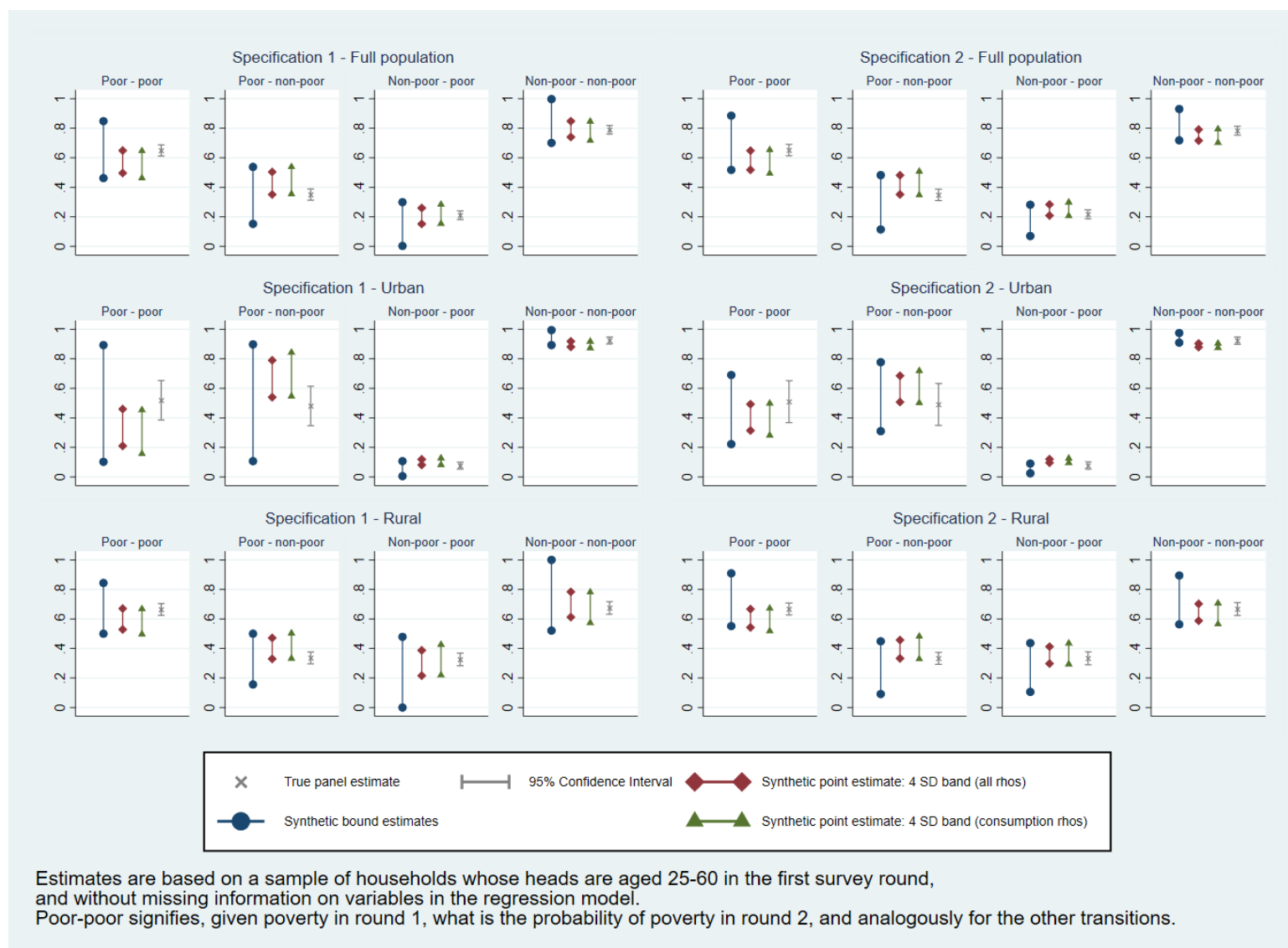


Figure 7 Conditional poverty transition probabilities

4. Application: Malaysian poverty dynamics

This section discusses how the method performs in a context where we have no choice but to rely on repeated cross-sectional data to obtain information about poverty dynamics. Our aim is to show how the precision of results improves compared to the bounds method and, by extension, how that increases their relevance to policy makers. The analysis in this section builds on Rongen & Lanjouw (2024).²⁵ Below we summarize our implementation of the synthetic panel based on Malaysian data; full details can be found in the publication cited.

4.1. Data and implementation

In this application, we use cross-sectional data from the Household Income Survey (HIS), collected by the Department of Statistics Malaysia (DOSM). The data were collected in eight rounds over the years 2004 to 2022; we estimate transitions for all seven pairs of consecutive survey rounds. Sample sizes of the surveys are large, ranging from 36,000 to over 80,000 household-level observations per survey round. The dependent variable in the income model is the natural logarithm of monthly household pre-tax income per capita. Income data have been adjusted for inflation and spatial price differences. Explanatory variables are age, ethnicity and gender of the household head, in addition to regional and urban/rural dummies. The synthetic panel is estimated for a subsample of households with heads in the age range of 25–60. We employ an absolute poverty line of RM 527 per capita per month to determine the poverty status of a household. Where we analyze vulnerability, this is defined as facing a conditional probability of 40 percent or higher of falling into poverty the next period.

One complication in applying the proposed formula by Dang & Lanjouw (2023) to transform an income correlation to a correlation of residuals is that the proposition may yield a negative residual correlation if the variation in the cross sectional data at hand does not match well with the sampled income correlation from a different source. As the residual correlation is the only free parameter in the proposition, it is then forced to become negative, which violates our assumption about the correlation range. Hence, these draws were excluded for those intervals where they yielded a negative residual correlation. In this application to Malaysia, this was the case for one specific draw from the full sample of correlations, which led to a negative

²⁵ Specifically, we present new versions of the Figures 1, 2, 7, 8 and 9 in Rongen and Lanjouw (2024).

residual in five intervals (but not in two others).²⁶ It did not occur in the sample drawn from income correlations only.

4.2. Results

First, we provide estimates at the aggregate level: joint poverty transition probabilities, conditional poverty transition probabilities, and broader mobility patterns. In order to assess robustness, we show estimation results for two sets of bootstrap samples: a sample drawn from income correlations only, and a more conservative sample drawn from both consumption and income correlations. It will become clear that the differences between the results from these two sets of samples are small. Second, we zoom in on selected probabilities for specific subgroups, in order to assess how the bootstrap estimates perform in smaller groups, and to see whether the difference between the two samples of correlations is small for subgroups as well. We discuss differences in the poverty entry shares of geographically-defined subgroups, and chronic poverty and poverty exit shares for a wider selection of groups over a specific interval. All figures provide bootstrap point estimates wrapped in a band of two standard deviations, both up and down. This 4SD band is the focus of our analysis.

Figure 8 and Figure 9 below show, first and foremost, that our 4SD band yields poverty mobility estimates that are much narrower than the 0/1 bound estimates. This is especially notable for the conditional probabilities displayed in Figure 9. For example, our new estimate for the probability of remaining poor over the 2019-2022 interval is between 25-40 percent, while the 0/1 bounds were 18-75 percent. Second, estimation results are not greatly affected by which set of correlations is used for the bootstrap sample. At the aggregate level, sensitivity of the point estimates to this choice of correlations is negligible. The 4SD band for the full sample estimates is slightly wider, as expected: this sample spans a wider range of correlations, leading to more variation in transition estimates. Nevertheless, the bands largely overlap and both bands perform much better than the 0/1 bound estimates.

²⁶ For example, using the estimated coefficient vector and variances for the Malaysian data in the interval 2009-2012, the proposition turns approximately into $\rho_\varepsilon = -0.8 + 1.7 \rho_y$, implying that any sampled income correlation smaller than 0.48 leads to a negative residual correlation.

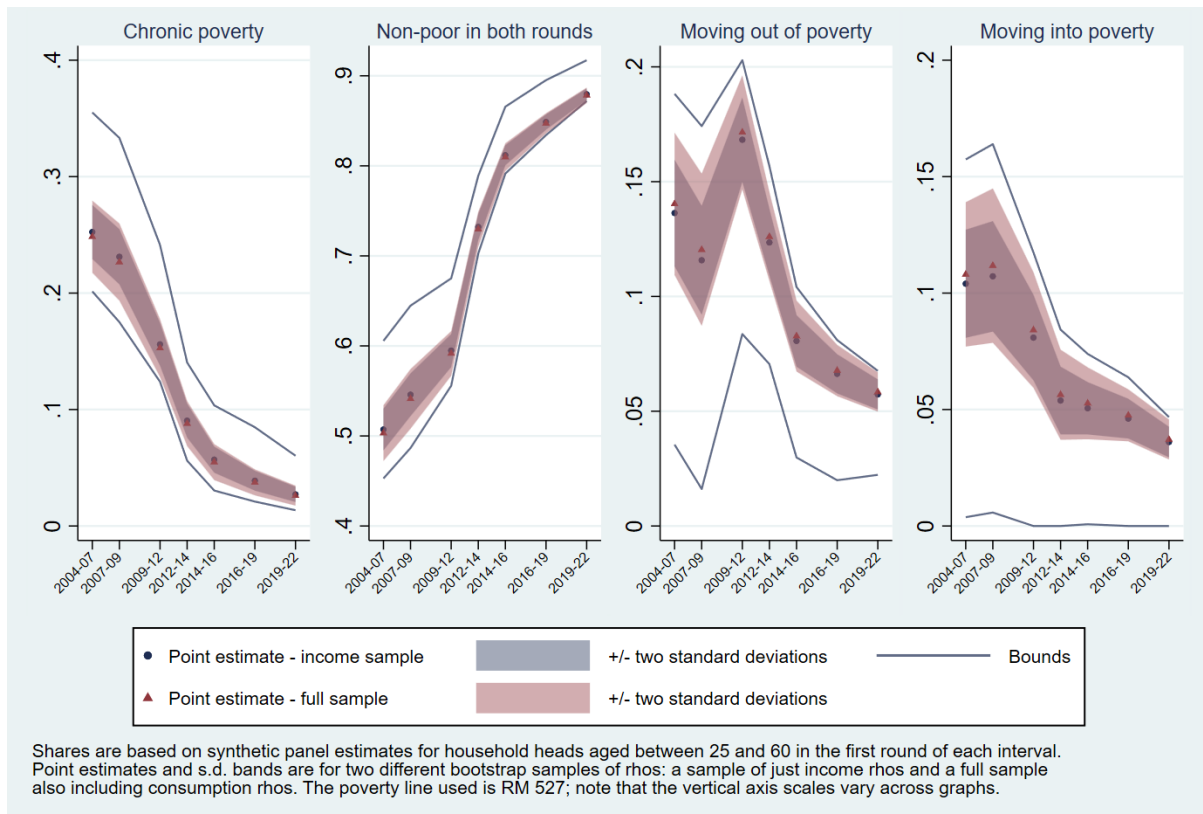


Figure 8 Poverty transition population shares

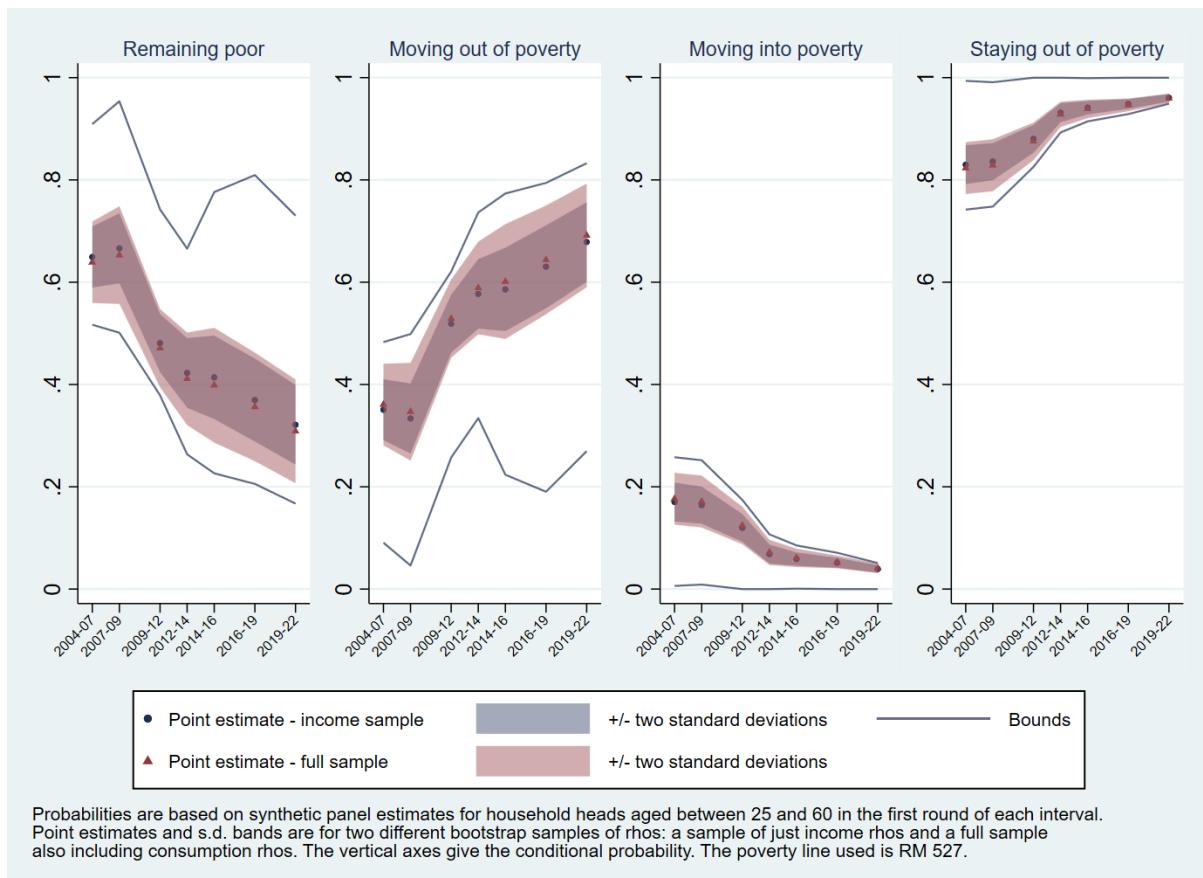


Figure 9 Conditional poverty transition probabilities at population level

Furthermore, using bootstrap estimates improves the precision of our broader set of mobility measures, which focus on upward and downward mobility, vulnerability, and economic security – displayed in Figure 10. Estimates improve across the board.

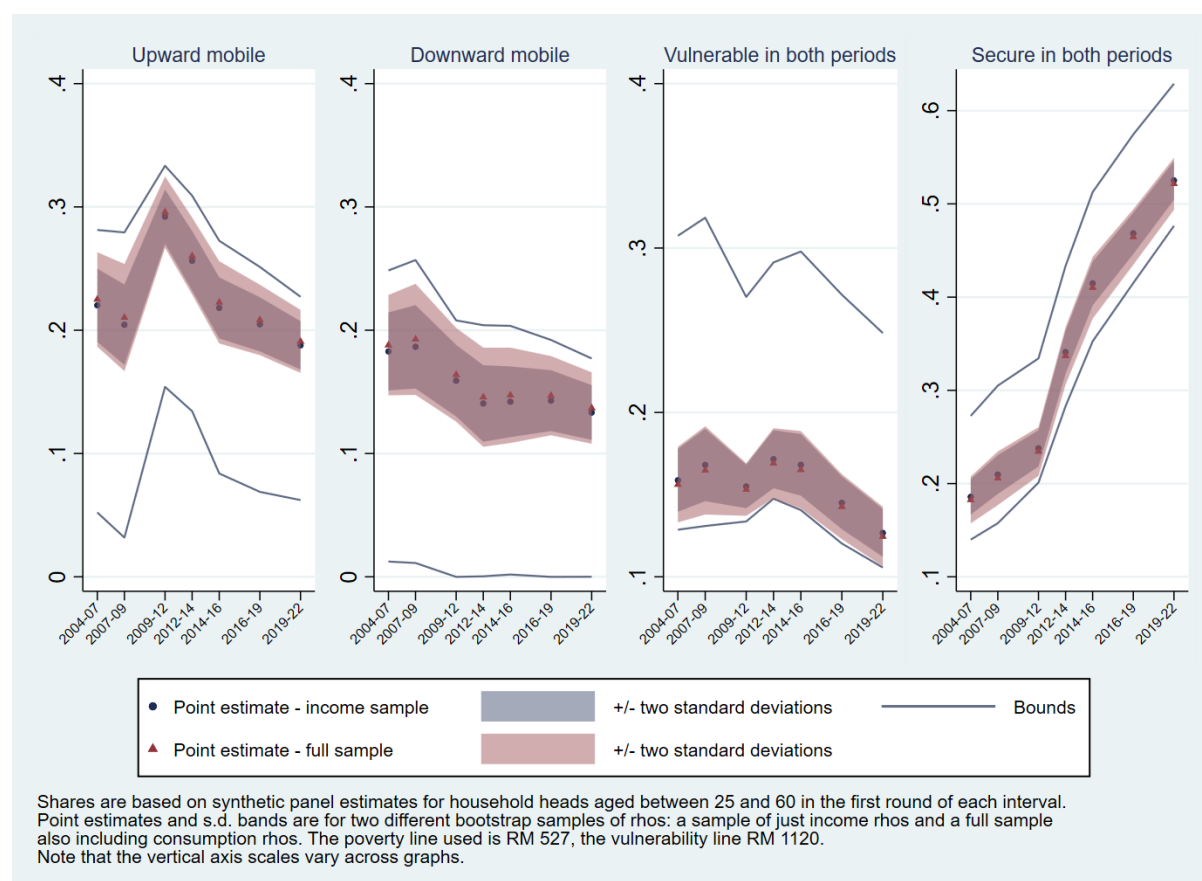


Figure 10 Broader economic mobility measures at population level

Zooming in on the share of selected subgroup populations who move into poverty in Figure 11, we observe improvements in precision similar to those at the aggregate level. These particular subgroups were chosen because they represent some of the starkest subgroup differences in poverty mobility. Our approach clearly narrows down estimates and improves, therefore, their suitability as a basis for policy action. For example, the estimate for the share of rural East Malaysians who fell into poverty over the 2019-2022 interval narrows from 0 – 18 percent to 8 – 14 percent. In addition, Figure 12 displays the two conditional probabilities of remaining poor and falling into poverty over the 2019 – 2022 interval for a range of subgroups.²⁷ It

²⁷ All of these subgroups are included as dummy variables in the income model, but that is not a necessity. Probabilities can also be aggregated on the basis of characteristics not included in the regression.

contrasts estimates based on the income sample with those derived from the full sample of correlations. As at the aggregate level, the two sets of estimates do not differ meaningfully from each other.

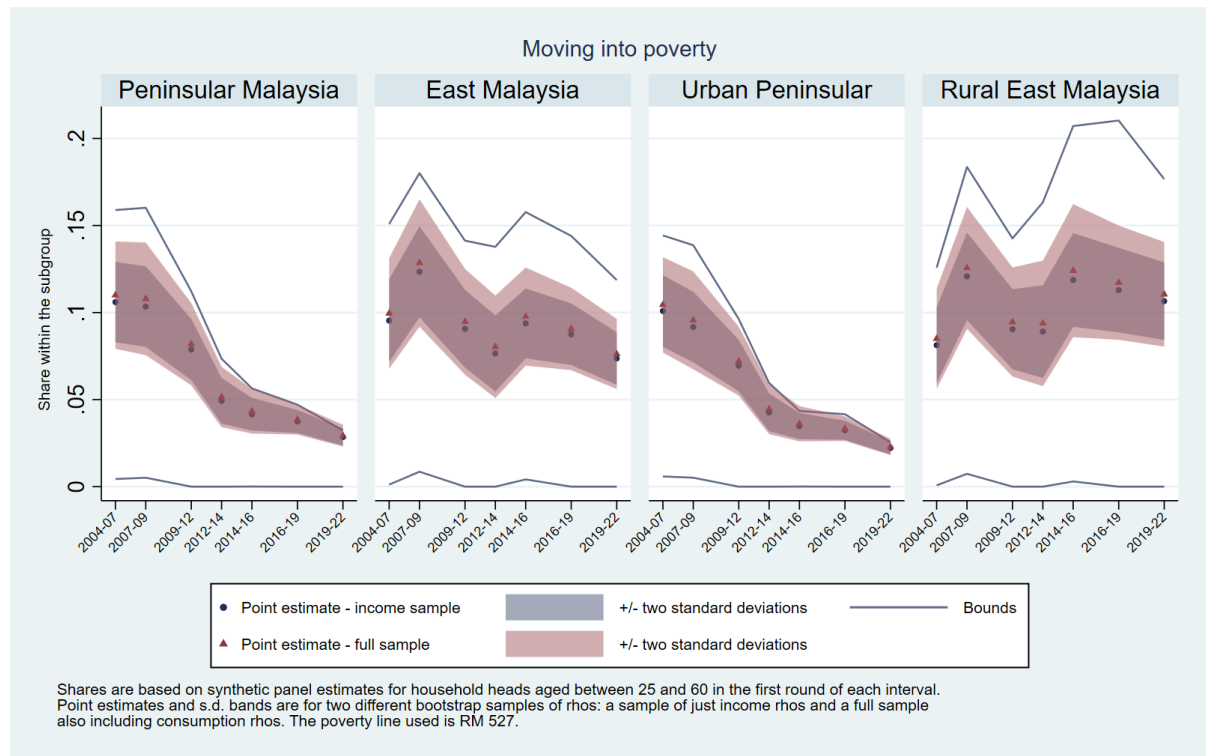


Figure 11 Share of selected subgroups moving into poverty

A different source of uncertainty is to what extent poverty transition estimates are affected by estimation error of the income model parameters. Given large sample sizes, such errors are usually assumed to be small in comparison to the width of the synthetic panel bounds. The question becomes more relevant now that we are estimating narrower bands. Hence, we assessed the variability of point estimates by fixing the correlation coefficient (at the mean correlation in our sample, see Table 1), and estimating transition probabilities across 100 bootstrap samples of the Malaysian data, which enables the computation of bootstrap standard errors. At the aggregate level, standard errors of the four joint probabilities range between 0.2 - 0.4 percentage points for the 2004-2007 interval, and between 0.1 – 0.2 percentage points for the 2019-2022 interval, which has a larger sample. The respective ranges for the conditional probabilities are 0.3 – 0.4 and 0.1 – 0.5 percentage points. Across the urban and rural subgroups, standard errors are slightly larger, as expected. These findings imply that model error is still comparatively small, such that the 4SD band can arguably be assumed to reflect the resulting uncertainty.

Taken together, the results on Malaysian income mobility in this section demonstrate that the 4SD band estimates considerably improve the precision and relevance of the synthetic panel method. This is most evident for the conditional poverty transition probabilities and the indicators of broader economic mobility.

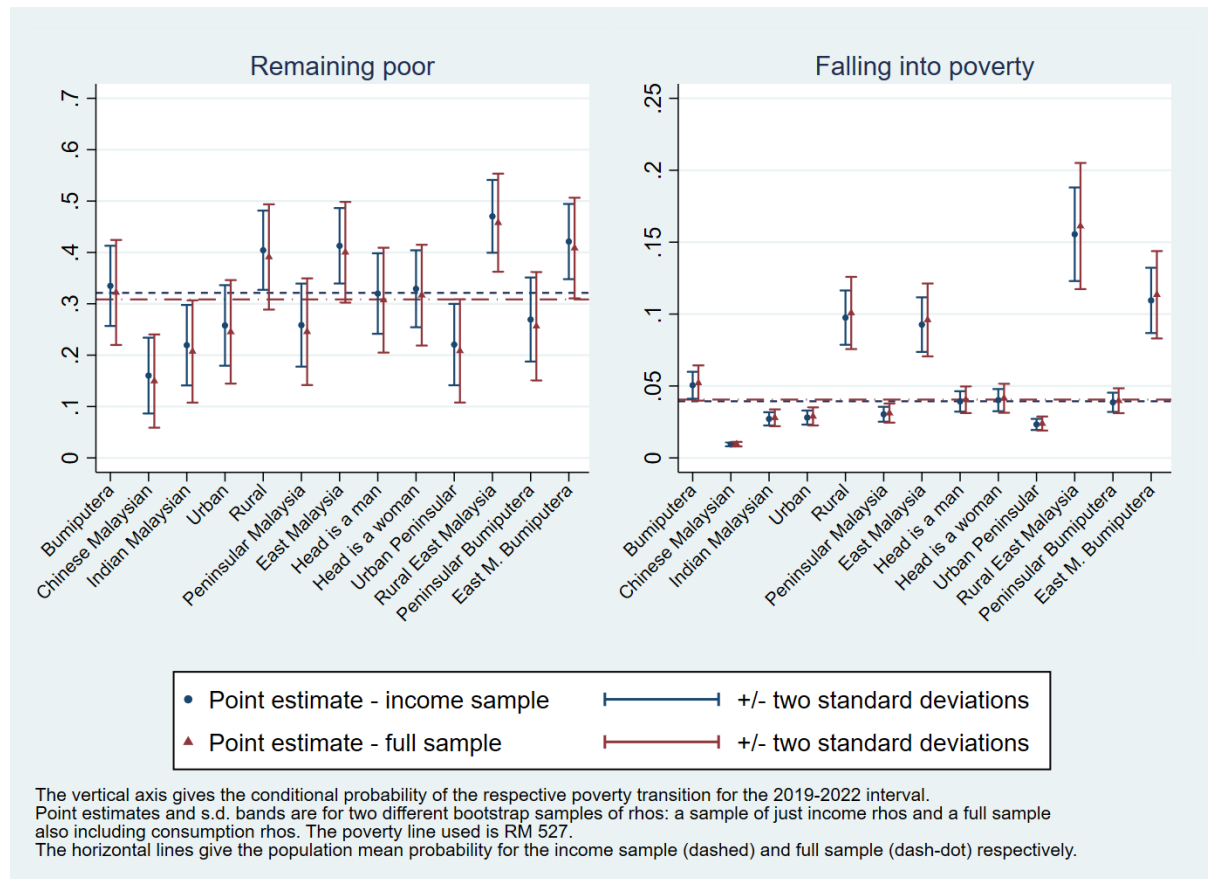


Figure 12 Subgroup conditional probabilities 2019-2022

5. Conclusion

We have presented a novel approach for obtaining point estimates of household poverty transitions and income mobility in a context where panel data are unavailable and we need to rely on cross-sectional data. These bootstrap point estimates for synthetic panels, in combination with a four standard deviation band around the point estimates, narrow down the estimation bandwidth considerably compared to existing synthetic panel bound estimates.

We do not know how household incomes are correlated over time if we have only cross-sectional data. Our approach starts from the idea that such correlations spring from a certain distribution, which spans a range that is narrower than the 0 – 1 interval, and that we may approximate this distribution by means of correlations observed in true panel data. Hence, we

have collected a sample of empirical correlation coefficients that describe the correlation of household income or consumption expenditure over time. Importantly, these correlations from various sources are consistent with each other and are clearly limited in range: 0.43 – 0.82 in our sample.

A validation using Tanzanian data and an application to Malaysia have demonstrated how this approach can be put to use to improve precision of poverty transition estimates, such as the share of the population and subgroups that is chronically poor, or the conditional probability of falling into poverty. These analyses show how bootstrap synthetic panel estimates can provide a stronger evidence base for policy.

Building a larger sample of correlation coefficients by including correlations from additional actual panels will further strengthen the method. In addition, such a larger sample may make it feasible to estimate a model that explains the variation in correlation in terms of country characteristics. That would then enable prediction of the correlation coefficient for a target country without panel data, as input for synthetic panel estimation. A further study of interest would be to simulate a range of panel datasets, each reflecting a different average serial correlation. Based on these simulated data, one could analyze at which extremes in the distribution of actual correlations the performance of our proposed method starts to worsen.

References

- Aikaeli, J., Garcés-Urzainqui, D. and Mdadila, K. (2021) “Understanding Poverty Dynamics and Vulnerability in Tanzania: 2012–2018”, *DEEP Working Paper 09*, Data and Evidence to End Extreme Poverty Research Programme, Oxford. <https://doi.org/10.1111/rode.12829>.
- Ball, Christopher. (2016). Estimating income dynamics from cross-sectional data using matching techniques (tech. rep.). Working Papers in Public Finance 06.
- Bentze Thomas and Philip Wollburg. (2024). “A Longitudinal Cross-Country Dataset on Agricultural Productivity and Welfare in Sub-Saharan Africa”. *Policy Research Working Paper* 10976. World Bank: Washington D.C.
- Colgan, Brian. (2023). “EU-SILC and the potential for synthetic panel estimates”. *Empirical Economics* 64. pp. 1247–1280. <https://doi.org/10.1007/s00181-022-02277-7>.
- Colgan, Brian. (2026). “Long run poverty dynamics in new EU member states”. *Empirical Economics*. (forthcoming)
- Dang, Hai-Anh, Peter Lanjouw. (2017). “Welfare Dynamics Measurement: Two Definitions of a Vulnerability Line”. *Review of Income and Wealth*, 63(4): 633-660.
- Dang, Hai-Anh, Peter Lanjouw. (2023). “Measuring Poverty Dynamics with Synthetic Panels Based on Repeated Cross Sections”. *Oxford Bulletin Of Economics And Statistics*, 85(3): 599-622.
- Dang, Hai-Anh, Peter Lanjouw, Jill Luoto, and David McKenzie. (2014). “Using Repeated Cross-Sections to Explore Movements in and out of Poverty”. *Journal of Development Economics*, 107: 112-128.
- Garcés Urzainqui, David (2023). *The Distribution of Development: Essays on mobility, inequality and social change*. PhD Thesis. Rozenberg Publishers and Tinbergen Institute. <https://doi.org/10.5463/thesis.101>
- Garcés-Urzainqui, D., Lanjouw, P., and Rongen, G. (2021). “Constructing synthetic panels for the purpose of studying poverty dynamics: A primer”. *Review of Development Economics*, 25, 1803– 1815. <https://doi.org/10.1111/rode.12832>
- Hérault, Nicolas and Stephen Jenkins. (2019). “How valid are synthetic panel estimates of poverty dynamics?”. *Journal of Economic Inequality* 17: 51-76.
- National Bureau of Statistics Tanzania (NBS) (2015). National Panel Survey - Wave 4, 2014 - 2015. Ref.TZA_2014_NPS-R4_v03_M. Dar es Salaam, Tanzania: NBS. Downloaded from <https://microdata.worldbank.org> on 20 January 2026.

National Bureau of Statistics Tanzania (NBS) (2021). National Panel Survey 2020-21, Wave 5.

Ref: TZA_2020_NPS-R5_v02_M. Dar es Salaam, Tanzania: NBS. Downloaded from

<https://microdata.worldbank.org> on 20 January 2026.

Rongen, Gerton. (2026). “Poverty dynamics in Tanzania, 2014–2020: A validation of synthetic panel bootstrap point estimates”. *DEEP Working Paper 41*. Data & Evidence to End Extreme Poverty (DEEP) programme, Oxford Policy Management.

Rongen, Gerton and Peter Lanjouw. (2024). “The Only Way Is Up?: Economic Mobility in Malaysia in the 21st Century.” *Policy Research Working Paper 10991*. World Bank: Washington D.C. <http://hdl.handle.net/10986/42490>.

Rongen, G., Ahmad, Z.A., Lanjouw, P. and Simler, K. (2024). “Regional and ethnic inequalities in Malaysian poverty dynamics”. *Journal of Economic Inequality*, 22: 101-130. <https://doi.org/10.1007/s10888-023-09582-w>

Appendix 1 – Additional tables

Table 3 Sample of panel income and consumption correlations and their source

Country	Start interval	End interval	Rho (standardized)	Income dummy	Study
Austria	2010	2012	0.770	1	Colgan (2023)
Bosnia and Herzegovina	2001	2004	0.521	0	Dang & Lanjouw (2023)
Bulgaria	2014	2016	0.686	1	Colgan (2026)
Czechia	2014	2016	0.822	1	Colgan (2026)
Ethiopia	2010	2012	0.430	0	Bentze & Wollburg (2024)
Ethiopia	2012	2014	0.509	0	Bentze & Wollburg (2024)
Ethiopia	2017	2021	0.542	0	Bentze & Wollburg (2024)
France	2010	2012	0.810	1	Colgan (2023)
France	2014	2016	0.784	1	Colgan (2026)
Greece	2010	2012	0.630	1	Colgan (2023)
Greece	2014	2016	0.661	1	Colgan (2026)
Hungary	2014	2016	0.698	1	Colgan (2026)
Ireland	2010	2012	0.680	1	Colgan (2023)
Italy	2010	2012	0.720	1	Colgan (2023)
Italy	2014	2016	0.780	1	Colgan (2026)
Lao PDR	2002	2007	0.604	0	Dang & Lanjouw (2023)
Niger	2011	2014	0.649	0	Bentze & Wollburg (2024)
Nigeria	2010	2012	0.566	0	Bentze & Wollburg (2024)
Nigeria	2012	2015	0.604	0	Bentze & Wollburg (2024)
Nigeria	2015	2018	0.665	0	Bentze & Wollburg (2024)
Norway	2010	2012	0.750	1	Colgan (2023)
Peru	2004	2005	0.749	0	Dang & Lanjouw (2023)
Peru	2004	2006	0.790	0	Dang & Lanjouw (2023)
Peru	2005	2006	0.749	0	Dang & Lanjouw (2023)
Poland	2010	2012	0.770	1	Colgan (2023)
Poland	2014	2016	0.713	1	Colgan (2026)
Slovak Republic	2010	2012	0.780	1	Colgan (2023)
Spain	2010	2012	0.700	1	Colgan (2023)
Sweden	2010	2012	0.800	1	Colgan (2023)
Tanzania	2008	2010	0.669	0	Bentze & Wollburg (2024)
Tanzania	2010	2012	0.699	0	Bentze & Wollburg (2024)
Tanzania	2012	2014	0.681	0	Bentze & Wollburg (2024)
Tanzania	2014	2019	0.687	0	Bentze & Wollburg (2024)

Uganda	2009	2010	0.658	0	Bentze & Wollburg (2024)
Uganda	2009	2011	0.637	0	Bentze & Wollburg (2024)
Uganda	2010	2011	0.547	0	Bentze & Wollburg (2024)
Uganda	2011	2013	0.546	0	Bentze & Wollburg (2024)
Uganda	2013	2015	0.643	0	Bentze & Wollburg (2024)
Uganda	2015	2018	0.700	0	Bentze & Wollburg (2024)
Uganda	2018	2019	0.654	0	Bentze & Wollburg (2024)
United Kingdom	2010	2012	0.650	1	Colgan (2023)
United States	2005	2007	0.760	1	Dang & Lanjouw (2023)
United States	2007	2009	0.820	1	Dang & Lanjouw (2023)
Viet Nam	2004	2006	0.810	0	Dang & Lanjouw (2023)
Viet Nam	2006	2008	0.780	0	Dang & Lanjouw (2023)
Note: all correlation coefficients rho are for the logarithm of household income or consumption and have been standardized to a two-year interval. Data sources: see column Study above. For Colgan (2023, 2026) and Dang & Lanjouw (2023), data were taken directly from the publication or obtained from the author. For Bentze & Wollburg (2024), the rhos are authors' calculations based on this harmonized dataset of the LSMS-ISA panel; we excluded one outlier from our calculations: a household observation in Uganda Wave 3 whose consumption expenditure was more than twenty standard deviations away from the mean in that wave.					

Table 4 Consumption model Synthetic Panel Tanzania

Dependent variable: Log of real annual hh pc consumption (TZS, 2017 prices)				
	Model 1		Model 2	
	Wave 4	Wave 5	Wave 4	Wave 5
Age of head in wave 4	-0.024 (0.020)	-0.036 (0.027)	-0.003 (0.019)	-0.007 (0.024)
Age of head squared	0.000 (0.000)	0.000 (0.000)	0.000 (0.000)	0.000 (0.000)
Head is female	0.012 (0.051)	-0.046 (0.066)	-0.033 (0.043)	-0.032 (0.060)
Head has formal education	0.419 (0.057)***	0.214 (0.084)**		
Urban household	0.716 (0.057)***	0.763 (0.066)***	0.435 (0.057)***	0.431 (0.076)***
Some primary education			0.146 (0.065)**	0.167 (0.087)*
Completed primary			0.363 (0.055)***	0.191 (0.072)***
Secondary			0.667 (0.073)***	0.665 (0.109)***
Tertiary or diploma			1.066 (0.099)***	1.193 (0.126)***
Region Arusha			0.506 (0.118)***	-0.036 (0.164)
Region Kilimanjaro			0.632 (0.119)***	0.178 (0.145)
Region Tanga			0.572 (0.119)***	-0.019 (0.155)
Region Morogoro			0.181 (0.151)	0.055 (0.205)
Region Pwani			0.794 (0.201)***	0.098 (0.158)
Region Dar Es Salaam			0.764 (0.102)***	0.344 (0.136)**
Region Lindi			0.312 (0.233)	0.044 (0.144)
Region Mtwara			0.293 (0.151)*	-0.046 (0.139)
Region Ruvuma			0.145 (0.174)	-0.115 (0.136)
Region Iringa			0.302 (0.160)*	0.090 (0.159)
Region Mbeya			0.435 (0.113)***	-0.004 (0.137)
Region Singida			0.246 (0.131)*	0.014 (0.171)
Region Tabora			0.103 (0.102)	-0.205 (0.134)
Region Rukwa			-0.539	-0.464

			(0.139)***	(0.188)**
Region Kigoma			-0.080	-0.679
			(0.111)	(0.184)***
Region Lakes			0.213	-0.117
			(0.100)**	(0.128)
Region Mara			0.281	-0.266
			(0.148)*	(0.183)
Region Manyara			0.293	-0.030
			(0.177)*	(0.269)
Region Njombe			0.315	0.187
			(0.121)***	(0.204)
Region Katavi			-0.048	-0.361
			(0.146)	(0.185)*
Region Simiyu			-0.053	-0.374
			(0.138)	(0.220)*
Region Geita			0.013	-0.264
			(0.165)	(0.171)
Region Zanzibar			0.096	-0.195
			(0.114)	(0.148)
Constant	13.547	13.876	12.857	13.307
	(0.410)***	(0.542)***	(0.387)***	(0.476)***
R^2	0.291	0.201	0.459	0.369
N	1,168	1,205	1,157	1,143

* $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$

Sample is limited to synthetic panel age range: 25-60 in wave 4. The omitted education dummy is No formal education, the omitted region dummy Dodoma.

Appendix 2 – Alternative standardization approach

With respect to the standardization of observed correlation coefficients to a two-year interval, Colgan (2026) discusses two approaches that are feasible in our context. Colgan notes (2026, p. 10):

Both the logarithmic and log-linear extrapolations provide a good approximate of $\rho_{y_{i1}y_{i2}}$ for most country settings and time periods. Unsurprisingly both methods produce very similar predictions in the short-run, but as the panel is extended the difference between the two approaches widens.

Based on available New Zealand data that enables comparison of the two approaches (compiled by Ball (2016)), Colgan concludes that estimating a logarithmic trend for the decline of rho over time provides the best point estimates.

Nonetheless, as a sensitivity analysis, this appendix assesses the distribution of our sample of correlation coefficients when standardized to a two-year period based on an assumed log-linear relationship. This relation can be written as follows, denoting the interval length in years by δ and the fixed depreciation factor, expected to be negative, by λ :

$$\rho_{\delta} = \rho_{\delta-1} + \lambda \rho_{\delta-1} = (1 + \lambda) \rho_{\delta-1} ,$$

or more generally as:

$$\rho_{\delta} = (1 + \lambda)^{\tau} \rho_{\delta-\tau} ,$$

where $\delta - \tau$ is the length of a shorter interval overlapping with the longer interval of length δ , again in years.

For each instance of overlapping intervals in our sample of correlations, we can then calculate the depreciation factor as follows:

$$\lambda_i = \left(\frac{\rho_{i,\delta}}{\rho_{i,\delta-\tau}} \right)^{\frac{1}{\tau}} - 1 .$$

The mean over all these instances would then be our estimate of the depreciation factor. Empirically, we observe some cases of higher correlations over longer intervals, but on average correlations decrease over time, such that our estimates are indeed negative.²⁸

Subsequently, we can use the following expression to standardize coefficients for an interval of length δ to a two-year interval:

$$\rho_2 = (1 + \hat{\lambda})^{2-\delta} \rho_\delta.$$

Results: distribution of standardized coefficients

The tables below provide summary statistics on the resulting distributions, which can be compared to the distribution of the log-trend standardized coefficients (top table, copied from Table 1 in the main text). Note that the income sample consists fully of two-year intervals, so it is not affected by the method of standardization. Comparing the three sets of statistics, we can conclude that the choice of standardization method only marginally affects the distribution of our sample of correlation coefficients. The median of the consumption correlation sample is the only parameter that shows a somewhat clearer shift, albeit a small one.

Log-trend

Weighted							
	mean	median	sd	min	max	N	Countries
Consumption	0.639	0.649	0.100	0.430	0.810	25	9
Income	0.737	0.750	0.058	0.630	0.822	20	15
Full sample	0.700	0.700	0.090	0.430	0.822	45	24

Log-linear (all depreciation factors)

Weighted							
	mean	median	sd	min	max	N	Countries
Consumption	0.631	0.629	0.111	0.430	0.810	25	9
Income	0.737	0.750	0.058	0.630	0.822	20	15
Full sample	0.697	0.700	0.097	0.430	0.822	45	24

²⁸ For the income coefficients, the mean depreciation factor is -0.056 when including positive occurrences, and -0.067 when we exclude them. For the consumption coefficients, these numbers are -0.034 and -0.048, respectively. Estimates have been weighted by the number of country survey intervals in our sample. One explanation for observing a positive factor is that a shock occurred that made the correlation more random in the short run. Then, after the passing of more time, incomes return to the pre-shock ordering, leading to a higher correlation. Another explanation could be the presence of price cycles of agricultural or raw materials.

Log-linear (only negative depreciation factors)

Weighted							
	mean	median	sd	min	max	N	Countries
Consumption	0.637	0.639	0.107	0.430	0.810	25	9
Income	0.737	0.750	0.058	0.630	0.822	20	15
Full sample	0.700	0.699	0.093	0.430	0.822	45	24

Table 5 Comparison of distributions standardized coefficients

Results: adjustment to actual survey interval in synthetic panel

In the process of creating a synthetic panel based on two cross-sectional surveys, the standardization approach also affects how the rho in our bootstrap sample is adjusted to account for the length of the interval between those surveys.²⁹ To obtain an idea of the sensitivity of poverty transition estimates to that adjustment, Figure 13 below shows how a particular sampled rho is adjusted under the different methods, for various survey intervals.

This demonstrates that for shorter intervals, differences in adjustments are small. For consumption measures in general, which method of standardization is applied will have little effect on transition estimates. For longer intervals and income measures notably, there is a large difference in the size of adjustment, as the depreciation factor for income correlations is larger than the one for consumption. For example, when adjusting a two-year income correlation coefficient of 0.8 to a survey interval of eight year, this coefficient will be adjusted to 0.69 using the log-trend method, while the log-linear method will yield a correlation of 0.57 or 0.53 depending on the whether positive depreciation factors are included. Using a lower correlation coefficient in estimating poverty transitions will result in a higher estimated level of mobility (in and out of poverty).

Intuitively, it is not obvious that the correlation should continue to decrease at an equal rate every year. In that sense, the log-trend decline seems more plausible. Moreover, we have almost no data on longer intervals for income measures: most correlations are for intervals of one, two or three years, and only one observation on a four year interval, which is the maximum length in our sample.

²⁹ Recall that all correlation coefficients in the bootstrap sample have been standardized to a two-year interval.

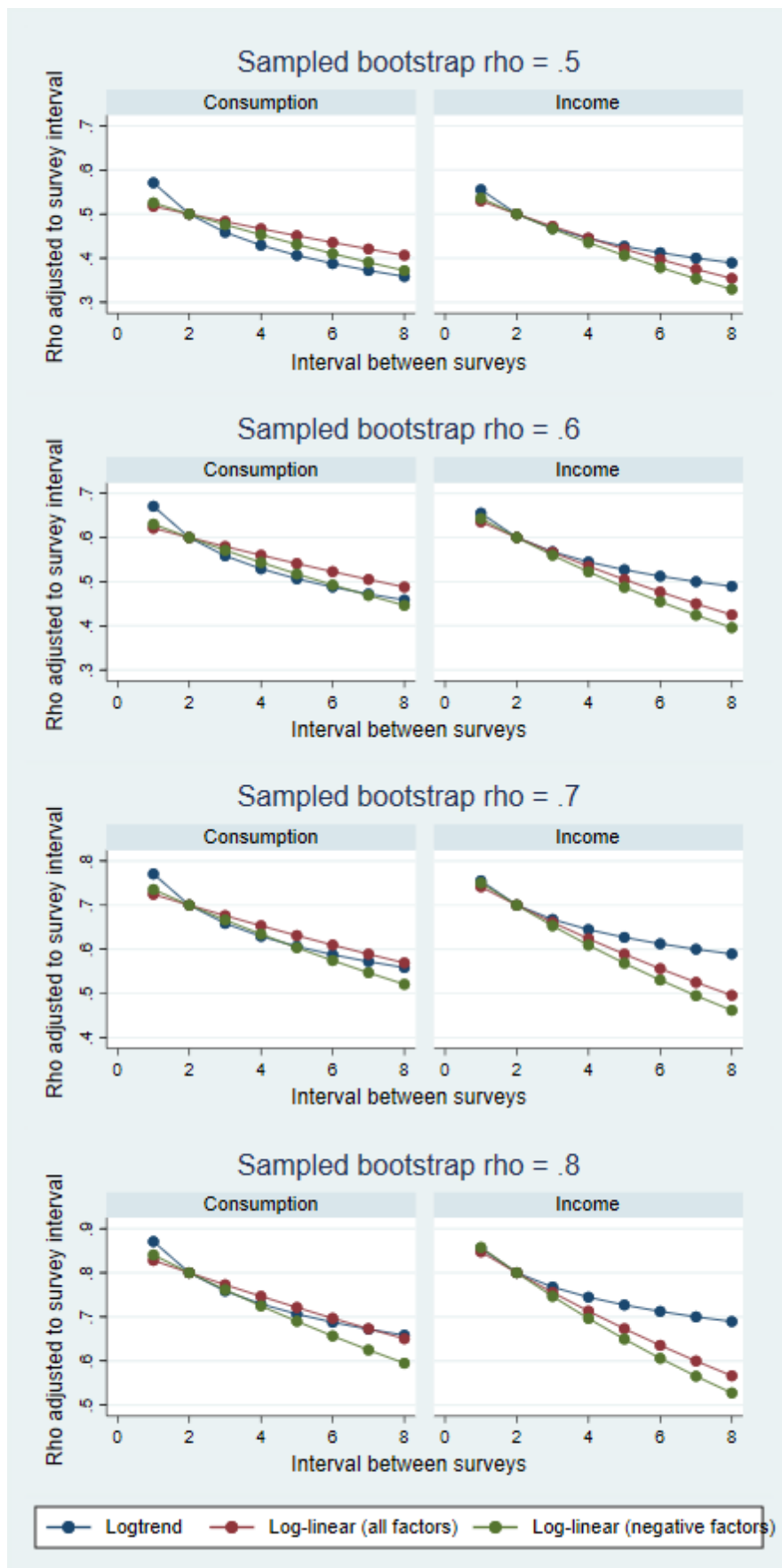


Figure 13 Comparison of adjustment to survey interval

Conclusion

The distribution of our sample of standardized coefficients is not affected much by which method is used to standardize observed coefficients to a two-year interval. However, in the construction of a synthetic panel, the method chosen may significantly affect transition results for longer intervals between surveys based on income measures. In this case, the log-trend adjustment seems more plausible. For consumption measures and shorter intervals, differences between the methods are small. These findings aligns with Colgan's assessment, which was for income measures. Hence, we consider the log-trend approach for standardization suitable and preferred.