

rivista di statistica ufficiale

REVIEW OF OFFICIAL STATISTICS

n. 1-2-3
2025

In this issue:

Monthly estimates for the Labour Force Survey:
an experimental study

*Alessio Guandalini, Silvia Loriga,
Mirjana Ogrizovi-Brasanac*

Non-response in the Italian Building Permits Survey:
a comparative analysis of imputation methods

Francesca Tuzi

Italian Tourism Satellite Account:
methodological approach

Sandra Maresca, Ilaria Piscitelli

rivista di statistica ufficiale

REVIEW OF OFFICIAL STATISTICS

n. 1-2-3
2025

In this issue:

Monthly estimates for the Labour Force Survey:
an experimental study

*Alessio Guandalini, Silvia Loriga,
Mirjana Ogrizovi-Brasanac*

5

Non-response in the Italian Building Permits Survey:
a comparative analysis of imputation methods

Francesca Tuzi

45

Italian Tourism Satellite Account:
methodological approach

Sandra Maresca, Ilaria Piscitelli

75

Editor:

Patrizia Cacioli

Scientific committee**President:**

Gian Carlo Blangiardo

Members:

Corrado Bonifazi	Vittoria Buratta	Ray Chambers	Francesco Maria Chelli
Daniela Cocchi	Giovanni Corrao	Sandro Cruciani	Luca De Benedictis
Gustavo De Santis	Luigi Fabbri	Piero Demetrio Falorsi	Patrizia Farina
Maurizio Franzini	Saverio Gazzelloni	Giorgia Giovannetti	Maurizio Lenzerini
Vincenzo Lo Moro	Stefano Menghinello	Roberto Monducci	Gian Paolo Oneto
Roberta Pace	Alessandra Petrucci	Monica Pratesi	Michele Raitano
Maria Giovanna Ranalli	Aldo Rosano	Laura Terzera	Li-Chun Zhang

Editorial board**Coordinator:**

Nadia Mignolli

Members:

Ciro Baldi	Patrizia Balzano	Federico Benassi	Giancarlo Bruno
Tania Cappadozzi	Anna Maria Cecchini	Annalisa Cicerchia	Patrizia Collesi
Roberto Colotti	Stefano Costa	Valeria De Martino	Roberta De Santis
Alessandro Faramondi	Francesca Ferrante	Maria Teresa Fiocca	Romina Fraboni
Luisa Franconi	Antonella Guarneri	Anita Guelfi	Fabio Lipizzi
Filippo Moauro	Filippo Oropallo	Alessandro Pallara	Laura Peci
Federica Pintaldi	Maria Rosaria Prisco	Francesca Scambia	Mauro Scanu
Isabella Siciliani	Marina Signore	Francesca Tiero	Angelica Tudini
Francesca Vannucchi	Claudio Vicarelli	Anna Villa	

Editorial support: Claudio Bava, Alfredina Della Branca, Marco Farinacci, Manuela Marrone

rivista di statistica ufficiale/review of official statistics

n. 1-2-3/2025

Four-monthly Journal: formerly registered at the Court of Rome, Italy (N. 339/2007 of 19th July 2007).

e-ISSN 1972-4829

p-ISSN 1828-1982

© 2025

Istituto Nazionale di Statistica

Via Cesare Balbo, 16 – Roma



Unless otherwise stated, content on this website is licensed under a Creative Commons License - Attribution - 4.0.

<https://creativecommons.org/licenses/by/4.0/deed.it>

Data and analysis from the Italian National Institute of Statistics can be copied, distributed, transmitted and freely adapted, even for commercial purposes, provided that the source is acknowledged.

No permission is necessary to hyperlink to pages on this website. Images, logos (including Istat logo), trademarks and other content owned by third parties belong to their respective owners and cannot be reproduced without their consent.

The views and opinions expressed are those of the authors and do not necessarily reflect the official policy or position of the Italian National Institute of Statistics - Istat.

The Scientific Committee, the Editorial Board and the authors would like to thank the anonymous reviewers (at least two for each article, on a voluntary basis and free of charge, with a double-anonymised approach) for their comments and suggestions, which enhanced the quality of this issue of the Rivista di statistica ufficiale/Review of official statistics.

Monthly estimates for the Labour Force Survey: an experimental study

Alessio Guandalini¹, Silvia Loriga¹, Mirjana Ogrizovi-Brasanac²

Abstract

The Labour Force Survey (LFS), conducted by the Statistical Office of the Republic of Serbia (SORS), provides quarterly labour market figures. The applied sample design allows the calculation of monthly estimates using a longitudinal sample structure. This paper presents the methodology for calculating monthly estimates of the main labour market indicators that comply with Eurostat requirements. The procedure for producing the monthly estimates is illustrated using data for the 4th quarter of 2019.

Keywords: Calibration estimator (CAL estimator), longitudinal constraints, Regression Composite Estimator, Adjusted Regression Estimator, sampling variance.

DOI: https://doi.org/10.1481/ISTATRIVISTASTATISTICAUFFICIALE_1-2-3.2025.01.

1 Alessio Guandalini (alessio.guandalini@istat.it); Silvia Loriga (siloriga@istat.it), Italian National Institute of Statistics - Istat.

2 Mirjana Ogrizovi-Brasanac (m.ogrizovic-brasanac@stat.gov.rs), Statistical Office of the Republic of Serbia - SORS.

The views and opinions expressed are those of the authors and do not necessarily reflect the official policy or position of the Italian National Institute of Statistics - Istat.

The authors would like to thank the anonymous reviewers for their comments and suggestions, which enhanced the quality of this article.

1. Introduction

Timely labour market information is essential to the national economy. It enables critical analysis of the workforce, including its behaviour, supply and demand, structural changes, and the impact of policies on it. In the European Union, official labour market information is collected through the Labour Force Survey (LFS). The LFS is based on Regulation (EU) 2019/1700³ of the European Parliament and Council. The Regulation establishes a common framework for European statistics relating to individuals and households. It also sets out the survey's main methodological requirements, content, and operational constraints. The main goals are the harmonisation of the LFS across the European Union, the improvement of data quality, and the comparability of results.

The survey is designed as a quarterly sample survey. The requirements for the sampling design, the estimation process, and the microdata delivery apply to quarterly data. The detailed information on participation in the labour market, collected through a very large set of variables, is reliable when analysed on a quarterly and annual basis.

For short-term analytical needs, Regulation (EU) 2019/2241 requires European countries to submit monthly unemployment data to Eurostat, limited to a few indicators. The deadline for data transmission falls between 25 and 27 days after the calendar month. To submit the data within the required period, a two-stage procedure is typically required. Provisional monthly estimates are computed from a partial sample for each month of the quarter. Final monthly estimates are computed from the full monthly sample, once data for the entire quarter are collected and quarterly estimates are computed. Final estimates differ slightly from provisional ones and are calculated to be consistent with quarterly estimates.

The Serbian LFS fulfils the European requirements for calculating quarterly estimates, and the adopted sampling design also enables the estimation of monthly labour market indicators. In this study, the procedure used to compute the monthly estimates is tested on Serbian LFS data for the 4th quarter of 2019.

³ Regulation (EU) 2019/1700 of the European Parliament and of the Council of 10 October 2019, amending Regulations (EC) No 808/2004, (EC) No 452/2008 and (EC) No 1338/2008 of the European Parliament and of the Council, and repealing Regulation (EC) No 1177/2003 of the European Parliament and of the Council and Council Regulation (EC) No 577/98.

The steps are as follows: computation of sampling weights for provisional and final monthly estimates; estimation of monthly figures for the main labour market indicators (employed, unemployed, and inactive); calculation of standard errors for these estimates; and implementation of the statistical procedures in R.

The paper is organised as follows: Section 2 provides an overview of the Serbian LFS; Section 3 addresses the production of provisional estimates; Section 4 examines variance estimation for the adopted estimator; Section 5 is dedicated to consistency with the quarterly estimates and the production of final monthly estimates. Each of Sections 3-5 contains three subsections: methodological aspects, application to the Serbian LFS, and focus on the R procedures. Finally, Section 6 discusses the main results.

2. Sample design of the Serbia LFS 2019

The Serbian LFS is a quarterly household sample survey. The target population comprises all private (non-institutional) households and persons in the territory of the Republic of Serbia, excluding the region of Kosovo and Metohija, which constitutes the usual population. Persons in collective households are excluded. The survey population is restricted to households living in the 2011 Census enumeration areas that contained at least 20 households at the time of the 2011 Census. The survey population is 1.5% smaller than the target population.

The sampling design is a two-stage stratified sample: enumeration areas (EAs, i.e., clusters of households) represent the primary sampling units (PSUs), and households are the final sampling units (SSUs). EAs are stratified by settlement type (urban/other)⁴ and by territory at the NUTS 3 level (25 areas).

A sample of enumeration areas is selected at the first stage, with a probability proportional to the number of persons aged 15 and over. In each selected EA, a sample of 10 households is selected using simple random sampling.

A 2-(2)-2 rotation scheme is adopted for the sample of households. Each household participates in the survey four times over an 18-month period; i.e., each household is in the sample for two consecutive quarters, then out of the sample for the next two quarters, and back in the sample for two further consecutive quarters.

Each quarter includes four independent subsamples (rotating groups): one new rotating group (households interviewed for the first time), two rotating groups from the previous quarter (households receiving the second and fourth interviews), and one rotating group from three quarters earlier (households interviewed for the third time).

Under this rotation scheme, the expected overlap of the sample is 50% between two consecutive quarters and between the same quarter in two consecutive years. This improves the efficiency of estimating changes relative to the previous quarter and to the same quarter of the previous year.

4 Official statistics in Serbia do not include a specific definition of rural settlements. Instead, an “administrative-legal” criterion is applied that designates settlements as either “Urban” or “Other”. Urban settlements are recognised as such by an act of the local self-government, with all other settlements falling into the category of “Other”.

Furthermore, data from previous rounds can inform current estimates.

Each subsample (rotation group) comprises 4,810 households across 481 enumeration areas. In each enumeration area, 10 households are selected.

The continuous conduct of the survey requires the sample distribution over time. Each subsample, consisting of one rotating group, is uniformly and randomly distributed across the 13 weeks of the quarter⁵.

Data are collected using a mixed-mode design combining Computer-Assisted Personal Interviewing (CAPI) and Computer-Assisted Telephone Interviewing (CATI): the first interview is always conducted by CAPI, and subsequent interviews are mainly conducted by CATI, except for households not found in previous rounds (phone numbers were not collected, etc.), which are interviewed by CAPI.

Data collection is fast. The time period to which the information on labour market participation refers (the reference week) is randomly assigned to each household. After each reference week, interviewers have two weeks to collect data. Owing to the use of computer-assisted techniques and the speed of data collection, the data are affected by few errors and are available in a timely manner for treatment.

The quarterly estimates are calculated using weights obtained through calibration (Deville and Särndal 1992). In the calibration procedure, design weights (for each stratum, the weight is the reciprocal of the product of the inclusion probabilities at each stage) are first corrected for non-response. The data are based on current demographic estimates: the distribution of the population by gender (two groups) and five-year age classes (14 groups) at the NUTS 3 level (25 groups), and the distribution of households by the number of household members (six groups) at the NUTS 3 level. The ReGenesee package (Zardetto 2015) in R is used to compute the calibrated weights. The logit function is used as the distance function. The same final weight is assigned to each individual in the household to ensure consistent estimates for both households and individuals.

5 When the year is composed of 53 weeks instead of 52, the sample of one quarter (first or fourth) is distributed over 14 weeks. This occurs once every four years.

3. Provisional estimates

To ensure data are published no later than 30 days after the end of the observed month, it is customary to publish provisional monthly estimates first. Thirty days after the end of the corresponding quarter, these provisional monthly estimates are usually replaced by final monthly estimates that are consistent with the quarterly estimates.

The data collection phase for the Serbian LFS usually ends two weeks after the end of the month, and the data are available for processing five days later. Almost all units involved in each reference month are interviewed and can be included in the calculation of the provisional monthly estimates.

The dataset used to produce provisional monthly estimates must include the variables (without missing values) listed in Table 3.1.

Table 3.1 - List and features of the variables in the input file for computing monthly estimates

VARIABLE NAME	Description
KEY_INDIV	Key variable for the individual
KEY_HH	Key variable for the household
KEY_PSU	Key variable for the Enumeration Area (EA)
STRATUM	Stratum variable (NUTS2 * urban/other * rotation group)
SEX	Sex
AGE	Age
NUTS2	Region code
NUTS3	Province code
ROTGR	Rotation group
EMP	Occupational status: employed (1=yes; 0=no)
UNEMP	Occupational status: unemployed (1=yes; 0=no)
INAC	Occupational status: inactive (1=yes; 0=no)

Source: Authors' processing

3.1 Design weights

In the LFS sample, selected households from each stratum are randomly assigned to the 13 weeks of the quarter. Because of this sampling feature, the households surveyed in a given month can be treated as a random sample and used to produce reliable monthly estimates.

The computation of the monthly sample design weight (dk) is the starting point in the computation of the weights. The design weight for the household k ,

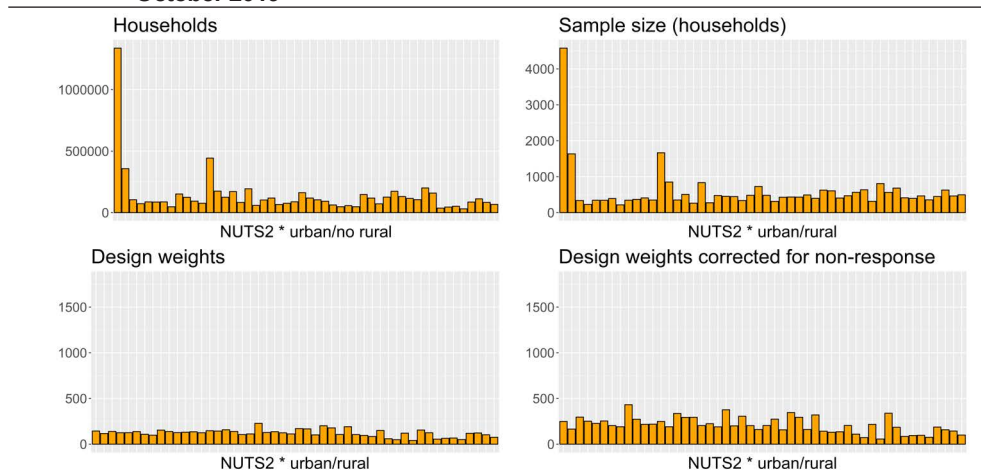
which is also shared by all individuals in the household, is equal to the inverse of the first-order inclusion probability, π_k . The first-order inclusion probability is computed at the NUTS 3 domain level for the type of settlement (urban/other), and is equal to the ratio of the planned number of households to be surveyed in the domain to the overall number of households in the population for the domain. This procedure approximates the real selection process for the monthly sample and is crucial for preventing extreme weights and enabling their calibration.

The design weight has to be adjusted for non-response among households. Non-response correction is applied at the domain level by multiplying the design weight by the inverse of the observed non-response rate in each domain. More complex weighting adjustments could be evaluated depending on the availability of auxiliary variables related to non-response (Bethlehem et al. 2011; Särndal and Lundström 2005; Little and Rubin 2002; Groves and Couper 1998).

The design weight corrected for non-response ($d_{k,nrc}$) represents the starting point for the monthly estimates and will be used in the work.

Figure 3.1 shows that the distribution of households in the sample is approximately equal to that in the domain. Looking at the values $d_{k,nrc}$ of the weights, it can be seen that in some domains their value after the correction for the non-response of households is significantly higher than the weights of the plan.

Figure 3.1 - Number of households, planned sample size, design weights, and design weights corrected for non-response at domain level (NUTS2 * urban/no urban). October 2019



Source: Authors' processing of the Serbian Labour Force Survey data

3.2 The Adjusted-Regression and the Regression Composite Estimators

3.2.1 Theoretical bases

The estimator proposed in the literature for producing monthly estimates is the regression composite (RC) estimator (Singh 1994 and 1996; Fuller 1999; Singh et al. 2001). The RC estimator incorporates auxiliary information from previous survey waves, making the estimates more stable. For this estimator, auxiliary information from previous survey waves is required for each unit in the sample. For the part of the sample for which these pieces of information are not available, data imputation should be performed. In the Canadian LFS, the RC estimator is used because there is a large overlap between the household samples for two successive months (5/6 of the sample overlaps), so imputation is needed for only 1/6 of households.

In the Serbian LFS, under the 2-(2)-2rotation scheme, the theoretical overlap between the samples referring to months m and $m-3$ is $\frac{1}{2}$, and the same is for the months m and $m-12$. When applying the RC estimator, it is necessary to impute half of the sample that does not overlap for months $m-3$ and $m-12$. Attrition and non-response reduce the planned overlap rate. In the 4th quarter of 2019, the overlap rate was 36% with the 3rd quarter of 2019, and 30% with the 4th quarter of 2018. The use of the RC estimator would introduce too much variability because it would require data imputation for 64-70% of the sample, respectively. Therefore, this estimator cannot be used, and a different estimator is recommended, such as the Adjusted-Regression (AR) estimator (Renssen and Nieuwenbroek 1997).

The AR estimator is known in the literature also as the Extended regression (E) estimator (Ballin et al. 2000; Rancourt 2001; Berger et al. 2009). For practical purposes and ease of use, its “calibrated version”, named AC estimator, has been defined (see, e.g., Särndal and Traat 2011; Ceccarelli and Guandalini 2012; Guandalini 2014). An AC estimator is used in this study to calculate the Serbian LFS monthly estimates.

The AC estimator lies on the calibration estimator (Deville and Särndal 1992), the most commonly used estimator in official statistics. The calibration estimator (CAL) may be written as follows

$$\hat{t}_{yCAL} = \sum_{k \in S} w_k y_k \quad (1)$$

and it represents a way to include auxiliary information in the estimation process based on a random sample (s) selected from a finite population. Through the calibration process, a system of weights, w_k , for the sample units is obtained, thereby making the sample consistent with the known distributions of structural variables in the reference population. This improves the efficiency and coherence of the estimates and can be used to produce all estimates. That is why it is widely used in official statistics (Särndal 2007).

The design weights or, the design weights corrected for non-response (as suggested in Deville et al. 1993), are modified so that estimates of totals for the P auxiliary variables ($x_1, x_2, \dots, x_P$) collected on the sample, are equal to the known totals, $\underline{t}_x = (t_{x1}, t_{x2}, \dots, t_{xP})$. The totals, $t_{x1}, t_{x2}, \dots, t_{xP}$, are usually retrieved from registers or administrative sources and are, by definition, assumed to be error-free.

The calibrated weights, w_k , are obtained by solving the following optimisation problem

$$\left\{ \begin{array}{l} \min_{w_k} \left\{ \sum_{k \in S} G(w_k, d_k) \right\} \\ \sum_{k \in S} w_k \underline{x}_k = \underline{t}_x \end{array} \right\} \quad (2)$$

where $\underline{x}_k$ is a vector of P auxiliary variables collected on the unit k . They are as close as possible to the d_k , an average sense, with respect to a pseudo-distance $G(\cdot)$, therefore, the CAL estimator keeps some good properties of the Horvitz-Thompson (HT) estimator, $\hat{t}_{yHT} = \sum_{k \in S} d_k y_k$ (Horvitz and Thompson 1952), which uses the design weights d_k . Furthermore, by definition, the calibrated weights enable to exactly reproduce the totals of auxiliary variables known in the population.

Different pseudo-distances $G(\cdot)$ can be used (see, e.g., Deville and Särndal 1992; Singh and Mohl 1996). Each metric determines a different system of calibrated weights with different ranges. The most used metric in official statistics is the truncated logarithmic distance,

$$\left(\frac{w_k}{d_k} - L\right) \ln\left(\frac{\frac{w_k}{d_k} - L}{1-L}\right) + \left(U - \frac{w_k}{d_k}\right) \ln\left(\frac{U - \frac{w_k}{d_k}}{U-1}\right)$$

because it prevents negative and extreme weights. The ratio w_k/d_k is defined between L (lower bound) and U (upper bound), and the bounds are usually set by the researcher.

When the Chi-Squared pseudo-distance is used, the CAL estimator is equivalent to the Generalised REGression (GREG) estimator (Bethlehem and Keller 1987; Cassel et al. 1976; Isaki and Fuller 1982; Särndal 1980; Särndal et al. 1989; Wright 1983)

$$\hat{t}_{y_{GREG}} = \hat{t}_{y_{HT}} + (\underline{t}_x - \hat{\underline{t}}_{x_{HT}})^T \hat{\underline{\beta}},$$

where the HT estimate is corrected by the difference between the vector of known totals and the vector of the totals of the same variables estimated on the sample with the HT estimator, weighted $\hat{\underline{t}}_{x_{HT}} = (\hat{t}_{x1_{HT}}, \hat{t}_{x2_{HT}}, \dots, \hat{t}_{xp_{HT}})$ by the vector of the regression coefficient $\hat{\underline{\beta}} = (\hat{\beta}_{x_1}, \hat{\beta}_{x_2}, \dots, \hat{\beta}_{x_p})^t$.

As long as the sample size and the population size are large, the CAL estimator is equivalent to the GREG estimator even when other pseudo-distances are used. This result is important because the sampling variance of CAL estimator is automatically equivalent to the sampling variance of GREG estimator:

$$\begin{aligned} AV(\hat{t}_{y_{CAL}}) &= \sum_{k \in S} \sum_{l \in S} \frac{\pi_{kl} - \pi_k \pi_l}{\pi_{kl}} (\hat{e}_k w_k)(\hat{e}_l w_l) \\ &\cong \hat{V}(\hat{t}_{y_{GREG}}) \end{aligned} \quad (3)$$

where π_{kl} are the second order inclusion probabilities and $\hat{e}$ are the estimated residuals of the superpopulation model, $\xi: y = \underline{X}\beta + e$. This model is implicitly assumed by the calibration estimator, therefore, it is possible to write the CAL estimator as a GREG estimator.

Because of the auxiliary totals are not all available from registers or administrative sources, survey-estimated totals can also be considered. In the case of surveys with rotating schemes with partial overlapping samples, their use it is even advocated. It enables to include in the estimation auxiliary variables collected on the same units on a different time. These variables are

usually highly correlated with the target variables and therefore, they help in increasing the efficiency of the estimates or at least in making them more stable and less volatile.

The expression of the CAL estimator in (1) holds, even if in this case Särndal and Traat (2011) suggest naming this estimator as AC estimator, pointing out the double nature of the considered auxiliary totals. The final weights can be obtained by solving the optimisation problem:

$$\left\{ \begin{array}{l} \min_{w_k} \left\{ \sum_{k \in S} G(w_k, d_k) \right\} \\ \sum_{k \in S} w_k \underline{u}_k = \tilde{\underline{t}}_u \end{array} \right. \quad (4)$$

where $\underline{u}_k = x_k, z_k$ is a vector of $P+M$ auxiliary variables in which the first P are $(x_1, x_2, \dots, x_P)$ and the last M are $(z_1, z_2, \dots, z_M)$. Totals of the x auxiliary variable, $\underline{t}_x = (t_{x1}, t_{x2}, \dots, t_{xP})$, are exactly known from registers or administrative data, while totals of the z auxiliary variables, $\tilde{\underline{t}}_z = (t_{z1}, t_{z2}, \dots, t_{zM})$, are estimated by the sample survey of another survey or previous waves of the same survey. The vector $\tilde{\underline{t}}_u = (\underline{t}_x, \underline{t}_z)$ consists of the $P+M$ components of which the last M are affected by sampling errors, while the first P are not.

The AC estimator has its equivalent GREG version, which is called the AR estimator

$$\hat{t}_{yAR} = \hat{t}_{yHT} + (\tilde{\underline{t}}_u - \underline{t}_{uHT}) \hat{\underline{\beta}} \quad (5)$$

where $\hat{t}_{uHT}$ are the HT estimates in the sample of the $P+M$ auxiliary variables and

$$\hat{\underline{\beta}} = (\hat{\underline{\beta}}_x, \hat{\underline{\beta}}_z)^t = (\hat{\beta}_{x1}, \hat{\beta}_{x2}, \dots, \hat{\beta}_{xP}, \hat{\beta}_{z1}, \hat{\beta}_{z2}, \dots, \hat{\beta}_{zM})^t$$

the vector of the regression coefficients in which the first P elements are related to the x auxiliary variables, while the last M are related to z auxiliary variables.

3.2.2. Application to the Serbian LFS

To produce monthly estimates that are efficient and stable, it is important to include in the estimation process a good set of auxiliary variables correlated with the target variables. This is crucial especially for monthly estimates because the sample size is $\frac{1}{3}$ of the quarterly sample, then such auxiliary variables are fundamental to increase the accuracy.

The AC estimator (or equivalently the AR) enables to entail in the estimation auxiliary variables for which the totals can be derived both from a register or estimated on a survey. Three kinds of auxiliary variables have been considered for the present purpose: household, individual and longitudinal auxiliary variables. In particular:

1. rotation group⁶;
2. number of components of the household⁷;
3. sex (SEX, M=male, F=female);
4. age class (AGE_CLASS, 0-14, 15-19, 20-24, 25-29, 30-34, 35-39, 40-44, 45-49, 50-54, 55-59, 60-64, 65-69, 70-74, 75+);
5. age class 25-74 (0=no, 1=yes);
6. employed a quarter before (1=yes, 0=no) [emp_1];
7. unemployed a quarter before (1=yes, 0=no) [unemp_1];
8. inactive a quarter before (1=yes, 0=no) [inac_1];
9. employed four quarters before (1=yes, 0=no) [emp_4];
10. unemployed four quarters before (1=yes, 0=no) [unemp_4];
11. inactive four quarters before (1=yes, 0=no) [inac_4].

Auxiliary variables in 1 and 2 are collected for the individuals in the sample but refer to household characteristics. The solution proposed by Lemaître and Dufour (1987) is adopted. In particular, the same value is assigned to all individuals belonging to the same household. Furthermore, since the auxiliary variables are originally constructed as dummy variables (1 = yes, 0 = no), they are divided by the number of household members. In this way, it is possible to simultaneously consider both household- and individual-level auxiliary variables. The rationale behind the use of these variables is that they can help address non-response and attrition problems, since these issues are strongly related to household composition and the number of interviews conducted within the household.

Auxiliary variables 3, 4, 5 are characteristics of each individual; variable 5 is useful to introduce at least this auxiliary information on age 25-74 at

6 Four modalities: first rotation group [RG1], second rotation group [RG2], third rotation group [RG3], fourth rotation group [RG4].

7 Six modalities: household with one component [CLASS_COMP1], two components [CLASS_COMP2], three components [CLASS_COMP3], four components [CLASS_COMP4], five components [CLASS_COMP5], six or more components [CLASS_COMP6].

the NUTS 3 level. Age and sex are very important variables for explaining the occupational status. Moreover, for NSIs, it is important to disseminate estimates consistent with the main population distributions. For household and individual auxiliary, variables totals are considered known.

Auxiliary variables 6, 7, 8, 9, 10 and 11 are the so-called longitudinal auxiliary variables. That is, they are the occupational status observed on the same individuals three months before and twelve months before. These auxiliary variables are obtained via a simple linkage operation with the data collected in the previous quarter and four quarters before. For this purpose, just the individual identifier (KEY_INDIV) is sufficient.

Due to the rotating scheme, this information is available for almost $\frac{1}{2}$ of the quarter, being $\frac{1}{2}$ the theoretical overlapping fraction between the two following quarters and between the same quarter in two following years. As already said, this leans the choice towards the AC (or equivalently the AR) estimator instead of the RC estimator.

The totals of the longitudinal auxiliary variables are based on estimates obtained in previous rounds of the LFS. They play as the z variables. Their use is very strategic for increasing the accuracy of monthly LFS estimates, because they are strongly correlated with the current occupational status, that is the target variable. However, since their totals are estimated, they introduce additional variability in the estimates that must be accounted for.

The implicit model assumed by the AC estimator that involves listed auxiliary variables is defined for each NUTS 2 level and can be written as:

(AGE_CLASS*SEX)+	Individual	x
RG1+RG2+RG3+RG4+	Household	
CLASS_COMP1+CLASS_COMP2+		z
CLASS_COMP3+CLASS_COMP4+	Individual	
CLASS_COMP5+CLASS_COMP6+		z
(SEX+SEX*CLASS2574)*NUTS3+	Longitudinal	
SEX*(emp_1+unemp_1+inac_1+ emp_4+unemp_4+inac_4)		z
- 1	no intercept	

It should be pointed that the model is without an intercept.

3.2.3 *Population totals*

In the LFS, besides the variables on the occupational status, auxiliary information on sex (male and female), age and geographical areas of residence (NUTS2: regions, NUTS3: provinces) are collected. These auxiliary variables are twofold important because, including them in the estimator, makes the weights consistent with the population distribution by sex and age in specific geographical areas; furthermore, they help in improving the estimates, because the occupational status is strongly correlated with sex, age and geographical areas of the individuals. According to the auxiliary variables in Section 3.2.2, five distributions based on current demographic estimates that are updated every quarter are observed:

- number of households by rotation group at NUTS2 level (block1, bl1);
- number of households by number of members at NUTS2 level (block2, bl2);
- population by sex and age class at NUTS2 level (block3, bl3);
- population by sex at NUTS3 level (block4, bl4);
- population by sex and age class 25-74 at NUTS3 level (block5, bl5).

3.2.4 *Longitudinal totals*

Totals of longitudinal auxiliary variables are used in this context to increase the efficiency of the estimates, reduce the volatility of the monthly series and lighten the impact of the small sample size of the monthly samples. They relate to the occupational status of individuals referred three months before (one quarter) and twelve months before (four quarters). They are usually highly correlated with the current occupational status (especially for employment), and they can reduce the sampling variance of the estimates.

This information is not available for all the individuals in the sample. The number of units in the sample for which these auxiliary variables are available depends on the rotation scheme adopted. In the Serbian LFS is 2-(2)-2. Theoretically, this represents 50% of the current sample. However, attrition and non-response reduce the actual overlapping rate (for instance, in the fourth quarter of 2019, the overlapping rate was 36% with the third quarter of 2019 and 30% with the fourth quarter of 2018). In this context, the use of the RC estimator

would introduce too much variability, as it would require the imputation of auxiliary variable values for approximately 70–74% of the sample, respectively. Therefore, the use of the AC estimator (and, equivalently, the AR estimator) is more appropriate. It is worth noting that, in this case, the totals of these variables, corresponding to the estimates of the occupational status referred to three and twelve months before, cannot be directly used. They have to be adjusted according to the overlapping rate (o) and the population change rate of individuals aged 15 years and over, between the periods (rc) and $\tilde{t}_z = o \text{ rc } \hat{t}_z$. Since the population totals are computed at the NUTS 2 level, all these quantities are computed at the same territorial level for the corresponding month. To determine the sampling variance of the estimates, it shall be ensured that these totals are estimates, so that the sample variance has to be adjusted, see Section 4.1.

3.3 Calibration

The use of the AC estimator consists of defining a calibration system as (4). It is preferable to use the weights already corrected for non-response. The distance function used is the truncated logarithmic function, which prevents negative and extreme weights. The auxiliary totals on which the calibrated weights are from current demographic estimations and from previous rounds of the survey. Overall, 344 auxiliary totals (86 auxiliary variables by 4 domains) are used in the calibration (Figure 3.2).

Figure 3.2 - Calibration system for the provisional estimates

NUTS2 (region)	b11 1-4	b12 5-10	b13 11-38	b14 39-56	tot_long 57-62 63-68		b15 69-86
1	Households by rotation group at NUTS2 level	Households by number of components at NUTS2 level	Population by age class and sex at NUTS2 level	Populati on by sex at NUTS3 level	Population by occupational status and sex 3 months before at NUTS2 Level	Population by occupational status and sex 12 months before at NUTS2 level	Population by age class 25-74 and sex at NUTS3 level
2							
3							
4							

Household auxiliary variables - Totals at the NUTS2 level available from the demographics estimates
 Individual auxiliary variables - Totals at the NUTS2 level available from the demographics estimates
 Individual auxiliary variables - Totals at the NUTS3 level available from the demographics estimates
 Longitudinal auxiliary variables - Totals at the NUTS2 level available from previous rounds of the LFS

Source: Authors' processing

Table 3.2 presents the comparison between the design weights, the design weights corrected for non-response and the calibrated weights. The non-response treatment and the calibration increase the variability of the sampling weights. It can be seen according to the Kish index (Kish 1965) which passes from 1.07 to 1.11 (by correction for non-response) and finally to 1.78 (by calibration).

Table 3.2 - Comparison between the design weights, the design weights corrected for non-response and the calibrated weights. October 2019

SAMPLING WEIGHTS	Min	Percentile 25%	Median	Mean	Percentile 75%	Percentile 90%	Percentile 95%	Percentile 99%	Max
Design weights (dk)	195	458	518	492	568	605	605	606	606
Design weights corrected by non-response (dk_cnr)	187	403	551	538	686	699	737	921	1,796
Calibrated weights (dk_cnr.cal)	40	295	447	613	754	1,228	1,612	2,449	6,197

Source: Authors' processing

3.4 Point estimates

The calibrated weights (w_k , dk_cnr.cal), named as provisional weights, are used for deriving the provisional estimates with expression (1). Auxiliary variables and totals described in Section 3.2.2 are used for the calculation of the monthly estimates for the number of employed, unemployed, active, and inactive people by sex and by region. Table 3.3 shows the provisional monthly estimates at the national level for October 2019.

Table 3.3 - LFS Provisional monthly estimates. October 2019 (a)

Occupational Status	Sex	Estimate (in thousand)
Active		3,285
Employed		2,970
Unemployed		315
Inactive		2,627
Active	Male	1,834
Active	Female	1,450
Employed	Male	1,670
Employed	Female	1,299
Unemployed	Male	163
Unemployed	Female	151
Inactive	Male	1,018
Inactive	Female	1,609

Source: Authors' processing

(a) Please note that the estimates are not the official estimates. They are just related to this experimental study.

3.5 R procedures

The methodology and all steps described in Section 3 have been implemented in R software. The R package ReGenesees (Zardetto 2015) is used. This package enables to perform the calibration, the computation of the estimates of parameters and of the variance for sample surveys. The main stages of the R script are explained below.

After preparing the monthly dataset and the blocks of auxiliary totals, it is necessary to define the design weights and the adopted sampling design.

The first step is to define the sampling design for the dataset (data) through the function `e.svydesign`. The sampling design is a stratified two-stage sample where the PSUs are the EAs (KEY_PSU) and the SSUs (KEY_HH) are the households. The PSUs are stratified by region, type of settlement (urban/no urban) and rotation group (STRATUM). The sampling fractions at the first stage (`fpc1`, fraction based on the persons in the stratum) and at the second stage (`fpc2`, fraction based on the households in the stratum for the selected EAs) must be indicated. The design weights corrected for non-response (`dk_nrc`) are used. The output is a design object (`des`) that contains all the relevant information on the sampling design. The following syntax is used.

Invoking the object `des` displays the message explaining that a two-stage stratified sample plan is defined and that for the month October 2019 there are 200 strata, 724 PSUs that are the enumeration areas and 4426 SSUs that are the households. In the calculation there will be 733 PSUs introduced, instead of 724 which are involved in the survey.

```
> des <- e.svydesign(data=data,
+                 ids=~KEY_PSU+KEY_HH,
+                 strata=~STRATUM,
+                 weights=~dk_nrc,
+                 fpc=~fpc1+fpc2,
+                 check.data = TRUE)
> des
Stratified 2 - Stage Cluster Sampling Design
- [200] strata
- [724, 4426] clusters
```

Call:

```
e.svydesign(data = data, ids = ~KEY_PSU + KEY_HH, strata = ~STRATUM,
           weights = ~dk_nrc, fpc = ~fpc1 + fpc2, check.data = TRUE)
```

The reduction of the PSUs is due to the collapsing step, a necessary pre-arrangement for properly managing the variance computation for the first and second stage of the sample selection (Rust and Kalton 1987). Strata that contain one PSU (enumeration area) and PSUs that contain one SSU (household) are collapsed (jointed with the corresponding, closest stratum or primary sampling unit) to enable the computation of the variance within the stratum and within the PSU.

The function `collapse.strata` is used for collapsing strata with one PSU. The parameter `block` defines inside which domain the strata can be collapsed, while an ad-hoc function manages the case in which just one SSU is observed in one PSU. The output `des.c` is an update of the design object `des` in which the additional information on collapsed strata is added. PSUs are grouped by an ad-hoc procedure.

```
> des.c <- collapse.strata(des,block=~NUTS3)
```

In order to perform the calibration, the following three steps are required: i) definition of the template for the auxiliary totals (`pop.template`); ii) filling the template; iii) calibration (`e.calibrate`).

The calibration model is defined with the `pop.template` function. The auxiliary variables to be used are listed and the level at which the totals are provided is defined (`partition`).

The output, `tot.cal_AC`, is an empty dataframe with 4 rows (one for each domain according to the variables specified in `partition`) and 86+1 columns (the header row plus one column for each variable included in the model).

```
> contrasts.off()

# Contrasts treatment has been switched off:

$contrasts
  unordered  ordered
"contr.off" "contr.off"

> tot.cal_AC <- pop.template(data=des.c,
+                             calmodel=~(AGE_CLASS:SEX)+
+                             RG1+RG2+RG3+RG4+
+                             CLASS_COMP1+CLASS_COMP2+
+                             CLASS_COMP3+CLASS_COMP4+
+                             CLASS_COMP5+CLASS_COMP6+
+                             (SEX+SEX:CLASS2574):COUNT_PROV+
+                             SEX:(emp_1+unemp_1+inac_1+
+                             emp_4+unemp_4+inac_4)-1
+                             partition=~NUTS2)
```

The command `contrasts.off()` considerably simplifies the procedure of compiling the output of `pop.template`. When it is used, it is sufficient to add, beside the blocks of auxiliary totals defined in the previous Sections, the R command with the `cbind` function (function for collapsing together two data frames when the number of rows in both data frames is equal).

```
> tot.cal_AC[, -1] <- cbind(bl1[, -1], bl2[, -1], bl3[, -1], bl4[, -1],
+                           tot_long[, -1], bl5[, -1])
```

Calibrated weights are obtained using the function `e.calibrate`. For the use of this function, it is necessary to assign first the values to its parameters, as it can be seen in the program below. So, the design is defined by the object `des.c`, the data frame with auxiliary totals is `tot.cal_AC`, the calibration model is `calmodel`, the level at which the totals are planned is defined with the function `partition`, `calfun` is a logit distance function, with the lower and upper bounds given by the parameter `bounds` and `aggregate.stage` is the function which defines the level on which calibrated weights are equal.

In particular, `calfun="logit"` and `bounds=c(0.1, 9.6)` means that the truncated logarithmic distance with $L=0.1$ and $U=9.6$ is used. It is important to point out that bound cannot be defined a priori and depends by the sample and the system of constraints. While, `aggregate.stage=2` means that all the individuals in the same household (the units of stage 2) will have the same calibrated weights.

```
> cal_AC <- e.calibrate(design=des.c,
+                      df.population=tot.cal_RC,
+                      calmodel=~(AGE_CLASS:SEX)+
+                        RG1+RG2+RG3+RG4+
+                        CLASS_COMP1+CLASS_COMP2+
+                        CLASS_COMP3+CLASS_COMP4+
+                        CLASS_COMP5+CLASS_COMP6+
+                        (SEX+SEX:CLASS2574):COUNT_PROV+
+                        SEX:(emp_1+unemp_1+inac_1+
+                          emp_4+unemp_4+inac_4)-1,
+                      partition=~NUTS2,
+                      calfun = "logit",
+                      bounds = c(0.1, 9.6),
+                      aggregate.stage = 2,
+                      force=FALSE)
> check.cal(cal_AC)
```

All Calibration Constraints (344) fulfilled (at tolerance level $\epsilon = 0.0000001$).

When the calibration is successful, the message of `check.cal(cal_AC)` states that all the calibration constraints (344 in this case) are fulfilled and a calibrated object is created. It can be seen as an update of `des.c`.

All the information related to the calibration are included in this new object and, the calibrated weights (w_k , `dk_ncr.cal`) are added to the dataset. The weights `dk_ncr.cal` are the provisional weights and they will be used for computing the provisional estimates.

The function `svystat` is used for calculating the provisional estimates. Since that one should calculate the estimates for the total number of active (`act`), employed (`emp`), unemployed (`unemp`) and inactive (`inac`) persons by region (NUTS2) and (SEX), the parameter `kind` (type of estimator) is set to "TM" and the parameter `estimator` is "Total".

The parameter `combo=2` enables the computation of the estimates at national level, by region, by sex and by region and sex.

```
> out <- svystat(design=cal_AC,  
+               kind="TM", estimator="Total",  
+               y=~act+emp+unemp+inac,  
+               by=~NUTS2+SEX,  
+               combo=2)
```


4. Sampling variance estimation

4.1 Theoretical bases

The AC estimator used for producing the monthly estimates includes auxiliary variables for which totals can be desumed both from a register, then exactly known, and from another survey or from a previous round of the survey, then affected by sampling errors. Therefore, its sampling variance depends on two sources of variability. The first source is the survey itself and entails stratification, stadification, sample size, sampling allocation. While, the second source is the variance of the estimated auxiliary totals.

Expression (3) related to the calculation of the CAL estimator variance, which does not include the variance of the estimated auxiliary totals, is not recommended for the calculation of the AC estimator variance (Berger et al. 2009; Dever and Vaillant 2010), as it can lead to underestimation of the sampling variance.

The suitable estimator for the variance is proposed in Berger et al. (2009), Renssen and Nieuwenbroek (1997), Ceccarelli and Guandalini (2012) and Guandalini (2014). These authors use the link between the CAL estimator and the GREG estimator as well as the AC estimator and the AR estimator.

The variance estimator of the AC estimator is derived starting from the expression of the AR estimator in (5), by adding and subtracting the expression t_u

$$\begin{aligned}\hat{t}_{y_{AR}} &= \hat{t}_{y_{HT}} + (\underline{\hat{t}}_u - \underline{\hat{t}}_u + \underline{t}_u - \underline{\hat{t}}_{u_{HT}})\underline{\hat{\beta}} \\ &= \hat{t}_{y_{HT}} + (\underline{t}_u - \underline{\hat{t}}_{u_{HT}})\underline{\hat{\beta}} - (\underline{t}_u - \underline{\hat{t}}_u)\underline{\hat{\beta}}\end{aligned}\quad (6)$$

where $\underline{t}_u = (\underline{t}_X, \underline{t}_z)$ is the vector of the auxiliary totals, assuming that totals of the z auxiliary variables are not affected by sampling errors. In (6) $\hat{t}_{y_{HT}} + (\underline{t}_u - \underline{\hat{t}}_{u_{HT}})\underline{\hat{\beta}}$ is the GREG estimator, $\hat{t}_{y_{GREG}}$ when both x and z variables are considered, assuming that the related totals are not affected by sampling error.

Therefore, the AR estimator can be written as

$$\begin{aligned}\hat{t}_{y_{AR}} &= \hat{t}_{y_{HT}} + (\underline{t}_u - \underline{\hat{t}}_{u_{HT}})\underline{\hat{\beta}} - (\underline{t}_u - \underline{\hat{t}}_u)\underline{\hat{\beta}} \\ &= \hat{t}_{y_{GREG}} - (\underline{t}_u - \underline{\hat{t}}_u)\underline{\hat{\beta}}\end{aligned}$$

and the estimator of the sampling variance is equal to

$$\begin{aligned}\widehat{V}(\hat{t}_{y_{AR}}) &= \widehat{V}(\hat{t}_{y_{GREG}}) + \widehat{V}\left((\underline{t}_u - \underline{\hat{t}}_u)\underline{\hat{\beta}}\right) - 2\widehat{COV}\left(\hat{t}_{y_{GREG}}, (\underline{t}_u - \underline{\hat{t}}_u)\underline{\hat{\beta}}\right) \\ &= \hat{A}_1 + \hat{A}_2 - \hat{A}_3 \\ &\cong \widehat{AV}(\hat{t}_{y_{AC}})\end{aligned}\tag{7}$$

Expression (7) approximates the variance of the AC estimator. The variance determined by

$$\hat{A}_1 = \widehat{V}(\hat{t}_{y_{GREG}})\tag{8}$$

represents the variance of the GREG estimator (expression 3). The variance of the auxiliary variables x and z is:

$$\hat{A}_2 = \underline{\hat{\beta}}^t \widehat{V}(\underline{\hat{t}}_u) \underline{\hat{\beta}}\tag{9}$$

$$= \underline{\hat{\beta}}_z^t \widehat{V}(\underline{\hat{t}}_z) \underline{\hat{\beta}}_z\tag{10}$$

as, only the totals $\underline{\hat{t}}_z$ are estimated. The component $\hat{A}_3$ of the AC estimator variance is:

$$\hat{A}_3 = 2\widehat{COV}\left(\hat{t}_{y_{GREG}}, \underline{\hat{t}}_u \underline{\hat{\beta}}\right)\tag{11}$$

$$= 2\sqrt{\widehat{V}(\hat{t}_{y_{GREG}})}[\widehat{CORR}(\hat{t}_{y_{GREG}}, \underline{\hat{t}}_z)]\sqrt{\widehat{V}(\underline{\hat{t}}_z)}\underline{\hat{\beta}}_z\tag{12}$$

and represents the covariance between $\hat{A}_1$ and $\hat{A}_2$. It takes into account that the totals of the z auxiliary variables have been estimated on a not negligible fraction of the units surveyed in the current round. The covariance in expression (11) is computed in terms of correlation (expression (12), according to relation $\widehat{COV}(X, Y) = \widehat{CORR}(X, Y)\sqrt{\widehat{V}(X)}\sqrt{\widehat{V}(Y)}$.

Expression (7) summarised all the steps needed for properly assessment of the variance of the AC (equivalently AR) estimator. Expression (8) is the sampling variance when considering all the auxiliary totals exactly known and that are error-free.

For the auxiliary totals that are estimates based on sample and which are affected by the sampling error, variance (8) has to be increased for the value given by expression (9), respectively by expression (10), as the totals of x auxiliary variables have no sampling error.

4.2 Application to the Serbian LFS

For the interpretation of results based on a random sample, the estimate of their accuracy is of great importance. This is particularly relevant for monthly estimates, as analysing their standard errors can determine which part of the estimate is based on the change of an indicator and which part is based on the sampling error (see, for example, Wolter 2007).

The sampling variance of the Serbian LFS monthly estimates is estimated using expression (7). As the auxiliary totals referred three and twelve months before, are adjusted according to the overlapping rate (o) and the population change rate of persons aged 15 and over, between the periods (rc), then the sample variance is adjusted as well. According to the relation valid for the variance $\widehat{V}(abX) = a^2b^2\widehat{V}(X)$, and the requirement that $\tilde{t}_z = orc\hat{t}_z$, the variance of auxiliary totals $\tilde{t}_z$ is calculated based on the expression $\widehat{V}(\tilde{t}_z) = \widehat{V}(orc\hat{t}_z) = o^2r^2c^2(\widehat{V})(\hat{t}_z)$.

In Table 4.1 to the estimates at the national level presented in Table 3.3 attached are: standard error (SE), the coefficient of variation (CV), the coefficient of variation in percents (CV%) and limits (LB and UB) for 95% confidence interval.

To show the influence of longitudinal auxiliary variables on the sample variance, all components in expression (7) are compared, as well as the sample variance in the case when only x variables are used in the calculation. Four variance estimates were observed:

CAL x is the estimate when only the x variables are considered, the CAL estimator in (1) is used and the sampling variance is computed using expression (3).

CALxz considers only the component $\widehat{A}_1$ in (7), the z auxiliary variables are considered and their totals are assumed to be exactly known, as well as those of the x auxiliary variables.

ACxz (covariance = 0) takes into account the sampling error introduced in the estimates because of the estimated totals of the z auxiliary variables, but it does not consider the covariance term and only $\widehat{A}_1$ and $\widehat{A}_2$ in (7) in the expression are accounted.

ACxz (covariance $\neq 0$) considers all the components in expression (7), that is $\widehat{A}_1$, $\widehat{A}_2$ and $\widehat{A}_3$.

Table 4.1 - LFS monthly estimates (provisional) with accuracy measures. October 2019

(a) (b)

Occupational Status	Sex	Estimate (in thousand)	SE	VAR	CV	CV%	LB	UB
Active		3,285	31.040	963.479	0.009449	0.944941	3,224	3,346
Employed		2,970	35.171	1236.981	0.011842	1.184209	2,901	3,039
Unemployed		315	22.930	525.775	0.072793	7.279304	270	360
Inactive		2,627	31.027	962.706	0.011811	1.181125	2,566	2,688
Active	Male	1,834	18.166	329.995	0.009905	0.990519	1,798	1,869
Active	Female	1,450	22.371	500.444	0.015428	1.542833	1,407	1,494
Employed	Male	1,670	21.480	461.371	0.012862	1.286241	1,628	1,712
Employed	Female	1,299	23.412	548.116	0.018023	1.802349	1,253	1,345
Unemployed	Male	163	13.907	193.414	0.085321	8.532076	136	190
Unemployed	Female	151	14.614	213.567	0.096781	9.678099	122	180
Inactive	Male	1,018	18.184	330.640	0.017862	1.786211	982	1,054
Inactive	Female	1,609	22.339	499.047	0.013884	1.388386	1,565	1,652

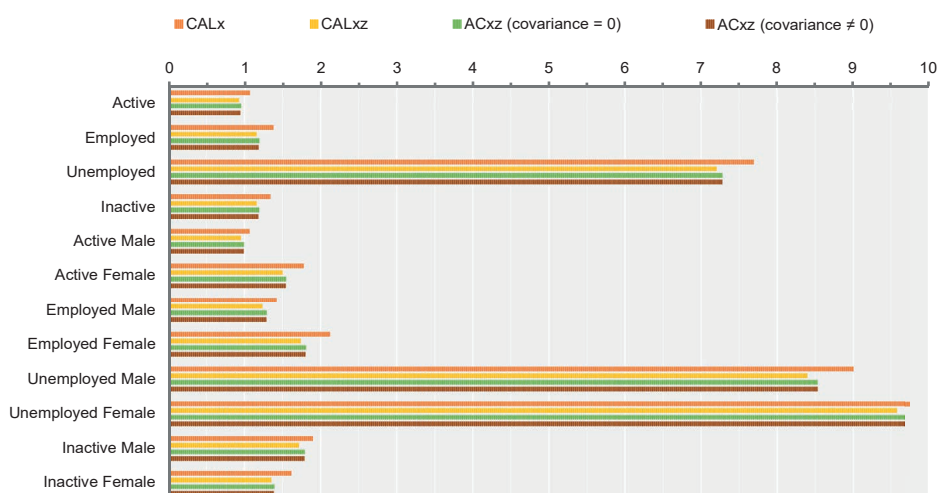
Source: Authors' processing

(a) (SE) standard error; (CV) coefficient of variation; (CV%) coefficient of variation in percents; (LB and UB) extremes of the confidence interval at 95%.

(b) Please note that the estimates are not the official estimates. They are just related to this experimental study.

Figure 4.1 shows the coefficients of variation for four variance estimates. It can be seen that including the z auxiliary variables in the estimation almost always increases the accuracy. Of course, the increase is higher when ignoring that the totals of the z auxiliary variables are estimates. But, even when the sampling variance of the AC estimator is properly managed, there is a general advantage in considering the longitudinal auxiliary variables. The possible disadvantage associated with some estimates is paid off by an important increase of the accuracy, for instance, for the unemployed status.

Figure 4.1 - Coefficient of variation of the LFS provisional estimates. October 2019
(percentage values)



Source: Authors' processing

4.3 R procedures

The computation of the sampling variance of the AC estimator in expression (7) has been implemented in R.

The software ReGenesees directly provides the component $\widehat{A}_1$. The output of the function `svystat`, besides the estimates (Y), provides the sampling variance (VAR). VAR is the variance of the CAL estimator (equivalent to the GREG estimator), that is, is computed with expression (3), assuming that totals for the x and z auxiliary variables are not affected by sampling error. Therefore, the components $\widehat{A}_2$ and $\widehat{A}_3$ have to be computed and included.

```
> out <- svystat(design=cal_RC,
+               kind="TM", estimator="Total",
+               y=~act+emp+unemp+inac,
+               by=~NUTS2+SEX,
+               combo=2)
```

For $\widehat{A}_2$ in expression (10) the regression coefficients $\underline{\beta}$ of the implicit superpopulation model $\xi : y = \underline{X} \underline{\beta} + e$ with $U = (X, Z)$ must be estimated on the sample. A model involving all the 344 auxiliary variables (296 x and

48 z) for each of the y variable (such as active, employed, ..., active male, active female, ...) has to be defined and for each of them, the regression coefficients have to be estimated. For making the computation faster, the regression coefficients are computed using the normalised calibrated weights (w) instead of the calibrated weights.

```
> b <- matrix(nrow=344,ncol=ncol(y))
> w <- (data$dk_nrc.cal/sum(data$dk_nrc.cal))*nrow(data)
> for(i in 1:ncol(y)){
+   model <- as.formula(paste(colnames(y)[i],
+                             "~((AGE_CLASS:SEX)+
+                             RG1+RG2+RG3+RG4+
+                             CLASS_COMP1+CLASS_COMP2+
+                             CLASS_COMP3+CLASS_COMP4+
+                             CLASS_COMP5+CLASS_COMP6+
+                             (SEX+SEX:CLASS2574):COUNT_PROV+
+                             SEX:(emp_1+unemp_1+inac_1+
+                             emp_4+unemp_4+inac_4)):NUTS2-1",
+                             sep=""))
+   beta <- lm(model,data,weights=w)$coefficient
+   b[,i] <- beta
+ }
> colnames(b) <- colnames(y)
> rownames(b) <- names(beta)
> b <- data.frame(b)
```

Only the regression coefficients related to the z auxiliary variables are needed. Since the sampling variance of $\tilde{t}_z$ has already been produced three and twelve months before, the component $\widehat{A}_1$ is obtained just by multiplying these quantities.

For $\widehat{A}_3$ in expression (11) the correlation should be estimated. There are several ways for estimating the correlation between two totals under complex sampling designs (see, e.g., Berger and Priam 2016). In this study, a simpler approximation is used in order to reduce the computational burden, and only the correlation between $\hat{t}_{yGREG}$ and $\tilde{t}_z$ was computed. $\widehat{V}(\hat{t}_{yGREG})$ and $\widehat{V}(\tilde{t}_{yGREG})$ have already been computed and $\widehat{A}_3$ is obtained just by multiplying these quantities as in expression (12).

```
> corrYZ <- matrix(nrow=ncol(y), ncol=ncol(z))
> for(i in 1:ncol(y)){
+   e <- y[,i]
+   for(j in 1:ncol(z)){
+     corrYZ[i,j] <- ifelse(is.na(cov(e,z[,j])),0,cov(e,z[,j]))
+   }
+ }
```

5. Final estimates consistent with the quarterly estimates

The main outputs of the LFS are the quarterly estimates. Quarterly data allow to produce very detailed information on the labour market, data are carefully checked and treated and the sample dimension allows to obtain more precise estimates related to the monthly estimates. For the easier interpretation of the results, and to stabilize monthly series it is important that the monthly estimates are consistent with the quarterly ones. Therefore, at the end of each quarter, there is a revision of the provisional monthly estimates and the dissemination of the final monthly estimates adjusted to be consistent with the quarterly ones.

Each quarterly estimate, $\hat{t}_y^Q$, can be seen as a weighted mean of the monthly estimates, $\hat{t}_y^{Mm}$ ($m=1, 2, 3$), where the weights are the number of the weeks of the month (4 or 5). For instance, in the 4th quarter 2019, the first month has 5 weeks, while the second and the third 4, then

$$\hat{t}_y^Q = \frac{5\hat{t}_y^{M1} + 4\hat{t}_y^{M2} + 4\hat{t}_y^{M3}}{13}$$

This result can be easily obtained through a calibration of the monthly estimates on the estimates referred to the quarter for the main labour market indicators. The consistence of the estimates is then guaranteed by definition. To avoid too much modification of the provisional monthly estimates most of the auxiliary variables already used for providing the provisional monthly weights are included in this calibration step. The auxiliary variables are mainly those related to the population distribution by sex, age classes and regions. The output is a system of weights that for each month enables to provide the final monthly estimates consistent with the quarterly ones.

5.1 Application to the Serbian LFS

To produce the final monthly estimates it is necessary to define a new calibration system as in expression (2), on which the calibrated weights that were used in computation for the provisional monthly estimates (provisional weights `dk_nrc.cal`) will be modified. All monthly data are treated together. Auxiliary variables included in the calibration process are the following:

12. sex (SEX, M=male, F=female);
13. age class (AGE_CLASS, 0-14, 15-19, 20-24, 25-29, 30-34, 35-39, 40-44, 45-49, 50-54, 55-59, 60-64, 65-69, 70-74, 75+);
14. month (month, M1=October, M2=November, M3=December);
15. age class 15-24 (0=no, 1=yes);
16. age class 25-74 (0=no, 1=yes);
17. number of the weeks in the month (nweek, 4, 5);
18. employed (emp=yes, 0=no) [emp];
19. unemployed (unemp=yes, 0=no) [unemp];
20. inactive (inac=yes, 0=no) [inac];
21. employed (emp*nweek/13=yes, 0=no) [Xemp];
22. unemployed (unemp*nweek/13=yes, 0=no) [Xunemp];
23. inactive (inac*nweek/13=yes, 0=no) [Xinac].

The auxiliary variables 1-9 refer to the observed month and serve to preserve the coherence of the monthly estimates with the corresponding reference population. This coherence was already assured by the provisional monthly weights and must be maintained in the final monthly weights.

The auxiliary variables 10-12 serve to reach the coherence between the monthly estimates of the main indicators and the corresponding quarterly estimates, through a weighted mean. A practical solution to obtain this coherence is a weighting of the auxiliary variables 7-9 by the number of the weeks in the observed month for obtaining auxiliary variables 10-12.

The implicit model assumed at the NUTS 2 level is:

(AGE_CLASS*SEX*month)+	Individuals
((Xemp+Xunemp+Xinac)*	occupational status
(CLASS1524+CLASS2574):SEX)	
-1	no intercept

The definition of this model enables to obtain monthly data for the occupational status (employed, unemployed and inactive), that is in accordance with corresponding monthly population, and all together they determine the parameter estimate that is equal to the already calculated one.

The totals used are (Figure 5.1):

- population by sex and age class at the NUTS2 level (block1, b11);
 - for the first month of the quarter (block1a, b11a);
 - for the second month of the quarter (block1b, b11b);
 - for the third month of the quarter (block1c, b11c);
- male 15-24 by status at NUTS2 level from quarterly estimates (block2 b12);
- female 15-24 by status at NUTS2 level from quarterly estimates (block3, b13);
- male 25-74 by status at NUTS2 level from quarterly estimates (block4, b14);
- female 25-74 by status at NUTS2 level from quarterly estimates (block5, b15).

Figure 5.1 - Calibration system for the final estimates

NUTS2 (region)	b11			b12	b13	b14	b15
	b11.a	b11.b	b11.c				
	1-28	29-56	57-84	85-87	88-90	91-93	94-96
1	Population by age class and sex at NUTS2 level Month 1	Population by age class and sex at NUTS2 level Month 2	Population by age class and sex at NUTS2 level Month 3	Male pop. 15-24 y.o. by occupational status at NUTS2 level	Female pop. 15-24 y.o. by occupational status at NUTS2 level	Male pop. 25-74 y.o. by occupational status at NUTS2 level	Female pop. 25-74 y.o. by occupational status at NUTS2 Level
2							
3							
4							

 Monthly auxiliary variables – Totals at NUTS2 level available from current demographic estimates.
 Quarterly auxiliary variables – Totals at NUTS2 level available from the quarterly estimates

Source: Authors' processing

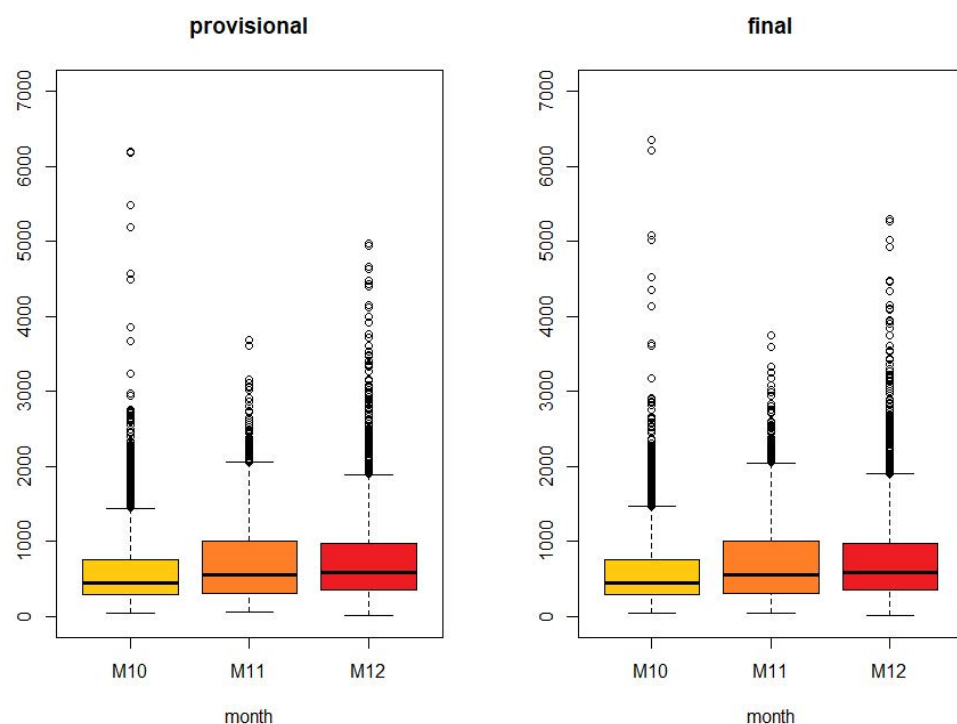
In this experimental study for the 4th quarter 2019, consistency of monthly and quarterly data is reached. The calibration process converges setting the limits $L=0.7$ and $U=1.2$ in logit distance function. The final weights (dk_cnr.cal.cal) are the calibrated weights obtained according to the provisional weights (dk_cnr.cal). This step of calibration slightly changes the provisional weights, which can be seen in the comparative overview in Table 5.1 and Figure 5.2.

**Table 5.1 - Comparison between the final weights of the three months of the quarter.
October-December 2019 (a)**

Weights	Month	Min	Perc. 25%	Mean	Median	Perc. 75%	Perc. 90%	Perc. 95%	Perc. 99%	Max
Provisional (dk_cnr.cal)	October	40	295	447	613	754	1,228	1,612	2,449	6,197
	November	49	302	544	730	1,004	1,609	1,903	2,480	3,688
	December	9	348	584	776	967	1,551	2,065	3,488	4,975
Final (dk_cnr.cal.cal)	October	35	293	447	613	759	1,227	1,591	2,532	6,349
	November	42	303	546	730	1,001	1,605	1,893	2,503	3,747
	December	9	348	582	776	969	1,543	2,014	3,542	5,292

Source: Authors' processing

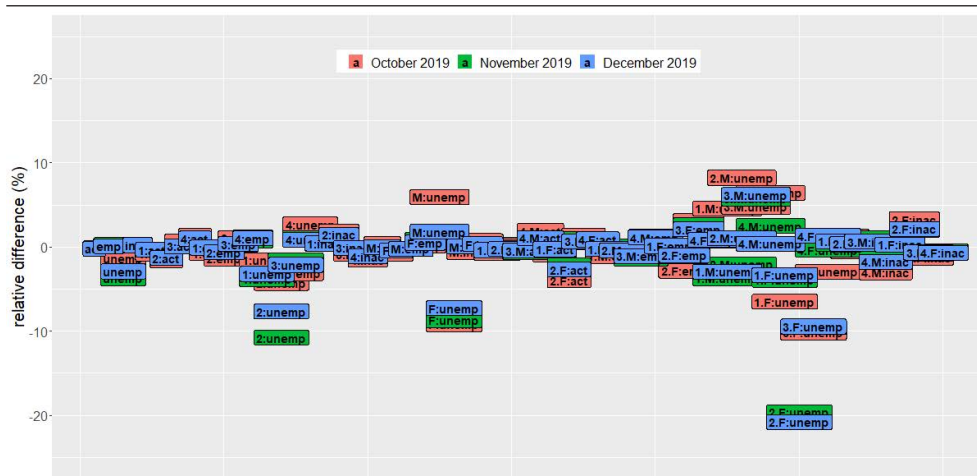
(a) Perc.=Percentile.

Figure 5.2 - Boxplot of the provisional weights (dk_cnr.cal) and the final weights (dk_cnr.cal.cal) for each month of the fourth quarter. Year 2019

Source: Authors' processing

Final monthly estimates are in average close to the provisional monthly estimates. This can be seen in Figure 5.3. Most of the estimates, i.e. 99 estimates over 180 estimates have a relative change lower than 1%, 64 estimates change less than 5%, while the remaining 17 estimates change more than 5%, and only the estimates of the number of unemployed women in region 2 are modified around the 20% in all months.

Figure 5.3 - Relative difference between final and provisional monthly estimates. October-December 2019 (percentage values)



Source: Authors' processing

Considering the accuracy of the final monthly estimates, it should be pointed out that only some estimates of totals are used in the calibration to achieve consistency with the quarterly data.

Contrary to the process, which was done for the provisional estimates, here their sampling error was not accounted for assessing the sampling variance of the final estimates.

Considering that quarterly estimates are reliable, it is assumed that this approximation has not a great impact on the sampling error of the final monthly estimates.

Comparing the coefficient of variation of the final monthly estimates with the coefficients of variation of the provisional monthly estimates, it can be


```

> contrasts.off()

# Contrasts treatment has been switched off:

$contrasts
  unordered   ordered
"contr.off" "contr.off"

> # definition of the calibration system
> tot.cal <- pop.template(data=des,
+                         calmodel=~(AGE_CLASS:SEX:month)
+                         +((Xemp+Xunemp+Xinac):(CLASS1524+CLASS2574):SEX) -1,
+                         partition=~NUTS2)
> tot.cal[, -1] <- tot[, -1]

```

Then the calibration can be performed by function `e.calibrate` explained in Section 3.5. The truncated logarithmic function, `calfun = "logit"` is used. The bounds for the weights L and U are set by the function `bounds = c(0.7, 1.2)`. The final weights are calculated to be equal for all the individuals in the same household, by the function `aggregate.stage = 2`.

```

> cal <- e.calibrate(design=des,
+                  df.population=tot.cal,
+                  calmodel=~(AGE_CLASS:SEX:month)+
+                  ((Xemp+Xunemp+Xinac):
+                  (CLASS1524+CLASS2574):
+                  SEX)
+                  -1,
+                  partition=~NUTS2,
+                  calfun = "logit",
+                  bounds = c(0.7, 1.2),
+                  aggregate.stage = 2,
+                  force=FALSE)
> check.cal(cal)
All Calibration Constraints (384) fulfilled (at tolerance level epsilon = 0.0000001).

```

The message of `check.cal(cal)` in the previous set of commands states that all the calibration constraints (384 in this case) are fulfilled and the calibrated object is created. The `cal` object contains the data for all the months of the quarter and the final weights to be used for providing the final estimates (`dk_ncr.cal.cal`). It means that provisional weights `dk_ncr.cal`, are calibrated.

After the calibration is performed, using the function `ext.calibrate`, the `des` object is defined and contains the information for the sample design plan for each monthly sample, the calibration model used for each month, and the calibration model used to achieve consistency with quarterly data.

In practice, this function is a combination of `e.svydesign` and `e.calibrate` and is demonstrated by the following set of commands.

```
> des <- ext.calibrated(data=cal$variables,
+                      ids=~KEY_PSUC+KEY_HH,
+                      strata=~STRATUM,
+                      weights=~dk_nrc,
+                      fpc=~fpc1+fpc2,
+                      check.data=TRUE,
+                      weights.cal=~dk_nrc.cal.cal,
+                      calmodel=~month:((AGE_CLASS:SEX)+
+                                       RG1+RG2+RG3+RG4+
+                                       CLASS_COMP1+CLASS_COMP2+
+                                       CLASS_COMP3+CLASS_COMP4+
+                                       CLASS_COMP5+CLASS_COMP6+
+                                       (SEX+SEX:CLASS2574):COUNT_PROV+
+                                       SEX:(emp_1+unemp_1+inac_1+
+                                           emp_4+unemp_4+inac_4)
+                                       )+
+                      ((Xemp+Xunemp+Xinac):
+                      (CLASS1524+CLASS2574):
+                      SEX)
+                      -1,
+                      partition=~NUTS2)
```

For computing the final monthly estimates, the object `des` is an input of the function `svystat`. The estimate of the totals for the number of active (act), employed (emp), unemployed (unemp) and inactive (inac) persons by region (NUTS 2) and sex (SEX) is performed by the parameter `kind="TM"` (Totals and Means) and the estimator="Total". Setting the parameter `group=~month` enables to immediately compute the estimates for all the three months. The output object is a list object, and each element contains the estimates for the observed month.

```
> out <- svystat(design=des,kind="TM",estimator="Total",
+               y=~act+emp+unemp+inac, by=~NUTS2+SEX, combo=2,
+               group=~month)
```

In the process of variance estimation, the impact of the quarterly estimates of totals in the last calibration step is ignored because it is assumed to be negligible, contrary to the impact of the longitudinal auxiliary totals. The procedure already shown in Section 4.3 has to be applied individually to each monthly data. The final estimates and the related measures of accuracy, such as sampling variance (VAR), standard error (SE), coefficient of variation (CV), coefficient of variation percent (CVpct) and the limits (LB, UB) of the 95% confidence intervals are computed and saved to be used for the calculation of the monthly estimates in the subsequent quarters.

6. Concluding remarks

The aim of this study was to define a methodology for calculating monthly Labour Force Survey estimates that incorporates data from previous survey waves (longitudinal data). In this way, the accuracy of the estimate would be improved, as the estimate is calculated on the basis of one-third of the quarterly sample size, while also using available historical data.

Considering the LFS sample plan, an estimation method was used that, in addition to the known auxiliary totals, introduces estimates of totals from previous waves of the survey into the calibration process (AC estimator). When defining the matrix of auxiliary calibration totals, it is necessary to correct the estimated totals according to the rate of overlap and the rate of change of the population aged 15 and over for the period between the previous and current survey. In addition, the AC method differs from the standard calibration procedure because the estimated variance includes the error of the estimated totals, which must be included in the estimates of the estimations of the accuracy.

The defined method refers to the estimates that would be calculated immediately after the end of the month and need to be published within the provided period. The study also defines the procedure for creating final monthly estimates after the end of the quarter, which are consistent with the quarterly estimates, differ slightly from the initial estimates, and are also published.

The practical goal of using the Serbian Labour Force Survey data to calculate quarterly estimates, as well as timely monthly estimates of key labour market indicators that are in accordance with Eurostat requirements, is achieved.

In calculating the provisional monthly estimates, 344 auxiliary variables were used, of which 296 relate to households and persons and their totals are known based on current demographic estimates, while 48 auxiliary variables were estimated based on the LFS three and twelve months earlier.

To calculate the final monthly estimates that are to be consistent with the quarterly estimates, 384 auxiliary variables were used, of which 336 related to persons and were also used in the calibration step to calculate the provisional monthly estimates, while 48 auxiliary variables were estimated based on LFS quarterly data.

The calculated provisional and final monthly estimates presented in this paper are not official, but their accuracy is in line with the accuracy of the monthly estimates of European countries published for monthly LFS data. The provisional monthly estimates are very close in value to the final monthly estimates, and it can be concluded that the applied methodology is appropriate and can be used for the calculation of the monthly estimates, especially when the volatility of the monthly estimates is taken into account.

References

Ballin, M., P.D. Falorsi, e A. Russo. 2000. “Condizioni di coerenza e metodi di stima per le indagini campionarie sulle imprese”. *Rivista di statistica ufficiale/Review of official statistics*, Volume 2/2000: 31-52.

Berger, Y.G., J.F. Muñoz, and E. Rancourt. 2009. “Variance estimation of survey estimates calibrated on estimated control totals - An application to the extended regression estimator and the regression composite estimator”. *Computational Statistics & Data Analysis*, Volume 53, N. 7: 2596-2604. <https://doi.org/10.1016/j.csda.2008.12.011>.

Berger, Y.G., and R. Priam. 2016. “A simple variance estimator of change for rotating repeated surveys: an application to the European Union Statistics on Income and Living Conditions household surveys”. *Journal of the Royal Statistical Society, Statistics in Society, Series A*, Volume 179, N. 1: 251-272. <https://doi.org/10.1111/rssa.12116>.

Bethlehem, J.G., F. Cobben, and B. Schouten. 2011. *Handbook of Nonresponse in Household Surveys*. Hoboken, NJ, U.S.: John Wiley & Sons. <https://doi.org/10.1002/9780470891056.ch14>.

Bethlehem, J.G., and W.J. Keller. 1987. “Linear weighting of sample survey data”. *Journal of Official Statistics*, Volume 3, N. 2: 141-153.

Cassel, C.M., C.-E. Särndal, and J.H. Wretman. 1976. “Some Results on Generalized Difference Estimation and Generalized Regression Estimation for Finite Populations”. *Biometrika*, Volume 63, N. 3: 615-620. <https://doi.org/10.2307/2335742>.

Ceccarelli, C., and A. Guandalini. 2012. “Increasing the accuracy of It-SILC estimates through the use of auxiliary information from Labour Force Survey”. *Statistica Applicata - Italian Journal of Applied Statistics*, Volume 24, N. 1: 103-115.

Dever, J.A., and R. Valliant. 2010. “A comparison of variance estimators for poststratification to estimated control totals”. *Survey Methodology*, Volume 36, N. 1: 45-56. <https://www150.statcan.gc.ca/n1/en/pub/12-001-x/2010001/article/11251-eng.pdf?st=L2HC2935>.

Deville, J.-C., and C.-E. Särndal. 1992. "Calibration Estimators in Survey Sampling". *Journal of the American Statistical Association*, Volume 87, N. 418: 376-382. <https://doi.org/10.2307/2290268>.

Deville, J.-C., C.-E. Särndal, and O. Sautory. 1993. "Generalized Raking Procedures in Survey Sampling". *Journal of the American Statistical Association*, Volume 88, N. 423: 1013-1020. <https://doi.org/10.2307/2290793>.

Fuller, W.A. 1999. "The Canadian Regression Composite Estimation". *Unpublished Manuscript*.

Groves, R.M., and M.P. Couper. 1998. *Nonresponse in Household Interview Surveys*. Hoboken, NJ, U.S.: John Wiley & Sons. <https://doi.org/10.1002/9781118490082>.

Guandalini, A. 2014. *Coerenza ed ottimalità dello stimatore calibrato*. PhD dissertation.

Horvitz, D.G., and D.J. Thompson. 1952. "A Generalization of Sampling Without Replacement From a Finite Universe". *Journal of the American Statistical Association*, Volume 47, N. 260: 663-685. <https://doi.org/10.2307/2280784>.

Isaki, C.T., and W.A. Fuller. 1982. "Survey Design Under the Regression Superpopulation Model". *Journal of the American Statistical Association*, Volume 77, N. 377: 89-96. <https://doi.org/10.2307/2287773>.

Kish, L. 1965. *Survey sampling*. New York, NY, U.S.: John Wiley & Sons.

Little, R.J.A., and D.B. Rubin. 2002. *Statistical Analysis with Missing Data*. Hoboken, NJ, U.S.: John Wiley & Sons. <https://doi.org/10.1002/9781119013563>.

Lemaître, G., and J. Dufour. 1987. "An Integrated Method for Weighting Persons and Families". *Survey methodology*, Volume 13, N. 2: 199-207. <https://www150.statcan.gc.ca/n1/en/pub/12-001-x/1987002/article/14607-eng.pdf?st=FcCf9qHs>.

Rancourt, É. 2001. "La régression étendue : un ensemble de pratiques d'estimation qui poussent constamment la théorie". In Droesbeke, J.-J., et L. Lebart (sous la direction de). *Enquêtes, modèles et applications*: 334-343. Malakoff, France: Dunod.

Renssen, R.H., and N.J. Nieuwenbroek. 1997. "Aligning Estimates for Common Variables in Two or More Sample Surveys". *Journal of the American Statistical Association*, Volume 92, N. 437: 368-374. <https://doi.org/10.2307/2291482>.

Rust, K., and G. Kalton. 1987. "Strategies for Collapsing Strata for Variance Estimation". *Journal of Official Statistics*, Volume 3, N. 1: 69-81. <https://www.proquest.com/openview/9648e5a196880c3765a96af6e4b50ae2/1?pq-origsite=gscholar&cbl=105444>.

Särndal, C.-E. 2007. "The calibration approach in survey theory and practice". *Survey Methodology*, Volume 33, N. 2: 99-119. <https://www150.statcan.gc.ca/n1/en/pub/12-001-x/2007002/article/10488-eng.pdf?st=YEcKmKFm>.

Särndal, C.-E. 1980. "On π -Inverse Weighting Versus Best Linear Unbiased Weighting in Probability Sampling". *Biometrika*, Volume 67, N. 3: 639-650. <https://doi.org/10.2307/2335134>.

Särndal, C.-E., and S. Lundström. 2005. *Estimation in Surveys with Nonresponse*. Hoboken, NJ, U.S.: John Wiley & Sons.

Särndal, C.-E., B. Swensson, and J.H. Wretman. 1989. "The Weighted Residual Technique for Estimating the Variance of the General Regression Estimator of the Finite Population Total". *Biometrika*, Volume 76, N. 3: 527-537. <https://doi.org/10.2307/2336118>.

Särndal, C.-E., and I. Traat. 2011. "Domain Estimators Calibrated on Information from Another Survey". *Acta et Commentationes Universitatis Tartuensis de Mathematica*, Volume 15, N. 2: 43-60.

Singh, A.C. 1996. "Combining information in Survey Sampling by Modified Regression". In *Proceedings of the Section on Survey Research Methods*, Volume 1: 120-129. Alexandria, VA, U.S.: American Statistical Association.

Singh, A.C. 1994. "Sampling design based estimating functions for finite population totals. Invited paper". In *Abstracts of the Annual Meeting of the Statistical Society of Canada, Banff*: 48. Calgary, Alberta, Canada, 8-11 May 1994.

Singh, A.C., B. Kennedy, and S. Wu. 2001. "Regression Composite Estimation for the Canadian Labour Force Survey with a Rotating Panel Design". *Survey Methodology*, Volume 27, N. 1: 33-44. <https://www150.statcan.gc.ca/n1/pub/12-001-x/2001001/article/5852-eng.pdf>.

Singh, A.C., and C.A. Mohl. 1996. "Understanding Calibration Estimators in Survey Sampling". *Survey Methodology*, Volume 22, N. 2: 107-115. <https://www150.statcan.gc.ca/n1/en/pub/12-001-x/1996002/article/2973-eng.pdf?st=B9mYnII6>.

Wolter, K.M. 2007. *Introduction to Variance Estimation*. New York, NY, U.S.: Springer, Series in Statistics. <https://doi.org/10.1007/978-0-387-35099-8>.

Wright, R.L. 1983. "Finite Population Sampling With Multivariate Auxiliary Information". *Journal of the American Statistical Association*. Volume 78, N. 384: 879-884. <https://doi.org/10.2307/2288199>.

Zardetto, D. 2015. "ReGenesees: an Advanced R System for Calibration, Estimation and Sampling Error Assessment in Complex Sample Surveys". *Journal of Official Statistics*, Volume 31, N. 2: 177-203. <https://doi.org/10.1515/jos-2015-0013>.

Non-response in the Italian Building Permits Survey: a comparative analysis of imputation methods

Francesca Tuzi¹

Abstract

Non-response bias is a relevant issue in survey data. When information is missing, inferences drawn exclusively from observed data may fail to accurately reflect the entire population structure, leading to potentially invalid conclusions. Depending on the nature of the underlying selection mechanism, traditional estimators may be biased. As a result, alternative methods robust to selection bias arising from incomplete data should be considered. This paper discusses various approaches to handling missing data in the Italian Building Permits Survey. The results of this study are part of a broader project, launched in 2020, aimed at harmonising imputation methods for both the structural and the short-term versions of the survey, which was implemented in June 2021.

Keywords: Unit non-response, cross-sectional imputation, longitudinal imputation

DOI: https://doi.org/10.1481/ISTATRIVISTASTATISTICAUFFICIALE_1-2-3.2025.02.

¹ Francesca Tuzi (fratuzi@istat.it), Italian National Institute of Statistics - Istat.

The views and opinions expressed are those of the author and do not necessarily reflect the official policy or position of the Italian National Institute of Statistics - Istat.

The author would like to thank the anonymous reviewers for their comments and suggestions, which enhanced the quality of this article.

1. Introduction

One of the main drawbacks of survey data is missing data. Depending on the nature of the selection mechanism underlying the missing data, statistical units with complete data may fail to accurately reflect the population structure, leading to violations of the properties of traditional estimators. The way missing data are handled depends on several factors, in particular, the nature of the underlying mechanism for missing data. If the mechanism is ignorable with respect to the parameters being estimated, traditional methods will yield consistent estimates. However, if this is not the case, adjustments must be made to account for the incomplete data.

In longitudinal surveys, non-response is particularly problematic because the same units are repeatedly observed over time. As a result, the non-response rate may increase with each subsequent wave of data collection. This study analyses non-response in the context of the Italian Building Permits Survey. This census survey is longitudinal and collects monthly data on new residential and non-residential buildings, as well as any enlargements of pre-existing buildings, based on building permits issued by municipalities. The survey aims to produce both annual structural building statistics and a smaller set of quarterly short-term indicators for transmission to Eurostat.

Historically, non-response in this survey has been addressed using different imputation methodologies, depending on the frequency of the statistics and the characteristics of the municipalities, particularly their population size. For structural statistics, a mixed approach combining longitudinal and donor techniques has been applied to Self-Representative (SR) municipalities, i.e., those with 50,000 or more inhabitants. For short-term statistics, a mixed approach was used for SR municipalities, while regression-based methods were employed for Non-Self-Representative (NSR) municipalities, i.e., those with fewer than 50,000 inhabitants. The decision to apply a mixed approach was based on differences in response propensities between SR and NSR municipalities, as well as on differences in response rates for structural and short-term statistics.

In recent years, following the introduction of stricter penalties for non-responding municipalities and a new data collection system by the Italian

national Institute of Statistics (Istat), response rates have improved significantly, suggesting the potential for a new approach to addressing missing data.

This paper presents experimental results from a comparative analysis of the short-term version of the Istat Building Permits Survey, assessing the performance of various imputation methods. These results form part of a broader project to introduce a new methodological framework for imputing missing data in the survey. The aim is to identify improved procedures from both statistical and practical perspectives and to develop a harmonised methodological approach across the survey's structural and short-term versions. The new methodological framework took effect in June 2021, with the dissemination of the short-term statistics for the first quarter of 2021 and the structural statistics for the year 2020.

The paper is organised into five Sections. The first Section introduces the analysis. The second Section discusses theoretical aspects of non-response treatment. The third Section outlines the main features of the survey, including the data, the imputation methods under comparison, the simulation framework, and the statistical measures used to assess the results. The fourth Section presents the main results of the simulation, both at the aggregate and disaggregated levels, together with a sensitivity analysis examining the impact of shocks on response rates. The final Section provides concluding remarks and further considerations.

2. Theoretical aspects of non-response treatment

The literature on missing data traditionally distinguishes between two types of non-response: item non-response and unit non-response, based on whether the information available for each unit is partially or entirely missing. In the context of the Istat Building Permits Surveys, item non-response is largely negligible, as respondents are typically contacted for follow-up when a questionnaire is partially completed. Conversely, unit non-response, which refers to total non-response within the reference period, is of greater importance, particularly when it exhibits systematic non-response (attrition).

The choice of the most appropriate imputation method depends on several factors, in particular the nature of the underlying missing-data mechanism. As this mechanism is usually unknown, valid inferences about the assumed model parameters for the data require certain assumptions.

Most imputation procedures are based on the Missing at Random (MAR) assumption, which states that, conditional on the observed data, the probability of non-response is independent of the missing values (Little and Rubin 2002). When the non-response mechanism is independent of both the missing and the observed values, the missingness is said to be Missing Completely at Random (MCAR).

The literature on missing data presents a variety of approaches to data imputation.

The present study evaluates the predictive performance of several methods from both the cross-sectional (Chen and Shao 2000; Manzari 2004) and longitudinal imputation (Eurostat 2014) classes.

2.1 Cross-sectional and Longitudinal Imputation Methods

Among methods for cross-sectional data, a straightforward approach to missing data is to exclude units with incomplete information from the analysis (complete-case analysis). However, there are two potential problems with complete-case analysis. First, if units with missing values systematically differ from fully observed units, they do not constitute a random sample of the entire data set (i.e., under an MCAR non-response mechanism); this strategy may lead to biased estimates. Second, even under the MCAR assumption,

analysing only complete-case respondents may yield less accurate estimates. In addition, depending on the extent of missing data, there may be too few complete cases to make valid inferences.

An alternative approach to handling missing data is to use all available information from both respondent and non-respondent units by constructing homogeneous imputation classes. These classes are subsets of units, both respondent and non-respondent, identified using fully observable variables. Within each class, the imputed values are homogeneous and consistent with those of the responding units. This strategy has been shown to help reduce both the bias introduced by non-response and the variability associated with imputation (Kalton and Kasprzyk 1986; Kovar and Whitridge 1995).

Within the class of cross-sectional data methods, another technique for predicting missing values is conditional imputation based on auxiliary variables. This approach introduces a notion of similarity between units, obtained by estimating an appropriate distance function that depends only on fully observable auxiliary variables (nearest neighbours). For each imputed unit, the donor is selected from units that minimise the distance function (Chen and Shao 2000).

Another approach to handling missing data uses regression techniques, either directly or indirectly (Bethlehem 1988). In this method, the values of the missing variable are predicted using a parametric model, represented by a regression equation that depends only on fully observable variables (regression imputation). The predicted values can then be used directly to replace the missing data, or indirectly through homogeneous adjustment cells of units (Brick and Montaquila 2009; Kalton and Flores-Cervantes 2003; Little 1986).

A second class of imputation methods applies to longitudinal data, in which respondents provide responses over time in a systematic way. For each unit i , with $i = 1, \dots, n$, data vectors are collected over t periods, with $t = 1, \dots, T$. For period t , a vector of cross-sectional data (observed units) is available, whereas for unit i , a vector of longitudinal data, which is likely to be highly correlated, is collected (Greene 2002; Wooldridge 2001). In this setting, non-response becomes a more important issue than in the traditional context, as it affects both the cross-sectional and longitudinal dimensions of the data. From a theoretical perspective, dealing with

non-response in panel data is more complex than in the classical case, as it involves managing the two dimensions of data: cross-sectional and longitudinal (Rubin 1987).

In panel data, longitudinal information offers at least two advantages. First, longitudinal data for a given statistical unit are typically reliable predictors of missing values, leading to more accurate imputations. Second, having both backward and forward information for imputation is particularly useful for variables that change over time, such as short-term statistics (Eurostat, 2014). However, even in a longitudinal context, cross-sectional imputation methods may be necessary in certain cases, such as for units responding for the first time in the panel or for completely non-compliant units. In such cases, alternative strategies may be employed, such as constructing prediction models for missing values based exclusively on information from responding units. The choice of imputation method depends on the characteristics of the observed variable. For instance, in a longitudinal context, the seasonal pattern of a variable may reveal important features that could influence the imputation process.

3. Application of a case study: the Istat Building Permits Survey

The Istat Building Permits Survey is a census-based survey that collects monthly data on new construction projects or the enlargement of pre-existing structures, both residential and non-residential, authorised by certified building permits (including building permits, the Start-of-Work Certified Report (SCIA), or Public Construction Work). Changes of use and renovations to pre-existing buildings that do not increase the buildings' volume are not included in the survey's scope. The observed population comprises municipalities applying for single-building works, either a completely new building (even if demolished and rebuilt) or an enlargement of a pre-existing building. Two or more building works carried out under the same building permit constitute two or more units of analysis, for which similar forms are completed.

The survey is part of the Italian National Statistical Programme and produces two types of statistics. Structural indicators are produced annually (Structural Statistics), and since 2003, quarterly economic indicators (Short-Term Statistics) have been transmitted to Eurostat within 90 days of the end of the reference quarter, in accordance with EU Regulation². Until June 2021, the treatment of non-response in the short-term version of the survey followed different methodologies depending on: 1) the size of the municipality in terms of population (SR and NSR municipalities); 2) the nature of the non-response within a three-month period (total or partial non-response). Total non-response in SR municipalities was treated using Minimum-Distance Donor (MDD) techniques, while partial non-response was addressed using longitudinal information from the same municipality. Regression-based imputation techniques based on propensity scores were applied to both total and partial non-response in NSR municipalities (Bacchini et al. 2006). In addition, because of lower non-response rates in SR municipalities, census-type imputation techniques were used, whereas sampling techniques were applied in NSR municipalities (Garozzo and Rallo 2008).

Rising response rates in recent years, particularly among NSR municipalities, have prompted a reconsideration of the imputation framework (Leo and Tuzi 2023). With high response rates, sampling-

2 Current edition: NSP 2020-2022. Updating for 2021-2022 approved by the Decree of the President of the Republic of 15 December 2022, published in the Ordinary Supplement N. 7 of the Official Journal - General Series - N. 44 of 21 February 2023.

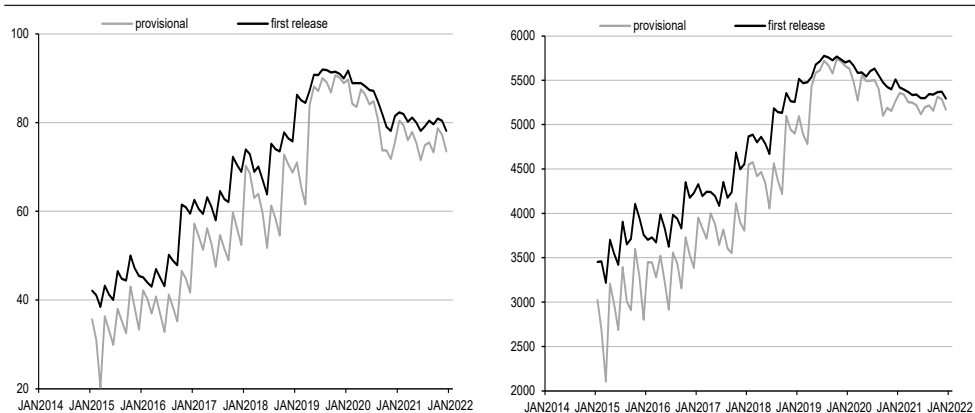
based imputation methods are no longer mandatory and may even be inconsistent with the structural version of the statistics. Moreover, in the context of panel data, longitudinal information could improve the quality of the imputed data.

3.1 The Data

The participation of responding units has increased significantly over the years, both in the proportion of the population and in the absolute number of municipalities. The most significant increases occurred in 2018 and 2019, particularly from the second quarter of 2018 onwards, when the provisional data line moved closer to the first revision, reflecting higher response rates in the first data release (Figure 3.1).

The annual average of weighted response rates at the provisional release (Table 3.1) rose from 56% in 2016 to over 90% in 2019, with the 2019 weighted rate reaching 93%. While the weighted and absolute rates were similar until 2018, the discrepancy became more pronounced in 2019, likely due to higher response rates in less populous municipalities.

Figure 3.1 - Fraction of respondents relative to total units (left side) and related population (right side). Provisional and first revision releases. Years 2014-2022 (a)



Source: Istat, Building Permits Survey

(a) On the right side, the population is divided by 10,000.

Due to the historically low response propensity of NSR municipalities, the short-term survey has generally restricted the use of longitudinal data to

imputing missing values in SR municipalities. However, with the recent improvement in response rates, one purpose of the experimental analysis was to assess whether extending longitudinal imputation procedures to include less-populated municipalities could improve the quality of the estimates.

Table 3.1 - Response rates by different data releases. Years 2016-2019 (absolute and weighted annual averages)

YEARS	Absolute rates		Weighted rates	
	Provisional	1 st Revision	Provisional	1 st Revision
2016	39.60	48.75	56.27	64.70
2017	52.64	63.46	62.67	71.36
2018	63.58	72.43	75.18	83.05
2019	82.72	89.31	90.25	93.52

Source: Istat, Building Permits Survey

Table 3.2 presents the weighted and absolute response rates for SR and NSR municipalities in 2019, by the number of months of collaboration and the geographical area to which the municipality belongs. The weighted response rate for municipalities that participated for at least 11 months is 78.6% in the provisional release, increasing by approximately 10 percentage points in the first revision. The proportion of non-respondents, defined as those who did not provide data, was 2.4% in the provisional data, compared with 2.3% in the first revision, reflecting a persistent non-response component.

Regarding the nature of the units, 87.2% of SR municipalities with at least 11 months of collaboration are included in the provisional data, compared with 73.7% of NSR municipalities. Overall, the proportion of municipalities with the highest number of months of participation is substantial. Among respondents in the 1-10 months category, the average number of months of collaboration is relatively high, with NSR municipalities averaging 7.8 months and SR municipalities averaging 8.4 months.

Response rates for SR municipalities in the 11-12 months category are very high across all Italian regions, even at the provisional release. However, significant differences are observed among NSR municipalities. The weighted response rate in the Northern regions is approximately 80% at the provisional release and exceeds 90% at the first revision.

In contrast, the response rate for the Central regions and Islands reaches 70% at the provisional release and rises to 80.0% at the first revision. The

Southern regions show lower response rates, at 61.5% at the provisional release and 73.0% at the first revision.

Table 3.2 - Weighted response rates by number of collaboration months, type of municipality, and Italian geographical area. Year 2019 (percentage values)

GEOGRAPHICAL AREAS	0 months		1-10 months		11-12 months	
	Provisional	1 st Revision	Provisional	1 st Revision	Provisional	1 st Revision
NSR Municipalities						
North-West	0.9	0.9	16.8	8.2	82.3	90.9
North-East	0.9	0.9	20.0	7.0	79.0	92.1
Centre	3.7	3.5	22.8	11.7	73.5	84.8
South	7.8	7.6	30.6	19.4	61.5	73.0
Islands	4.4	4.1	26.1	15.3	69.5	80.6
Total	3.4	3.2	23.0	11.8	73.7	84.9
Mean (Median)			7.8 (9.0)	7.5 (8.0)	11.8 (12.0)	11.9 (12.0)
SR Municipalities						
North-West	0.0	0.0	8.1	0.0	91.9	100.0
North-East	0.0	0.0	4.2	2.7	95.8	97.3
Centre	0.0	0.0	14.7	10.4	85.3	89.6
South	1.2	1.2	21.3	13.1	77.5	85.7
Islands	4.1	4.1	9.2	4.8	86.6	91.1
Total	0.8	0.8	12.0	6.9	87.2	92.4
Mean (Median)			8.4 (9.0)	8.5 (9.0)	11.9 (12.0)	11.9 (12.0)
Total	2.4	2.3	19.0	10.0	78.6	87.6
Mean (Median)			7.8 (9.0)	7.5 (8.0)	11.8 (12.0)	11.9 (12.0)

Source: Istat, Building Permits Survey

The evidence from the analysis suggests that extending longitudinal imputation procedures to include less populous municipalities is a reasonable approach. Furthermore, the geographical dimension should not be neglected when planning new procedures.

3.2 Alternative methodologies for missing data imputation

Due to higher non-response rates and greater variability in the estimates, the analysis was carried out on the subpopulation of NSR municipalities. Among the cross-sectional methods, three imputation procedures were tested, based on: (a) a similarity criterion between statistical units, (b) homogeneous imputation groups, and (c) the prediction of missing data using regression models. For purely theoretical purposes, a method based solely on the available data was also included.

Within the class of methods based on the criterion of similarity between units, two versions of the Minimum-Distance Donor (MDD) technique were tested: the simple version (Quis) and the normalised version (Quis_{norm}). The latter accounts for the municipality's population size. Missing data are replaced with the value of the minimum-distance donor, estimated using a set of auxiliary variables identified within homogeneous strata defined by geographical and social characteristics. These auxiliary variables, known for all units in the NSR subpopulation, include geographical area and the municipality's population. The strata for the simulation were obtained by combining five population classes (0-3,000; 3,001-7,000; 7,001-13,000; 13,001-25,000; 25,001-50,000) with five Italian geographical areas (North-West, North-East, Centre, South, Islands).

Let $(x_1, \dots, x_q)$ denote the vector of q auxiliary variables, and $D(i, k)$ the distance function between the missing unit i and the potential donor k . The minimum-distance donor k for the non-respondent unit i is selected by minimising the following distance function with respect to the auxiliary variable x :

$$D(i, k) = \sum_{j=1}^q |x_{ji} - x_{jk}| \quad (1)$$

A second useful approach within the class of cross-sectional methods approximates missing values using statistical measures, such as the mean (MSTR, MSTR_{norm}) and the median (MED_{STR}), estimated within homogeneous strata of responding units. The missing value is replaced by the average or median of the observed values within the stratum to which the non-responding unit belongs. The strata are those described in the previous section.

Let n_h^r denote the number of respondents in the h stratum to which the non-respondent municipality i belongs. The imputed value for the non-respondent unit i is then given by:

$$y_h^i = \frac{1}{n_h^r} \sum_{j=1}^{n_h^r} y_j \quad (2)$$

In addition to the simple average, a normalised version is estimated using the same weighting structure as described for the MDD technique. Furthermore, to account for potential outliers, median-based estimators were also tested.

The last class of cross-sectional methods comprises adjustment-cell models, which are based on the definition of groups of sample units (cells) with approximately equal response probabilities for a specific item (Eltinge and Yansaneh 1977; Bethlehem 2002). Adjustment cells are formed by grouping sample units according to their estimated response probabilities, derived either directly or indirectly via regression equations. The resulting adjustment estimator for a population average or total will have non-response bias close to zero, provided that the within-cell covariances between survey items and response probabilities are approximately zero (Falorsi et al. 2005). In addition to the Fabi estimator (used until 2021), which imputes missing data by identifying homogeneous non-responding cells, a classical regression imputation technique (Regr) was also tested, in which the estimated probabilities are used directly in the imputation procedure. The Fabi estimator, employed for imputing non-response in NSR municipalities, is an adjustment-cell estimator:

$$\tilde{Y}_t = \sum_{k \in r_t} y_{t,k} w_{t,k} = y_{t,k} \frac{1}{\pi_{t,k|at} \varphi_{k|s_t} + (1 - \pi_{t,k|at}) \varphi_{k|\bar{s}_t}} \quad (3)$$

where $y_{t,k}$ is the variable of interest for the k -th unit at time t , $w_{t,k}$ is the corresponding weight for the k -th unit, derived from the estimated parameters of the logistic regression; $\pi_{t,k|at}$ and $(1 - \pi_{t,k|at})$ represent the inclusion and non-inclusion probabilities, respectively, while $\varphi_{k|s_t}$ and $\varphi_{k|\bar{s}_t}$ denote the conditional probabilities of response for the two subsamples. The last two terms are approximated by the proportions of responding units in each subsample relative to the corresponding totals.

The use of repeated time-series data for the same unit is particularly suitable for imputing missing data in a longitudinal context. This study proposes a set of procedures that combine longitudinal information in different ways. The first version employs short-term information, namely the three months preceding the missing data, aggregated using a simple average (LT). The second version uses longer-term information, namely the twelve months

preceding the missing data, which are aggregated in several ways: simple average (LA), simple average excluding the extreme values of the distribution ($LA_{\min\max}$), and the median (LA_{median}). Let $y_{i,t}$ be the observed value of the y variable for the generic i -th unit at time t , and let $r_{i,t}$ be the dichotomous variable indicating response and non-response for the i -th unit at the same time t , such that:

$$r_{i,t} = \begin{cases} 0 & \text{se } y_{i,t} = 0 \\ 1 & \text{se } y_{i,t} \geq 1 \end{cases} \quad (4)$$

If $y_{i,t}$ is missing but longitudinal information is available, its value will be replaced with the average observed in the previous period. Let the reference period be defined by the extremes $t-m$ and $t-l$. The imputed value will then be derived from the following expression:

$$\tilde{y}_{it} = \frac{\sum_{k=t-m}^{t-1} r_{i,k} y_{i,k}}{\sum_{k=t-m}^{t-1} r_{i,k}} \quad (5)$$

The imputed value is approximated by the average of the observed $y_{i,k}$ values from $t-m$ to $t-l$, weighted by the dichotomous variable $r_{i,k}$. As noted in the literature, a drawback of these methods is that they are not applicable to units that are either completely non-respondents or have limited longitudinal information. In such cases, the procedure is not applicable, and the MDD imputation method must be used.

For purely theoretical purposes, an estimator based on complete information (AvCs) is proposed. A total of twelve estimators were tested: two based on donor imputation (Quis and Quis_{norm}); three based on simple and normalised averages (MSTR, MSTR_{norm}) and the stratum-wise median (MED_{STR}); two based on regression imputation (Fabi, Regr); one based on available cases only (AvCs); and two based on longitudinal imputation methods (LA, LT). Within the longitudinal methods, three versions of data aggregation were empirically tested: first, the simple mean (LA_{mean}); second, the median (LA_{median}); and finally, the simple mean pre-treated to exclude outliers ($LA_{\min\max}$).

3.3 The simulation design

The aim of the analysis was to evaluate the performance of different imputation methods in estimating key variables from the short-term version of the survey, specifically the number of residential dwellings and the non-residential surface area. The simulation was carried out using monthly data from the first quarter of 2018 to the last quarter of 2019. The quality of each imputation method was assessed using the provisional version of the data, which was characterised by lower response rates. SR municipalities were excluded from the simulations, both because of their high response rates and because longitudinal methods were already in use for them. The simulation was carried out on the fully responding municipalities within each of the 24 months. For each month, 100 random samples of units were generated, consistent with the geographical distribution of response and non-response in the reference month. Missing units were selected from each sample according to the observed proportions in the true distribution, while preserving the original geographical structure. The samples were drawn within strata defined by the five Italian geographical areas: North-West (NW), North-East (NE), Centre (C), South (S) and Islands (I). It should be noted that the number of sample units may vary from month to month, depending on the number of respondents and the territorial distribution of responses and non-responses. However, it remains consistent across samples within the same month.

In 2019, for instance, 5,469 of 7,750 NSR municipalities responded in January, 5,050 in February, and 6,881 in December (Table 3.3). The 100 samples for January were drawn from the respondent subpopulations (5,469 for January, 5,050 for February, and 6,881 for December). Each sample was simulated to reflect the territorial distribution of the entire population (respondents and non-respondents) for that month. The proportion of non-respondent units sampled within each geographical area matched the true distribution. For instance, in January 2019, 25.7% of residents in the NE were classified as non-respondents, while the remaining 74.3% were respondents. A similar distribution was observed in the NW, where 16.8% of residents were non-respondents, and the remaining 83.2% were respondents.

The different imputation methods used require further explanation. Regarding the longitudinal methods, samples were drawn from the subpopulation of respondents likely to have a higher propensity to respond

in earlier periods. To minimise potential bias in the estimates due to self-selection in the samples, the true backward information has been replaced with a simulated version based on the reference population's distributions. In the case of the LT method, it covers the three months prior to the missing data.

Table 3.3 - Distribution of simulated municipalities by Italian geographical area and month. Year 2019 (absolute and percentage values) (a)

MONTHS	North-West		North-East		Centre		South		Islands		Total
	n	Sampling fraction	n	Sampling fraction	n	Sampling fraction	n	Sampling fraction	n	Sampling fraction	
January	1,147	16.8	2,203	25.7	610	34.8	999	42.4	510	30.9	5,469
February	1,094	20.6	2,027	31.6	559	40.3	908	47.6	462	37.4	5,050
March	1,051	23.7	1,929	34.9	511	45.4	821	52.7	413	44.0	4,725
April	1,238	10.2	2,696	9.0	745	20.4	1,234	28.8	571	22.6	6,484
May	1,290	6.4	2,797	5.6	796	15.0	1,324	23.6	611	17.2	6,818
June	1,266	8.1	2,768	6.6	792	15.4	1,299	25.1	610	17.3	6,735
July	1,300	5.7	2,824	4.7	821	12.3	1,393	19.7	628	14.9	6,966
August	1,280	7.1	2,802	5.5	806	13.9	1,380	20.4	620	16.0	6,888
September	1,250	9.3	2,759	6.9	789	15.7	1,319	23.9	593	19.6	6,710
October	1,300	5.7	2,832	4.5	822	12.2	1,428	17.6	633	14.2	7,015
November	1,297	5.9	2,825	4.7	814	13.0	1,416	18.3	622	15.7	6,974
December	1,279	7.2	2,790	5.9	797	14.9	1,391	19.8	624	15.4	6,881
TOTAL	1,378		2,964		936		1,734		738		7,750

Source: Istat, Building Permits Survey

(a) n = Sample size of always respondents.

Regarding the MDD procedure, the possibility of encountering empty donor cells was a primary concern. This arose from using multiple strata, resulting from integrating a five-dimensional population variable with a five-dimensional geographic variable. However, given the high response rates observed in recent years, this is considered unlikely.

The Fabi estimator was simulated by partitioning the municipalities into sampled and non-sampled units in line with the original population distribution. Stratification was carried out by constructing 25 strata based on combinations of demographic and geographic variables. For the longitudinal information used to estimate the logistic models, the response/non-response vectors at times $t-1$ and $t-2$ were simulated using the actual distributions from the previous two years. This approach reduces the risk of self-selection that could arise when sampling only from the initial subpopulation of respondents. The way the samples were selected ensures that, because the process is random, the two backward vectors for

the same municipality may differ across replications. However, the overall distribution structure remains consistent with respect to the response/non-response pattern.

3.4 Statistical measures of accuracy

The performance of each imputation method was assessed by comparing the distributions of the simulated and true data. Several statistical measures were used: bias and variability of estimates, and the sensitivity of the methods to exogenous shocks. With the exception of the regression-based methods, where the discrepancy between the two distributions can be assessed only at the aggregate level (the sum of the variables), all other methods were compared at the highest level of disaggregation of the data.

Simple and absolute distance measures were used to estimate the deviations between simulated and true values for the donor technique, methods based on homogeneous grouping of units, and methods using longitudinal information. Simple means and Standard Deviations (SDs) were then calculated for these measures. These statistics are intended to provide insight into the bias of the estimates, its direction, and the variability introduced by the different imputation procedures. The similarity between the distributions of the original data, $Y = (Y_1, Y_2, \dots, Y_n)$, and those of the imputed data, $\tilde{Y} = (\tilde{Y}_1, \tilde{Y}_2, \dots, \tilde{Y}_n)$, was assessed using various distance measures. The first measure was the simple distance, which approximates the difference between the original and the imputed data by estimating the simple average:

$$d_1(\tilde{Y}, Y) = \frac{1}{n} \sum_{i=1}^n (\tilde{Y}_i - Y_i) \quad (6)$$

A second measure approximates the difference between the imputed and true values in absolute terms:

$$d_2(\tilde{Y}, Y) = \frac{1}{n} \sum_{i=1}^n |\tilde{Y}_i - Y_i| \quad (7)$$

To assess the degree of similarity between the simulated and true distributions, the Kolmogorov-Smirnov (K-S) distance was also calculated. This statistical measure quantifies the degree of similarity between the simulated and theoretical distributions by the maximum distance between

their respective cumulative distribution functions³. Denoting $F_{\tilde{Y}_n}$ and F_{Y_n} as the cumulative distribution functions of the simulated and true data, respectively, the K-S distance is given by the following formula:

$$d_{KS}(F_{\tilde{Y}_n}, F_{Y_n}) = \max_t (|F_{\tilde{Y}_n}(t) - F_{Y_n}(t)|) \quad (8)$$

In addition to the methods described above, the simulated and true totals were compared across all imputation procedures using the Mean Absolute Percentage Error (MAPE). This is a widely used indicator in the field of Short-Term Statistics. The MAPE measures the relative discrepancy between the true and simulated totals, expressed as a percentage. Let $\tilde{Y}$ and Y represent the distributions of the simulated and true data, respectively, and let $n = (n_1, n_2)$ be the data vector, where n_1 is the subset of units with complete information and n_2 is the subset of units with missing data.

The approximate measure of the distance between the observed and simulated totals is:

$$d_3(\tilde{Y}, Y) = 100 \cdot \frac{|\sum_{i=1}^n Y_i - \sum_{i=1}^n \tilde{Y}_i|}{\sum_{i=1}^n \tilde{Y}_i} \quad (9)$$

Finally, to evaluate the robustness of the imputation methods described in Section 3.2, a sensitivity analysis was conducted with respect to exogenous shocks to response rates in a single month. The analysis assessed the effects of these shocks on various statistical indicators at the monthly, quarterly, and annual levels.

3 The Kolmogorov-Smirnov test does not rely on any assumptions regarding the underlying distribution. It has the advantage that the distribution of the test statistic is independent of the cumulative distribution function being tested.

4. Results

This Section focuses on estimating key variables in the residential and non-residential sectors, using 2,400 simulated samples. The findings from the cross-sectional and longitudinal methods were assessed at both the disaggregated and aggregated levels.

By contrast, estimates derived from weighting techniques were evaluated only at the aggregated level.

4.1 Micro-level analysis

Statistical measures were computed for each year, month, and sample using a pair of 12x100-element matrices, which were then synthesised through successive aggregations. Deviations between the imputed and true values were estimated at the highest level of data disaggregation (the municipality). Averages over the 12 months were then computed for each sample, and an overall average was derived. For each group defined by month, sample, and field of interest (residential and non-residential), both simple and absolute distances between imputed and true values were estimated. The statistical measures described in the previous section were applied to both the full sample of respondent and non-respondent units and to the subsample of imputed units only. The statistical indicators based on absolute distance are intended to highlight inconsistencies between the distributions, while the simple version of the distance reveals the presence of bias and its direction.

All imputation techniques produced better results in 2019 than in 2018, primarily due to higher response rates in 2019. In line with Istat's recent policy to increase municipal participation in surveys, future trends are more likely to resemble those observed in 2019 than in 2018 (Tables 4.1 and 4.2). Focusing on 2019, the results show very small differences between methods, particularly for the total sample. Across imputation methods, the mean and standard deviation of the absolute distance measures exhibit a trade-off: methods that perform better on the mean-deviation metric tend to show greater variability, and vice versa. This relationship is particularly evident in the residential context.

Table 4.1 - Absolute deviations between imputed and true values by imputation method.
Years 2018 and 2019 (mean and standard deviation (SD) of average values) (a) (b)

IMPUTATION METHODS	2018				2019			
	Full sample		Imputed values		Full sample		Imputed values	
	Mean	SD	Mean	SD	Mean	SD	Mean	SD
RESIDENTIAL CONTEXT								
Quis	0.304	2.154	0.851	3.568	0.135	1.264	0.784	3.118
Quis _{norm}	0.304	2.158	0.852	3.575	0.135	1.265	0.784	3.121
MSTR	0.267	1.553	0.749	2.549	0.119	0.898	0.686	2.169
MSTR _{norm}	0.266	1.550	0.744	2.545	0.118	0.895	0.682	2.163
MED _{STR}	0.176	1.612	0.493	2.693	0.079	0.945	0.454	2.338
AvCs	0.177	1.625	-	-	0.079	0.952	-	-
LA	0.235	1.555	0.663	2.571	0.108	0.894	0.619	2.166
LA _{mean}	0.178	1.574	0.498	2.627	0.079	0.906	0.454	2.236
LA _{minmax}	0.272	1.652	0.768	2.724	0.126	0.987	0.737	2.388
LT	0.258	1.752	0.723	2.905	0.115	1.005	0.663	2.437
NON-RESIDENTIAL CONTEXT								
Quis	0.808	10.658	2.259	17.610	0.360	5.986	1.956	14.236
Quis _{norm}	0.809	10.658	2.260	17.611	0.360	5.987	1.957	14.237
MSTR	0.726	6.766	2.024	11.111	0.336	4.087	1.821	9.577
MSTR _{norm}	0.723	6.779	2.019	11.134	0.335	4.089	1.816	9.582
MED _{STR}	0.415	6.651	1.160	11.053	0.194	4.111	1.050	9.749
AvCs	0.415	6.651	-	-	0.194	4.111	-	-
LA	0.667	7.137	1.868	11.812	0.312	4.551	1.735	10.751
LA _{mean}	0.420	6.631	1.175	11.018	0.197	4.096	1.065	9.709
LA _{minmax}	0.681	7.371	1.912	12.202	0.310	4.374	1.735	10.289
LT	0.725	8.617	2.031	14.398	0.327	5.365	1.828	12.625

Source: Istat, Building Permits Survey

(a) Deviations estimated at the municipality level.

(b) Non-residential values in hundreds.

Table 4.2 - Simple distance between imputed and true values by imputation method.
Years 2018 and 2019 (mean and standard deviation (SD) of average values) (a) (b)

IMPUTATION METHODS	2018				2019			
	Full sample		Imputed values		Full sample		Imputed values	
	Mean	SD	Mean	SD	Mean	SD	Mean	SD
RESIDENTIAL CONTEXT								
Quis	-0.003	2.176	-0.008	3.670	-0.001	1.272	-0.009	3.215
Quis _{norm}	-0.003	2.180	-0.007	3.677	-0.001	1.273	-0.009	3.217
MSTR	-0.001	1.577	-0.003	2.662	-0.001	0.907	-0.003	2.278
MSTR _{norm}	-0.001	1.574	-0.002	2.657	-0.001	0.904	-0.002	2.270
MED _{STR}	-0.170	1.613	-0.476	2.696	-0.076	0.945	-0.440	2.341
AvCs	-0.177	1.625	-	-	-0.079	0.952	-	-
LA	-0.025	1.573	-0.067	2.657	-0.006	0.902	-0.033	2.253
LA _{mean}	-0.141	1.578	-0.394	2.646	-0.062	0.908	-0.355	2.254
LA _{minmax}	0.030	1.675	0.093	2.831	0.020	0.996	0.138	2.495
LT	-0.008	1.771	-0.023	2.996	-0.004	1.012	-0.008	2.525

Source: Istat, Building Permits Survey

(a) Distance estimated at the municipality level.

(b) Non-residential values in hundreds.

Table 4.2 cont. - Simple distance between imputed and true values by imputation method. Years 2018 and 2019 (mean and standard deviation (SD) of average values) (a) (b)

IMPUTATION METHODS	2018				2019			
	Full sample		Imputed values		Full sample		Imputed values	
	Mean	SD	Mean	SD	Mean	SD	Mean	SD
NON-RESIDENTIAL CONTEXT								
Quis	0.026	10.693	0.071	17.777	-0.006	5.999	-0.030	14.378
Quis _{norm}	0.027	10.693	0.072	17.779	-0.006	5.999	-0.029	14.379
MSTR	0.010	6.812	0.022	11.333	-0.002	4.105	-0.009	9.779
MSTR _{norm}	0.011	6.826	0.025	11.354	-0.002	4.107	-0.007	9.783
MED _{STR}	-0.414	6.651	-1.159	11.053	-0.194	4.111	-1.050	9.749
AvCs	-0.415	6.651	-	-	-0.194	4.111	-	-
LA	-0.035	7.173	-0.095	11.982	-0.020	4.563	-0.048	10.902
LA _{mean}	-0.389	6.633	-1.088	11.030	-0.182	4.097	-0.984	9.719
LA _{minmax}	-0.021	7.407	-0.052	12.374	-0.025	4.387	-0.055	10.449
LT	0.006	8.652	0.018	14.559	-0.015	5.375	-0.002	12.766

Source: Istat, Building Permits Survey

(a) Distance estimated at the municipality level.

(b) Non-residential values in hundreds.

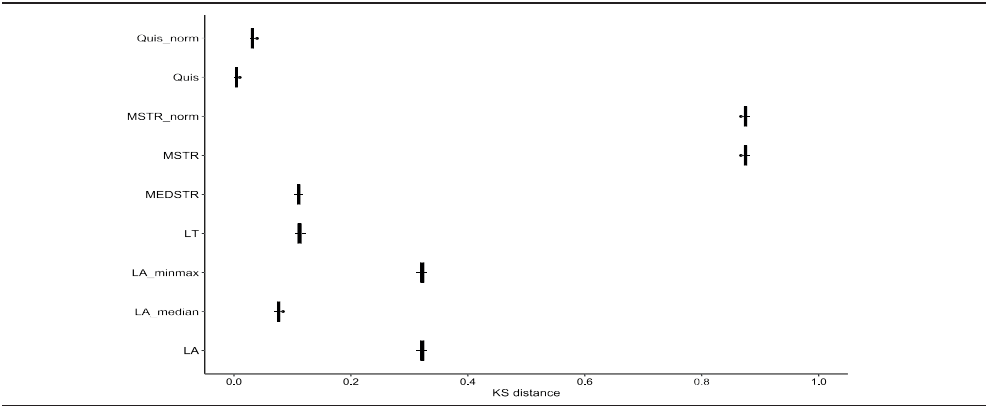
Across both the residential and non-residential sectors, the best results, in terms of average absolute distances, are those based on the median approximation, whether in a cross-sectional form (MED_{STR}) or in a longitudinal context (LA_{median}). In addition, the most convenient standard deviations for both fields are those obtained using simple and normalised means by stratum (MSTR and MSTR_{norm}). Very similar results were obtained for the residential domain with the LA estimator and for the non-residential domain with the LA_{median} estimator. Overall, in terms of absolute distance measures, the traditional annual longitudinal technique (LA) achieves the highest performance in the residential case, whereas the annual longitudinal technique with aggregation via the median (LA_{median}) yields the best results in the non-residential case.

Simple deviation means are generally underestimated by all methods except the annual longitudinal technique (LA_{minmax}), which excludes extreme values. This technique slightly overestimates in the residential context. Smaller biases are observed for donor techniques (Quis), for those based on stratum means (MSTR), and for longitudinal techniques based on quarterly information (LT) or on monthly information aggregated by means (LA). Using the median to impute missing values, whether by cross-sectional (MEDST) or longitudinal (LA_{median}) methods, results in greater bias in the estimates. Furthermore, MDD techniques show increased instability, as evidenced by the standard deviations

of the simple distance measures. The results suggest that methods based on stratum means are superior in both contexts, and that the traditional annual longitudinal (LA) technique is also superior in the residential domain.

The Kolmogorov-Smirnov distance measures for both the residential and non-residential sectors perform optimally when the MDD method is used. In contrast, the stratified means (MSTR) preserve the characteristics of the observed data (Figure 4.1) to a lesser extent than alternative techniques. This result is closely related to the zero-inflated nature of the observed data distribution, which is better approximated by the donor technique. As with previous indicators, longitudinal imputation yields intermediate results.

Figure 4.1 - Box Plot of the Kolmogorov-Smirnov (K-S) distance. Year 2019



Source: Authors' processing

4.2 Macro-level Analysis

Mean Absolute Percentage Error (MAPE) indicators were estimated for each sample, month and imputation method. Aggregate measures for each imputation technique were obtained by first averaging the sample-level MAPEs for each month, then calculating the overall average at the global level. The twelve indicators were then compared across time frequencies: annual, quarterly and monthly (Table 4.3)⁴.

⁴ For simplicity, the results of the normalised versions have been omitted, as they showed very similar results to the non-standardised versions.

Across all imputation techniques, the annual MAPE for 2019 was lower than for 2018, indicating improved response rates over time. Median-based imputation techniques (MED_{STR} and LA_{median}) produced less accurate results in both residential and non-residential domains. By contrast, the most efficient technique, indicated by the lowest MAPE, was imputation by stratum mean (MSTR).

In the residential domain, imputation methods based on regression techniques (Fabi, Regr), those based on average values stratified by stratum (MSTR), and those using annual longitudinal data (LA) yield small, very similar MAPE values. By contrast, the donor technique (Quis) and the quarterly longitudinal approximation (LT) are slightly less effective. All other methods are considered less suitable for this context. In the non-residential field, the imputation method based on average values differentiated by stratum (MSTR) performs best, followed by regression-based estimators and the annual longitudinal method (LA) or its variant that excludes extreme values (LA_{minmax}). Conversely, the donor techniques (Quis) and those based on quarterly longitudinal data (LT) perform comparatively less well.

Table 4.3 - Mean Absolute Percentage Errors (MAPEs) by imputation method and domain. Years 2018 and 2019

RESIDENTIAL CONTEXT									NON-RESIDENTIAL CONTEXT									Non-resp. Rate
Quis	MSTR	MED _{STR}	LA	LA _{median}	LA _{minmax}	LT	Fabi	Quis	MSTR	MED _{STR}	LA	LA _{median}	LA _{minmax}	LT	Fabi			
YEAR																		
2018	6.3	4.9	32.9	6.1	27.1	9.1	5.8	4.7	11.2	8.3	30.9	8.1	28.9	9.5	9.4	8.5	24.8	
2019	3.1	2.3	15.6	2.6	12.5	5.1	3.0	2.4	5.3	4.0	14.2	5.1	13.3	5.1	5.6	4.5	9.7	
QUARTER																		
2018-I	6.9	5.9	28.8	5.4	24.3	4.8	5.0	5.5	11.0	7.3	26.8	8.8	25.4	9.6	7.9	7.4	25.3	
2018-II	6.8	5.0	38.2	10.9	33.0	5.3	7.6	5.0	9.6	7.6	37.6	6.6	35.4	7.3	10.3	7.6	29.0	
2018-III	6.8	5.3	38.0	4.9	30.2	13.0	7.6	5.0	18.4	14.3	34.8	11.3	32.4	12.2	12.6	14.5	27.4	
2018-IV	4.4	3.4	26.7	3.3	21.0	13.4	3.2	3.3	5.7	4.2	24.4	5.7	22.7	8.7	6.8	4.3	17.6	
2019-I	5.1	4.0	30.3	4.6	24.2	6.9	5.2	4.1	9.3	7.2	29.0	9.9	27.0	9.8	11.4	8.6	18.6	
2019-II	2.4	1.8	12.5	2.1	10.1	3.7	1.9	1.8	4.5	3.3	10.7	3.2	10.1	3.2	3.9	3.6	8.3	
2019-III	2.4	1.7	10.4	2.2	8.3	5.5	2.9	1.8	3.7	2.7	9.3	3.8	8.6	3.7	4.3	2.9	6.5	
2019-IV	2.5	1.7	9.3	1.7	7.3	4.1	1.9	1.8	3.8	2.7	7.8	3.4	7.3	3.7	3.0	2.9	5.6	

Source: Istat, Building Permits Survey

Table 4.3 cont. - Mean Absolute Percentage Errors (MAPEs) by imputation method and domain. Years 2018 and 2019

	RESIDENTIAL CONTEXT								NON-RESIDENTIAL CONTEXT								Non-resp. Rate
	Quis	MSTR	MED _{STR}	LA	LA _{median}	LA _{minmax}	LT	Fabi	Quis	MSTR	MED _{STR}	LA	LA _{median}	LA _{minmax}	LT	Fabi	
2019																	
Jan.	3.8	3.1	25.8	2.4	19.9	9.5	3.1	3.2	12.3	11.3	24.1	9.7	23.0	9.8	11.0	11.4	15.7
Feb.	5.7	4.0	30.4	2.9	23.5	8.0	3.4	4.4	6.4	4.6	30.3	10.1	27.3	7.6	12.4	5.0	19.1
Mar.	5.9	4.9	34.7	8.4	29.2	3.1	9.1	4.8	9.3	5.8	32.7	9.9	30.7	11.9	10.8	9.5	20.9
Apr.	2.7	1.7	14.6	1.6	11.6	6.4	2.2	1.8	6.8	4.4	13.2	4.7	12.6	4.8	4.6	4.3	10.1
May	1.9	1.7	10.9	2.6	9.3	1.7	1.9	1.6	3.3	2.9	8.2	2.6	7.7	2.5	3.6	3.9	7.6
June	2.5	2.0	11.8	1.9	9.4	3.1	1.6	2.0	3.4	2.5	10.8	2.4	10.1	2.4	3.4	2.6	7.2
July	2.4	1.8	9.4	2.4	8.4	2.0	1.7	1.8	3.6	2.7	8.5	2.5	8.1	2.6	2.6	2.6	5.4
Aug.	2.4	1.7	9.7	3.0	7.0	8.4	5.3	1.7	3.1	2.5	8.9	5.7	8.1	5.3	6.9	2.7	6.2
Sep.	2.3	1.7	12.0	1.4	9.6	6.0	1.7	1.7	4.5	2.9	10.4	3.2	9.8	3.3	3.2	3.4	7.7
Oct.	1.9	1.6	8.1	1.9	7.0	2.1	2.0	1.7	4.1	2.7	6.9	2.6	6.5	2.7	2.5	3.2	5.1
Nov.	3.0	1.7	9.5	1.4	7.3	4.5	1.6	1.7	2.5	2.2	7.9	4.4	7.3	5.1	3.2	2.3	5.5
Dec.	2.5	1.9	10.2	1.7	7.7	5.8	2.0	1.9	4.8	3.2	8.7	3.2	8.2	3.5	3.1	3.3	6.4

Source: Istat, Building Permits Survey

The greater accuracy of the MSTR imputation method in both the residential and non-residential contexts is confirmed at the quarterly frequency. This is followed by the longitudinal method approximated by the annual average (LA). A similar analysis at the monthly frequency yields comparable results in the non-residential context. However, in the residential context, as evidenced by the greater prevalence of months with the lowest MAPE throughout the year, the annual longitudinal (LA) method generally yields better results than the other methods. In the residential domain, the lack of convergence between results at the annual and quarterly frequencies and those at the monthly frequency is likely influenced by the March 2019 period, which was characterised by an exceptionally high non-response rate.

The MAPE results suggest that the imputation method based on stratified means (MSTR) is preferable when the aim is greater stability at the aggregate level, even in the presence of high non-response rates. Conversely, the longitudinal method (LA_{minmax}), despite its lower precision under low response rates, generally yields better results in terms of the number of months for which the MAPE is minimised.

4.3 Sensitivity analysis

To assess the robustness of the different imputation techniques, a sensitivity analysis was conducted with respect to changes in the response rates for a specific month, reducing them to 70%. The impact of this perturbation was assessed at monthly, quarterly, and annual frequencies, based on estimates of the indicators presented in Section 3.4⁵.

Table 4.4 presents the estimated means and standard deviations of the absolute values of the distance measures and the MAPEs across different data frequencies.

Table 4.4 - Mean and standard deviations of absolute distances, and MAPE by imputation method. Perturbation in response rates. Year 2019

IMPUTATION METHODS	YEAR			QUARTER			MONTH		
	Mean	SD	MAPE	Mean	SD	MAPE	Mean	SD	MAPE
RESIDENTIAL CONTEXT									
Quis	0.172	1.397	3.7	0.227	1.642	4.8	0.523	2.708	10.1
MSTR	0.152	1.009	2.9	0.202	1.197	3.9	0.465	2.031	8.4
MED _{STR}	0.101	1.067	20.6	0.131	1.252	29.4	0.302	2.123	69.9
LA	0.139	1.002	2.8	0.182	1.159	2.4	0.424	1.970	3.5
LAm _{edian}	0.101	1.018	16.6	0.128	1.192	23.8	0.294	2.028	56.7
LAm _{inmax}	0.166	1.112	8.1	0.232	1.313	16.2	0.544	2.259	40.9
LT	0.148	1.126	3.1	0.192	1.291	2.5	0.448	2.198	3.6
Fabi	-	-	2.9	-	-	4.2	-	-	8.8
NON-RESIDENTIAL CONTEXT									
Quis	0.440	6.443	6.4	0.482	5.603	8.2	1.076	7.510	15.7
MSTR	0.408	4.406	4.8	0.435	3.545	5.8	0.976	5.376	11.7
MED _{STR}	0.236	4.432	19.5	0.250	3.560	29.1	0.567	5.408	71.8
LA	0.401	4.932	7.9	0.502	4.096	14.8	1.190	6.512	38.7
LAm _{edian}	0.239	4.414	16.6	0.252	3.535	23.8	0.569	5.354	56.7
LAm _{inmax}	0.400	4.742	8.3	0.517	3.965	16.6	1.218	6.259	43.7
LT	0.412	5.908	7.3	0.486	5.091	9.8	1.146	8.811	23.9
Fabi	-	-	5.3	-	-	6.0	-	-	11.5

Source: Istat, Building Permits Survey

For both the residential and non-residential sectors, the most significant effects of the disturbance are observed at the monthly and quarterly levels in the month and quarter in which the shock occurs, whereas the effects at the annual level are minimal. The methods with the lowest absolute averages are those based on the median (LA_{median} or MED_{STR}), as well as the annual

⁵ For the analysis, only the most relevant results are presented in this document, although all indicators have been calculated.

longitudinal technique (LA) for the residential domain, which also exhibits the lowest variability. In the residential domain, the least effective techniques, in terms of both the mean and the standard deviation of absolute distances, are those based on the donor method (Quis) or on longitudinal data without outliers ($LA_{\min\max}$). A similar pattern of results can be observed at the annual level in the non-residential area.

In the non-residential domain, there is less convergence towards a single method, as shown by the results at the quarterly and monthly levels. The largest deviations are observed for the annual longitudinal method (LA) and its outlier-free variant ($LA_{\min\max}$), whereas the donor technique (Quis) and the method based on short-term longitudinal data (LT) show greater variability.

For the macro-level indicators, the results are consistent across data frequencies and both domains. In the residential domain, the lowest MAPEs are achieved by longitudinal methods using annual or quarterly averages (LA, LT). In the non-residential domain, the most accurate results are obtained using methods based on stratum-level averages (MSTR) and regression estimators (Fabi). Methods based on the median value (MED_{STR} , LA_{median}) and on longitudinal data without extreme values ($LA_{\min\max}$) are not recommended. However, results vary by data frequency: the MAPE of the $LA_{\min\max}$ estimator at the annual level is significantly lower than that of its quarterly and monthly versions.

The results of the sensitivity analysis to response-rate shocks provide useful insights into both the statistical superiority of certain techniques over others and the possibility of excluding some low-performing methods. In general, the first category includes the LA technique (for dwellings) and the MSTR method (for both dwellings and non-dwellings). Although the latter does not produce the best results in absolute terms, it offers an intermediate solution across the indicators used. Methods based on MDD techniques belong to the second group. Moreover, the choice of the most appropriate imputation method is closely related to data frequency, especially in the non-residential context. It is therefore crucial to take into account the time domain underlying the statistics when making this decision.

5. Concluding remarks

The increase in response rates observed in recent years in the Building Permits Survey has prompted a re-examination of the methodological framework for imputing missing data. The use of different imputation methods has become less effective as response rates have increased, even in smaller municipalities. In addition, the existing literature suggests that, in panel surveys, longitudinal information has strong predictive power for imputing missing data.

This study evaluated several imputation methods, both cross-sectional and longitudinal, for the NSR subset of municipalities. The results showed that, compared with the procedures used before the innovations introduced in June 2021, better results can be achieved not only through alternative cross-sectional methods but also by using longitudinal information for the same unit.

In both application fields, the outcomes of the micro, macro, and sensitivity analyses do not converge to an optimal imputation technique in absolute terms. Significant differences occur at the domain level, with lower convergence in the non-residential field due to its greater variability.

Regarding estimate quality, there are significant differences between the two application areas. In the residential context, there is near-complete convergence towards the annual longitudinal technique and the stratified mean. The Kolmogorov-Smirnov distance measure is the only technical limitation in this respect. However, strict distributional similarity is not a crucial factor in selecting the most efficient procedure, since building permit statistics are produced in aggregate form. In the non-residential domain, greater dispersion in the results is observed. The mean by stratum is identified as the most appropriate technique, followed by longitudinal methods based on the median or those excluding extreme values from the distribution. In contrast, donor techniques generally yield suboptimal results.

The experimental analysis confirms the predictive power of longitudinal information for imputing missing data in the Building Permits Survey. The findings also suggest that different approaches should be applied depending on the field of interest, residential and non-residential. In addition, the sensitivity analysis provides useful insights into the limitations of certain methods relative to others and highlights the importance of the time frequency of the statistics in choosing the most appropriate method.

References

Bacchini, F., R. Iannaccone, and E. Otranto. 2006. "Imputation of missing values for longitudinal data: an application to the Italian Building Permits". *Rivista di statistica ufficiale/Review of official statistics*, N. 1/2006: 29-42. Roma, Italy: Istat. https://www.istat.it/wp-content/uploads/2011/06/1_2006.pdf.

Bethlehem, J.G. 2002. "Weighting Nonresponse Adjustments Based on Auxiliary Information". In Groves, R.M., D.A. Dillman, J.L. Eltinge, and R.J.A. Little (Eds.). *Survey Nonresponse*: 275-288. New York, NY, U.S.: Wiley & Sons.

Bethlehem, J.G. 1988. "Reduction of Nonresponse Bias through Regression Estimation". *Journal of Official Statistics*, Volume 4, N. 3: 251-260. <https://www.scb.se/contentassets/ca21efb41fee47d293bbee5bf7be7fb3/reduction-of-nonresponse-bias-through-regression-estimation.pdf>.

Brick, J.M., and J.M. Montaquila. 2009. "Nonresponse and Weighting". In Pfeffermann, D., and C.R. Rao (Eds.). *Sample Surveys: Design, Methods and Applications*. Handbook of Statistics 29A: 163-185. Amsterdam, The Netherlands: North-Holland Publishing Company. [https://doi.org/10.1016/S0169-7161\(08\)00008-4](https://doi.org/10.1016/S0169-7161(08)00008-4).

Chen, J., and J. Shao. 2000. "Nearest Neighbor Imputation for Survey Data". *Journal of Official Statistics*, Volume 16, N. 2: 113-131. <https://www.scb.se/contentassets/ca21efb41fee47d293bbee5bf7be7fb3/nearest-neighbor-imputation-for-survey-data.pdf>.

Eltinge, J.L., and I.S. Yansaneh. 1997. "Diagnostics for Formation of Nonresponse Adjustment Cells, With an Application to Income Nonresponse in the U.S. Consumer Expenditure Survey". *Survey Methodology*, Volume 23, N. 1: 33-40. <https://www150.statcan.gc.ca/n1/en/pub/12-001-x/1997001/article/3103-eng.pdf?st=baTm247f>.

Eurostat, ESSnet Project Memobust - Methodology of Modern Business Statistics. 2014. "Imputation for Longitudinal Data". In Eurostat, ESSnet Project Memobust. *Memobust Handbook on Methodology of Modern Business Statistics*. Luxembourg: Eurostat.

Falorsi, P.D., G. Alleva, F. Bacchini, and R. Iannaccone. 2005. “Estimates based on preliminary data from a specific subsample and from respondents not included in the subsample”. *Statistical Methods and Applications*, Volume 14, N. 1: 83-99. <https://doi.org/10.1007/BF02511576>.

Garozzo, S., e G. Rallo. 2008. “La rilevazione dei permessi di costruire: il controllo e la correzione dei dati”. *Contributi Istat*, N. 13/2008: 169-184. Roma, Italia: Istat. https://www.istat.it/wp-content/uploads/2018/07/13_2008.pdf.

Greene, W.H. 2002. *Econometric Analysis*. Saddle River, NJ, U.S.: Prentice Hall.

Kalton, G., and I. Flores-Cervantes. 2003. “Weighting Methods”. *Journal of Official Statistics*, Volume 19, N. 2: 81-97. <https://www.scb.se/contentassets/ca21efb41fee47d293bbee5bf7be7fb3/weighting-methods.pdf>.

Kalton, G., and D. Kasprzyk. 1986. “The treatment of Missing Survey Data”. *Survey Methodology*, Volume 12, N. 1: 1-16. https://www150.statcan.gc.ca/n1/en/pub/12-001-x/1986001/article/14404-eng.pdf?st=KSPE_BGv.

Kovar, J.G., and P.J. Whitridge. 1995. “Imputation of Business Survey Data”. In Cox, B.G., D.A. Binder, B. Nanjamma Chinnappa, A. Christianson, M.J. Colledge, and P.S. Kott (Eds). *Business Survey Methods*: Chapter 22: 403-423. Hoboken, NJ, U.S.: John Wiley & Sons. <https://doi.org/10.1002/9781118150504.ch22>.

Leo, A., e F. Tuzi (a cura di). 2023. *Novità metodologiche e di diffusione degli indicatori dei permessi di costruire*. Letture Statistiche - Metodi. Roma, Italia: Istat. <https://www.istat.it/produzione-editoriale/novita-metodologiche-e-di-diffusione-degli-indicatori-dei-permessi-di-costruire/>.

Little, R.J.A. 1986. “Survey Nonresponse Adjustments for Estimates of Means”. *International Statistical Review*, Volume 54, N. 2: 139-157. <https://doi.org/10.2307/1403140>.

Little, R.J.A., and D.B. Rubin. 2002. *Statistical Analysis with Missing Data*. Hoboken, NJ, U.S.: Wiley & Sons. <https://doi.org/10.1002/9781119013563>.

Manzari, A. 2004. “Combining Editing and Imputation Methods: An Experimental Application on Population Census Data”. *Journal of the Royal Statistical Society, Series A (Statistics in Society)*, Volume 167, N. 2: 295-307. <http://www.jstor.org/stable/3559862>.

Rubin, D.B. 1987. *Multiple Imputation for Nonresponse in Surveys*. Hoboken, NJ, U.S.: Wiley & Sons. <https://doi.org/10.1002/9780470316696>.

Wooldridge, J.M. 2001. *Econometric Analysis of Cross Section and Panel Data*. Cambridge, MA, U.S.; London, UK: The MIT Press. <https://ipcid.org/evaluation/apoio/Wooldridge%20-%20Cross-section%20and%20Panel%20Data.pdf>.

Italian Tourism Satellite Account: methodological approach

Sandra Maresca¹, Ilaria Piscitelli¹

Abstract

Although official tourism statistics provide extensive information on the tourism sector, a Tourism Satellite Account (TSA) integrates these data into a methodological and conceptual framework derived from the National Accounts. The methodology developed for the Italian TSA uses official sources and adopts a mixed approach for the demand perspective and a bottom-up approach for the supply perspective. The ITSA Tables follow international standards for both characteristic tourism products and industries and for country-specific tourism goods, as selected from the list provided in the guidelines. A TSA measures the direct effects of tourism consumption within the reference economy; it does not capture the full potential of the tourism sector, but it is nevertheless the first step in measuring indirect and induced effects.

Keywords: Tourism Satellite Account (TSA), tourism, National Accounts, tourism expenditure.

DOI: https://doi.org/10.1481/ISTATRIVISTASTATISTICAUFFICIALE_1-2-3.2025.03.

¹ Sandra Maresca (maresca@istat.it); Ilaria Piscitelli (piscitelli@istat.it), Italian National Institute of Statistics - Istat.

The views and opinions expressed are those of the author and do not necessarily reflect the official policy or position of the Italian National Institute of Statistics - Istat.

The author would like to thank the anonymous reviewers for their comments and suggestions, which enhanced the quality of this article.

1. Introduction

The TSA is an international, standardised accounting framework, structured into ten Tables, suitable for measuring the tourism sector within the reference economy. For this reason, it is built on the principles of national accounting, representing an extension and a thematic focus (Eurostat 2009 a and b).

Tourism statistics emphasise the activities undertaken by visitors, with particular attention to descriptive variables such as mode of transport, main purpose of travel, main destination, and expenditure (Eurostat 2014). This sectoral description of tourism has long been internationally integrated through the TSA, which consistently highlights the tourism sector's direct role in the broader national economy.

A TSA (Eurostat 2002):

- makes the tourism sector visible, even though its contribution is already included in the central National Accounts;
- measures tourism on the demand side (tourism expenditure and consumption, collective tourism consumption, fixed gross capital formation), on the supply side (production, intermediate consumption, value added, employment), and provides non-monetary indicators (overnight stays, number of trips, number of establishments, etc.);
- sets up an accounting structure to compare monetary and non-monetary estimates, in order to draw the measure of the tourism sector consistently with National Accounts.

The Italian TSA compiles 8 of the ten recommended Tables: Tables 1-7 and Table 10, which contain the essential information for measuring the size of the tourism sector that directly serves visitors. This information is expressed in terms of the main aggregates of national accounting (added value, gross domestic product, employment, etc.) and is compiled mainly using its methodology and evaluation criteria, thus enabling comparison with other productive sectors and/or the rest of the economy.

The TSA measures the direct effects of tourism on the reference economy. However, broader measures of indirect and induced effects are expressly excluded from the TSA's purposes.

2. The Tourism Satellite Account

In National Accounts, the economy is described through a detailed framework that covers the supply of goods and services – both domestically produced and imported – and their demand for intermediate and final consumption; the value added generated by economic activities; and the interrelations among activities and products.

A TSA represents a focus and a deepening of such a framework, in which the analysis of all aspects of demand for goods and services that might be associated with tourism is highlighted using the structure and principles of national accounting, while also developing specific issues and concepts suitable for making the tourism sector as a whole emerge from the framework of National Accounts in which it is included.

2.1 Key concepts of National Accounts and of TSA

2.1.1 *Residence*

One of the foundational concepts of National Accounts (NA) is residence. Residence is attributed to individuals on the basis of the family (unit) to which they belong, which in turn is established by whether they have and use, within the economic territory of reference, accommodation considered by the family itself as its main residence. When an individual travels outside the economic territory of reference, the place where they have their usual residence is decisive in determining their residence.

Since tourism is linked to the mobility of people, residence and abode are key elements in establishing the origins of different kinds of tourism flows and in accurately recording their activities. In tourism statistics, international tourists are classified by their country of residence, while domestic tourists are classified by their usual residence.

In the context of the NA, residence refers to travellers who travel outside their country of residence. In the context of the TSA, residence refers to the usual environment, used to distinguish travellers from visitors and to classify travellers by origin.

2.1.2 Production and value added

In the NA schemes, production and value added highlight each productive sector's contribution to economic performance. The TSA evaluates the tourism sector in the same way, using the same concepts and definitions.

NA production includes, in addition to the goods and services produced by resident production units, certain imputed items, such as the services provided by owners' own use of the dwelling – the so-called imputed rents. From a TSA perspective, such imputed rents apply only to services associated with holiday accommodation on one's own account, excluding, by definition, the main residence and referring only to second homes. As in NA, imputed rents are accounted for twice in a TSA: on the supply side – production – and on the demand side – tourism consumption.

2.1.3 Evaluation criteria

The evaluation criteria for TSA aggregates, on both the demand and supply sides, are the same as those used in national accounts: basic production prices and purchase prices for tourism consumption. Subsequent reconciliation requires adding elements such as taxes, imports, and trade and transport margins to the basic production prices to convert them to purchase prices.

A special issue in the TSA is the net evaluation of tour operators' production. The organising services, typically representing the main activity of a tour operator and included in the final price of a package tour, must be treated separately from the other services. In NA, a package tour is evaluated on a gross basis, whereas in the TSA, including the Italian one, it is evaluated on a net basis, as it covers only the value of the organising service actually performed by the tour operator.

2.2 Tourism on the demand side

A central element in defining the tourism sector is the visitor, without whom tourism flows and the associated monetary transactions would not exist. For this reason, tourism is regarded as a demand-side activity.

Visitors are a subset of travellers, and tourism is more limited than travel. A traveller is someone who moves between different geographic locations for any purpose and for any duration. A visitor moves outside their usual environment for less than a year and for a main purpose other than employment by a resident entity in the place visited (UNSD and UNWTO 2010).

2.2.1 Visitor, usual environment, purpose of trip

The definition of a visitor is based on three fundamental aspects:

1. usual environment, defined as the geographical area (though not necessarily contiguous) within which individuals conduct their regular daily routines (UNSD and UNWTO 2010). This fundamental concept is characterised by factors such as distance from the usual residence, frequency and duration of movement, and crossing internal or international borders (UNSD and UNWTO 2010);
2. duration of trip, less than 12 months;
3. purpose of trip, articulated as follows.
 - *Personal reasons*: holiday, leisure and recreation, visiting friends and relatives, education and training, health and medical care, religion, shopping, transit, other (UNSD and UNWTO 2010).
 - *Business and professional reasons*, including the activities of the self-employed and employees, provided they do not constitute an implicit or explicit employer-employee relationship with a resident producer in the country or place visited. Cross-border workers, i.e., those who cross an internal or international border daily to go to their workplace, as well as seasonal workers, are excluded.

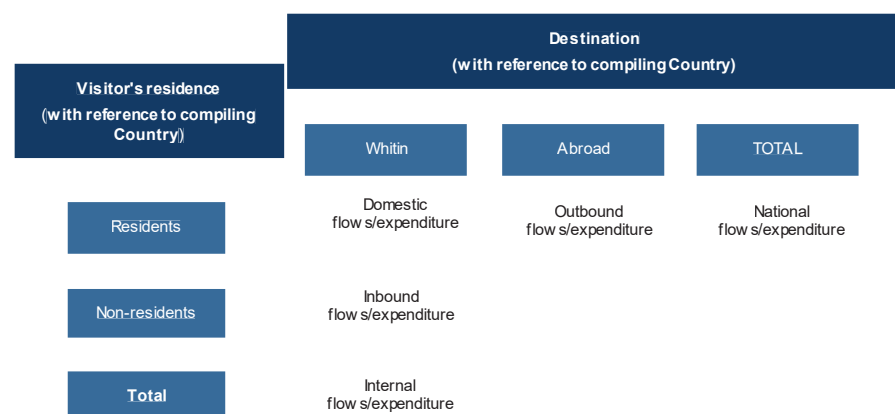
Visitors are classified as follows:

1. Tourist: overnight visitor.
2. Excursionist: same-day visitor.

2.2.2 Tourism flows, tourism expenditure and tourism consumption

From the perspective of the compiling country, the interaction between the concepts of residence and visitor gives rise to different types of tourism flows, which are essential for evaluating tourism expenditure. These expenditures, in turn, constitute crucial components of internal tourism consumption, a key concept for measuring tourism demand (Figure 2.1).

Figure 2.1 - Tourism flows and expenditures



Source: Authors' processing

Inbound flows and expenditures: referring to non-resident visitors travelling in the country of reference.

Domestic flows and expenditures: referring to resident visitors travelling in the country of reference.

Outbound flows and expenditures: referring to resident visitors travelling outside the country of reference.

National flows and expenditures: referring to resident visitors travelling in and out of the country of reference (domestic + outbound).

Internal flows and expenditures: referring to non-resident and resident visitors travelling within the country of reference (domestic + inbound).

The purchase of products by visitors, both residents and non-residents, is highly relevant to the economy of reference, as it forms part of the monetary aggregates used to measure the tourism sector on the demand side: tourism expenditure and internal tourism consumption.

Tourism expenditure is the amount spent on consumption goods and services, as well as valuables, for personal use or for giving away during tourism trips (United Nations Commission of the European Communities *et al.* 2010).

Tourism consumption extends beyond tourism expenditure. In addition to the monetary transactions that tourism expenditure focuses on, it also includes services associated with holiday accommodation on one's own account, tourism social transfers in kind, and other imputed consumption (United Nations Commission of the European Communities *et al.* 2010). These additional items are recorded in the specific column in Table 4 of the TSA, namely "Other components of tourism consumption".

Tourism expenditure and consumption refer to the broader definition of final consumption expenditure adopted in the NA, which defines it as "supported by resident institutional units for goods and services used for the direct satisfaction of individual or collective needs of members of a community" (Eurostat 2013). Actual final consumption, by contrast, refers to the "acquisition of consumer goods and services by a sector". The difference between the two concepts lies in how goods and services financed by governments or Non-Profit Institutions Serving Households (NPISHs) are treated when provided to households as in-kind social transfers (Eurostat 2013).

While expenditure considers only the monetary outlays directly incurred by families for the purchase of goods and services, actual consumption encompasses their overall consumption, including that for which families benefit without having directly incurred the expense. In the latter case, this concerns the benefits in kind provided by the government and NPISHs.

The comparison of expenditure and consumption in NA and TSA excludes employer-paid business-trip expenditure for employees (Table 2.1). From a TSA perspective, these expenses are tourism consumption for accommodation, transport services, and travel agencies' services; however, in NA, they are intermediate consumption for companies rather than final consumption by workers.

Table 2.1 - Actual final consumption in NA vs. tourism consumption in TSA

National Accounts (NA)	Tourism Satellite Account (TSA)
1. Households' final consumption	1. Visitors' expenditure
2. Individual final consumption of Government and NPISHs	2. Individual final consumption of Government and NPISHs
	3. Other components of tourism consumption
Actual households' final consumption	Tourism consumption

Source: Authors' processing

2.2.3 Other components of tourism consumption

The “Other components” of tourism consumption are mostly made up of:

- housing services provided by vacation homes on own account;
- business travel.

Each family has its own main residence, whose location identifies the country of residence and the place of usual residence. Any other homes (owned or leased by the family) are considered second homes.

For the purposes of the TSA, the only second homes to be considered relevant are those used for tourism, even if located near or within their usual environment, in derogation of the principles of frequency and duration established for identifying the tourism sector. The use of second homes for tourism is a distinctive element in quantifying the related services within the TSA. Another peculiarity is that second homes can be located in the economic territory of another country, in which case they affect the Balance of Payments (BoP) (IMF 2004).

Business travel refers to the expenses incurred by companies for their employees' business trips.

Expenses borne by companies – mainly accommodation, transport and travel agencies' services – are recorded in NA as intermediate consumption, whereas under a TSA perspective they form part of “Other components” of tourism consumption. Other expenses always borne by companies but covered through lump-sum payments are treated in NA as wage subsidies; therefore, such expenses, even in the context of business trips, are treated as final rather than intermediate. In their tourist transposition, these final expenditures feed tourism expenditure rather than tourism consumption.

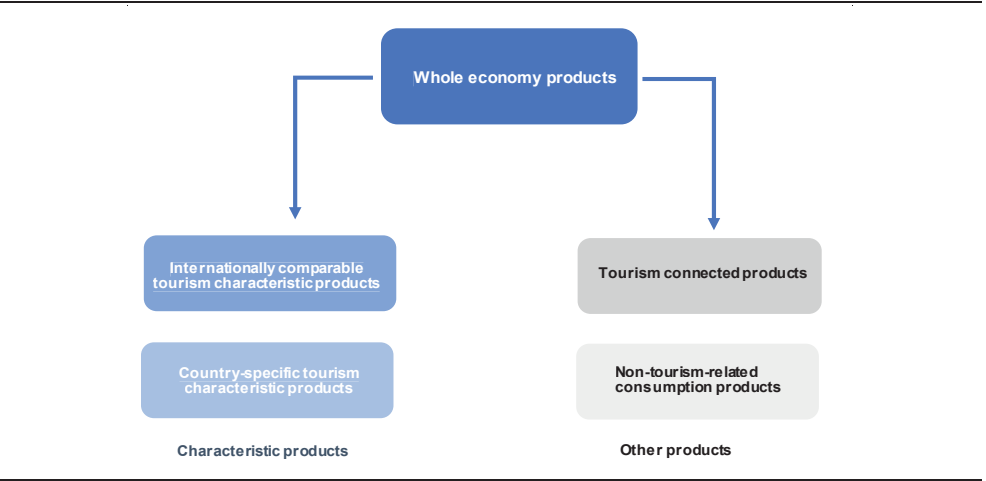
From the compiling country’s perspective, expenditures on business travel reported under “Other components” relate only to resident visitors, whereas those related to non-resident flows, which are traced back to the intermediate consumption of non-resident companies, do not support the economy of the TSA of reference.

2.3 Tourism on the supply side

On the supply side, one of the most significant aspects of the macroeconomic analysis conducted through the TSA is the description of tourism’s productive contribution, considered a sector in its own right, relative to the rest of the national economy. From this point of view, tourism, once again dominated by the crucial figure of the visitor, is defined as the set of productive activities (industries) that produce goods and services (products) for consumption in tourism.

Tourism products can be grouped into two macro-categories: (1) characteristic tourism products and (2) other products (Figure 2.2) (UNSD 2008).

Figure 2.2 - Classification of tourism products



Source: Authors’ processing

Characteristic products are those whose supply would cease or be significantly reduced in the absence of visitors; the others, whether connected

or not, represent residual categories whose link with tourism lacks the prevalence typical of the former.

The internationally comparable characteristic products are the top 10 listed in Table 2.2.

Table 2.2 - Tourism products and tourism industries in Italian TSA

Tourism products	Tourism industries
1. Accommodation services for visitors	1. Accommodation for visitors
2. Food- and beverage-serving services	2. Food- and beverage-serving activities
3. Railway passenger transport services	3. Railway passenger transport
4. Road passenger transport services	4. Road passenger transport
5. Water passenger transport services	5. Water passenger transport
6. Air passenger transport services	6. Air passenger transport
7. Transport equipment rental services	7. Transport equipment rental
8. Travel agencies and other reservation services	8. Travel agencies and other reservation
9. Cultural services	9. Cultural services
10. Sports and recreational services	10. Sports and recreational services
11. Shopping	11. Retail trade of country-specific tourism characteristic goods
12. Other products	12. Other industries

Source: Authors' processing

For specific characteristic products, tourism attribution varies significantly across countries or areas of reference. Each country can select and develop this basket of specific characteristic products from an extensive list of goods and services proposed in international manuals (UNSD and UNWTO 2010).

From a productive perspective, tourism-characteristic industries are those whose supply would cease to exist in significant quantities in the absence of visitors (UNSD and UNWTO 2010). This link is to be taken literally, in the sense that industries whose “main productive activities directly serve visitors or that are in contact with them” are considered tourist industries (UNSD and UNWTO 2010). Therefore, those that contribute to the production of products purchased by the visitor but which, however, reach him only through the channel of mediation are excluded from the meaning of the tourism industry.

The overall supply generated by the tourism industries and other (non-tourism) industries serves both the demand of visitors – residents and non-residents – and that of non-visitors (consumers) (Maresca et al. 2014).

In Table 5 of TSA (T5), which presents the production accounts for tourism and other industries, one of the main supply-side indicators is generated: the

value added of tourism industries (VATI), in addition to the tourism value added (VAT) estimated in Table 6 of TSA. The VATI does not provide any assessment of the actual tourist destination; it represents the accounting quantification of output produced by characteristic tourism industries, regardless of how that output is distributed for final use between visitors and consumers. Furthermore, because it is limited to the tourism industry, the VATI does not account for the productive contribution of other industries that also contribute to the direct satisfaction of tourism demand through their (secondary) supply of tourism goods and services.

3. The sources of Italian TSA

The implementation of the Italian TSA (ITSA) relies on official sources, thereby providing a pool of information that complies with international methodological standards (MAP et al. 2003).

They are divided into two main types:

- Surveys on tourism;
- National Accounts data.

The former provide monetary and non-monetary data for specific tourism aspects, while the latter are used both as basic data and at an aggregate level to assess the consistency of tourism estimates with those of the National Accounts.

3.1 Tourism statistics

Within this type of source, surveys are distinguished by their scope of analysis and the information they produce.

The main surveys used for the compilation of ITSA are:

- Survey on International Tourism (*Indagine sul turismo internazionale*) (by the Bank of Italy)² (Alivernini et al. 2014);
- Trips and Holidays Survey (THS) (*Indagine CAPI su viaggi e vacanze*) by Istat³ (2025c);
- Survey on Capacity of tourism accommodation establishments (*Capacità degli esercizi ricettivi*) (by Istat);
- Survey on Occupancy of tourist accommodation establishments (*Movimento dei clienti negli esercizi ricettivi*) by Istat (2025d).

3.1.1 Survey on International Tourism

The extensive sample Survey on International Tourism conducted by the Bank of Italy as part of its institutional activity of compiling Italy's BoP is a

2 For details on this survey, see <https://www.bancaditalia.it/statistiche/tematiche/rapporti-estero/turismo-internazionale/index.html?com.dotmarketing.htmlpage.language=1&dotcache=refresh>.

3 For details on this survey, see <https://www.istat.it/en/microdata/trips-and-holidays-2/>.

crucial source for compiling Tables 1 and 3 of the TSA, providing information on tourism expenditure in Italy by non-resident visitors and on tourism expenditure abroad by visitors who are resident in Italy.

The survey is conducted at a sample of road, rail, airport and port border points and captures both inbound and outbound tourism. The survey provides analytical information on tourism spending, the number of overnight stays, and other aspects of travel that are of interest for analysing tourism occurrences.

3.1.2 Trips and Holidays Survey

Since 2014, TH has been included as a focus in the broader Households Budget Survey, which has undergone structural and methodological changes. It measures: the number and characteristics of trips made by residents for holiday or business purposes, both in Italy and abroad; the expenditure incurred by families on travel; the reasons people have not travelled; and same-day visits (excursions). The estimates provided are internationally comparable and enable analysis of the evolution of individuals' tourist behaviour.

The survey mainly aims to quantify residents' tourism flows, both in Italy and abroad, and to provide a reference framework for the behaviour and travel methods of Italian tourism demand, including expenditure.

3.1.3 Survey on the Capacity of tourism accommodation establishments

The two Istat surveys on accommodation facilities provide the non-monetary data required for Table 10 of the Italian TSA. These data, in addition to describing some qualitative characteristics of tourism, provide an essential basis for calculating per capita values and for any case in which reliable indicators are needed to correct the estimate.

It is an annual census survey that detects:

- the number of establishments, beds, rooms and bathrooms at the municipal level;
- the number of establishments and beds at the municipal level.

The former is recorded for hotel establishments, which include all typological specifications and/or types of accommodation that, under regional

regulations, are classified as hotel-type structures (e.g., hotel villages, motels, period residences, *meubl  * or *garn  * hotels, historic residences, wellness centres, beauty farms, etc.). The latter is recorded for complementary establishments (divided into campsites, tourist villages, campsites and tourist villages in mixed form, rental accommodation managed in an entrepreneurial form, agrotourist and rural tourism accommodation, youth hostels, holiday homes, alpine refuges and other complementary accommodation facilities not elsewhere classifiable) and for private rental accommodation (divided into bed and breakfasts and other private accommodation). Information on private rental accommodation (“other private accommodation”) is also collected but not disseminated, as it is not considered exhaustive.

With regard to hotels only, since 2008, accommodation capacity has also been recorded by structure size class (fewer than 25 rooms, 25 to 99 rooms, or 100 rooms or more).

3.1.4 Survey on Occupancy of tourism accommodation establishments

The survey on customer movements in accommodation establishments is a monthly census that records the arrivals and overnight stays of resident and non-resident visitors in Italian accommodation facilities, classified by region and country of residence.

Arrivals are the number of customers who check in at accommodation establishments during the period considered; overnights are the number of nights they spend there.

Also, in this case, the survey units consist of accommodation establishments, that is, hotels, complementary structures and private rental accommodation.

The information collected mainly concerns arrivals and nights spent, by type of accommodation (using the same classification as the Survey on the Capacity), by country of residence for non-resident visitors, and by region of residence for visitors who are residents in Italy.

3.2 National Accounts data

Measuring the tourism sector within the economy of reference requires close integration with the National Accounts system. This connection is evident in the centrality of accounting data, which are used both as a direct source for compiling the TSA Tables and as a quantitative constraint to reconcile tourism estimates with the accounting aggregates.

As a data source, the National Accounts provides its information starting with the Supply and Use Tables (SUTs) (Istat 2025a).

On the supply side, the reconstruction of tourism industries and other industries of T5, in accordance with the “Tourism Satellite Account: Recommended Methodological Framework 2008” (TSA: RMF 2008) (United Nations Commission of the European Communities et al. 2010), required the use of basic data derived from company surveys at the 4-digit level of classification of economic activity (class). On the demand side, National Accounts data provided both the constraints for the estimates derived from the survey data and the information base for developing indicators.

3.2.1 *Supply and Use Tables (SUTs)*

The SUTs highlight the relationship between the branches of economic activity, i.e., the groupings of Units of Economic Activity (UEAs), and the branches of homogeneous production (products) by accurately describing the internal production processes and transactions associated with those products.

For the purposes of compiling the TSA, a non-disseminated internal version of the SUT matrices was used at the highest level of detail.

The Supply Table shows the total availability of resources at basic prices, distinguishing between domestic and imported production. Within this framework, the production matrix provides a detailed picture of the domestic supply of goods and services, classified by product and by industry. From this Table, tourism industries and products were selected.

For the same level of detail, the Use Table specifies the uses of goods and services by intermediate consumption, final consumption, gross fixed capital formation and exports. For each branch of economic activity, it describes the

structure of its costs, given by the combination of product inputs and the value added at purchase prices, in all its components.

Using this matrix, the costs of the tourism industries have been estimated. In addition, the duly processed data on imports, net taxes, and distribution margins were used to estimate, in Table 6 of TSA (T6), all the valuation components required therein. Finally, the Use Table was used to estimate the number of items that fall under the “Other components of tourism consumption” category (business travel).

3.2.2 *Survey data underlying the NA schemes*

The compilation of the ITSA cannot disregard the use of the basic data underlying the National Accounts. These include archives containing economic surveys of Italian companies, household final consumption, the housing census, and central government expenditure by the Classification Of the Functions Of Government (COFOG).

Frame-SBS and economic surveys of companies

Two surveys of companies are used for the Italian National Accounts compilation:

- the Survey on the System of Accounts of Enterprises (from now on, SCI);
- the Survey on Small and Medium-sized Enterprises and on the exercise of arts and professions (from now on, PMI).

The Frame-SBS archive is based on extensive use of administrative data and is integrated with data from the two aforementioned surveys, thereby enabling the exploitation of administrative data’s informational potential for statistical purposes.

The SCI census survey covers all Italian companies with at least 250 employees operating in the industrial and service sectors (this threshold was 100 employees until 2016), except for certain divisions within monetary and financial intermediation, insurance, and domestic services. Companies are drawn from the Istat Statistical Archive of Active Enterprises (ASIA) (Istat 2025b), which the survey collects, among other things, the main economic variables used to calculate value added.

The PMI sample survey covers all Italian enterprises with 1-249 employees in the industrial, wholesale, and services sectors that were active in the reference year. For each company, the survey records key economic variables, including employees, wages and salaries, gross fixed capital formation, revenues, and production costs. The list of sampled units is drawn from ASIA.

Data derived from the Frame-SBS archive, and used for the tourism industry's production account, were extracted at the level of economic activity classes.

Household Consumption Surveys (HCS)

The National Accounts data on household consumption underpin the TSA's estimates of domestic tourism expenditure and consumption and, in turn, are derived from the underlying sample survey.

The Household Consumption Survey conducted from 1997 to 2013 was replaced by the Household Expenditure Survey. It aims to detect the structure and the amount of final consumption expenditure according to the main social, economic and territorial characteristics of resident households.

In particular, expenditure on food, alcoholic beverages and tobacco, clothing and footwear, housing, water, electricity, fuels, furniture, household goods and services, health services and expenditure, transport, communications, recreation, entertainment and culture, education, accommodation and catering services, and other goods and services is recorded.

The survey adopts the European Classification of Individual Consumption by Purpose (COICOP), which groups consumption by the type of need it satisfies. However, for the purposes of the Supply and Use scheme, it is necessary to shift from function-based estimation to product-based estimation, using a special consumption transition matrix.

Final consumption of the government's individual final consumption and NPISHs

Expenditure on individual final consumption by government and NPISHs falls within the scope of the TSA to the extent that it has direct tourism value and contributes to the determination of the "Other components of tourist consumption" in Table 4 of the TSA (T4).

Government expenditure is classified by function according to the COFOG. The reference universe is the institutional sector of central government, as

defined in the European System of Accounts 2010 (ESA 2010) (Eurostat 2013), and includes all institutional units that produce other non-market goods and services intended for collective and individual consumption. Through the functional classification of general government expenditure, it is possible to distinguish between individual and collective final consumption expenditure. For the purposes of T4, only individual final consumption was considered⁴.

As far as NPISHs are concerned, their contribution to internal tourism consumption refers to expenditure on thermal activities.

ASIA, the Statistical Archive of Active Enterprises

ASIA⁵ is a Register of companies and local units, updated annually by Istat through the integration of administrative and statistical sources. The Register includes all economic units engaged in arts and professions in the industrial and service sectors and provides identifying (name and address) and structural information (economic activity, employees and self-employed, legal form, start and end dates of activity, turnover).

The field of observation for ASIA-Local Units is the same as for ASIA-Enterprises. The information provided concerns the location (at municipal level), the economic activity and the number of employees of local units dependent on active enterprises. The territorial detail of local units is necessary for compiling Table 10d, “Number of structures in the tourist industries classified by size class of employees”, which extends part of Table 7 on employment in the tourism industries.

Data on employment

The employment data processed for the National Accounts are fundamental to compiling Table 7 of the TSA (T7). Its standard format requires distinguishing between employed and self-employed, and by sex, for each of the three required labour inputs: jobs, hours worked, and FTEs⁶.

For the compilation of T7, data at the class level of economic activity were processed; however, they were not disaggregated by sex.

⁴ Collective consumption falls within the scope of Table 9 and is excluded from the ITSA.

⁵ Enterprises that have carried out production activity for at least six months in the reference year are considered active.

⁶ They represent a measure of employment in which part-time jobs (part-time employment contracts and second jobs) are reported in Full-Time Equivalent units. The work units are calculated net of the redundancy fund.

4. Estimation methodology for Italian TSA

4.1 Selection of tourism activities and products

The reconstruction of the tourism sector on the supply side begins by applying the general principle that enterprises are defined as tourism enterprises “whose main activity is the characteristic tourism activity itself and which directly serves visitors” (UNSD and UNWTO 2010). The criterion of the main activity, measured by prevailing value added, is the same as that which guides the classification of economic activities. Specifically, within the hierarchy of indicators used to attribute a firm’s main activity, added value is placed at the top (among others, value of production, sales, labour cost, number of workers, number of hours worked (Istat 2009). Nevertheless, the international manual also adds that “the output of tourism industries may not consist exclusively of products characteristic of tourism and the output of non-tourism industries may include some products characteristic of tourism” (UNSD and UNWTO 2010).

The approach to reconstructing the tourism industry has been bottom-up. Starting from the highest level of detail provided by the Italian SUT and the underlying basic data, it was possible to extract from the Italian National Accounts the branches that contain tourism, as required by the TSA.

It should be emphasised that the accounting branch, even when it includes tourism activities, rarely aligns with the TSA’s tourism industry, as the two are based on different criteria for grouping economic activities. Once the branches were selected from the National Accounts, their analysis was conducted in three distinct phases:

- breakdown of the accounting branch into classes of economic activity, i.e., at the 4-digit level of the classification of activities;
- elimination of all classes conforming to the “accounting” concept of the branch but not to the “tourism” concept of the industry;
- inclusion in each tourism industry of only the related tourism classes.

In Table 4.1, the composition of the characteristic tourism industries by class, as well as their corresponding links to the accounting branch of origin, is detailed.

Table 4.1 - The reconstruction of the tourism industries by class of economic activity

INA's selected branches of economic activities	ISIC Rev. 4 code – classes included in the selected branches	ISIC Rev.4 code – classes included in tourism industries	IT5 tourism industries
Accommodation	5510-5520-5530-5590	5510-5520-5530-55901	
Buying and selling of real estate and real estate activities for third parties	6810-6831-6832	6810-6831-6832	1 - Accommodation for visitors
Rental and management of properties owned or leased	6820	6820	
Food and beverage serving activities	5610-5621-5629-5630	5610-5630	2 - Food and beverage serving activities
Railway transport	4910-4920	4910	3 - Railway passenger transport
Other land passenger transport	4931-4932-4939	4932-4939	4 - Road passenger transport
Maritime and inland water transport	5010-5020-5030-5040	5010-5030	5 - Water passenger transport
Air transport	5110-5121-5122	5110	6 - Air passenger transport
Renting and operating leasing activities	7711-7712-7721-7722-7729-7731-7732-7733-7734-7735-7739-7740	7711	7 - Transport equipment rental
Services activities of Travel Agencies, Tour Operators and related reservation services and activities	7911-7912-7990	7911-7912-7990	8 - Travel agencies and other reservation services activities
Creative, arts and entertainment activities	9001-9002-9003-9004	9001-9002-9003-9004	9 - Cultural activities
Libraries, archives, museums and other cultural activities	9101-9102-9103-9104	9102-9103-9104	
Renting and operating leasing activities	7711-7712-7721-7722-7729-7731-7732-7733-7734-7735-7739-7740	7721	10 - Sports and recreational activities
Lotteries, betting and casino-related activities	9200	9200	
Sports, amusement and recreation activities	9311-9312-9313-9319-9321-9329	9311-9319-9321-9329	

Source: Authors' processing

Unlike in other industries, the reconstruction of tourism products was based on their identification within the well-articulated framework provided by the Italian SUTs⁷ (Table 4.2).

⁷ 98 branches and 262 products.

Table 4.2 - The reconstruction of the tourism products from Italian SUTs

Selected products relevant to tourism	IT5 products
Hotel services and similar	
Other accommodation services	
Sale of real estate services with own property made services	
Real estate services for third parties	1 - Accommodation services for visitors
Administration services and property management for third parties	
Real residential housing services	
Imputed residential housing services	
Food and beverage sales services	2 - Food and beverage serving services
Interurban railway passengers transport services	3 - Railway passenger transport services
Other land passenger transport services	4 - Road passenger transport services
Shipping, cabotage and inland water passenger transport services	5 - Water passenger transport services
Air transport services for passengers	6 - Air passenger transport services
Services of renting and leasing of consumer goods for recreation and leisure	7 - Transport equipment rental services
Travel agencies services	8 - Travel agencies and other reservation services activities
Tour Operators services	
Other reservation services and related services	
Creative, art and entertainment services	9 - Cultural services
Services relating to gambling	
Management services of sports arenas and sports facilities	10 - Sports and recreational services
Sports entertainment and recreation services	
Sports, amusement and recreation activities	

Source: Authors' processing

4.2 Estimation of output in tourism industries and other industries

The starting point for estimating tourism industry production in T5 is the SUTs' production matrix, available for internal use at the highest level of detail, together with the underlying basic data, also extracted at the highest available level of detail. The production estimate is based on the 4-digit class of the Italian Classification of Economic Activities (Ateco 2007), which is derived directly from the International Standard Industrial Classification of All Economic Activities (ISIC Rev. 4) at the 4-digit level of detail. However, for the products, due to a lack of more specific information, the highest level of articulation available in the National Accounts matrices was used. In Table 4.3, the production of the accounting branch of "Catering" is shown by product.

Table 4.3 - Production of catering branch by product and class of activity. NA vs. TSA.
Year 2019 (values in million euros) (a) (b)

Product description	Product Flag	ISIC Rev. 4				€(mln)	
		56.10	56.21	56.29	56.30	Total output branch (NA)	Total Tourism Industry (TSA)
		T	NT	T	NT		
Meat and meat products	NT	0	0	0	0	1	0
Milk and dairy products	NT	2	0	0	1	3	3
Condiments, baby food and other food products	NT	3	0	0	1	4	4
Electrical equipment for intermediate purposes	NT	0	0	0	0	1	0
Repair services	NT	10	0	2	4	16	14
Buildings and civil construction work	NT	1	0	0	0	2	2
Wholesale services for third parties	NT	14	0	2	6	23	20
Retail services (excluding automotive fuel)	TS	488	13	78	204	783	692
Storage and warehousing services and services related to land, sea, water and air transport	NT	0	0	0	0	1	0
Hotel and similar services	T	272	7	43	114	436	386
Other accommodation services	T	0	0	0	0	1	1
Catering and beverage sales services	T	50,458	1,355	8,011	21,061	80,904	71,538
Other catering services	NT	4,444	119	706	1,857	7,125	6,300
Software produced for own use	NT	3	0	1	1	5	5
Information processing, hosting and related services	NT	0	0	0	0	1	1
Actual residential rents	T	5	0	1	2	9	8
Non-residential rents	NT	21	1	3	9	34	30
Scientific research and development services	NT	7	0	1	3	11	10
Rental and leasing services for recreational and leisure consumer goods	T	16	0	2	7	25	22
Licensing services for the use of intellectual property products and similar products, excluding copyrighted works	NT	8	0	1	3	13	11
Building and landscape maintenance services	NT	0	0	0	0	1	1
Photocopying, conference and trade fair organisation services and other business support services n.e.c.	NT	19	1	3	8	30	27
Creative, artistic and entertainment services	T	1	0	0	0	1	1
Gambling services	T	343	9	54	143	550	487
TOTAL		56,117	1,507	8,909	23,445	89,981	79,562

Source: Istat, National Accounts, TSA for Italy

(a) The total may not match the sum of the products shown in the Table because, for simplicity, those with production below one million euros have been excluded.

(b) Data not disclosed.

Although the articulation in 98 branches of the SUTs' matrix allows the catering sector to be isolated, it is not possible at this level of aggregation to distinguish only the economic activities to be included in the corresponding tourism industry. To proceed with the disarticulation of the branch's production by class of activity, it is necessary to use the level of detail provided by the basic data. This Table also presents the data broken down by class of economic activity.

With reference to this branch, the tourism content for the ITSA refers to classes 56.10 and 56.30, while the remaining classes have been excluded from the tourism industry because they relate to the provision of meals for canteens, hospitals, etc.

As mentioned, the lack of basic data at the same level of detail as economic activities has prevented proceeding in a similar way. Nevertheless, the level of articulation of the Italian SUTs has enabled the precise identification of the products within the scope of the TSA. In particular, tourism characteristics subject to international comparison (flag T) have been identified among the 262 items of the accounting matrices.

Country-specific characteristic products (flag TS) have been identified among items representing distribution margins. In addition to tourism products, secondary production (flag NT) is necessary to reconstruct the economy's overall production, as T5 is merely a tourism perspective on production.

4.2.1 Special issues

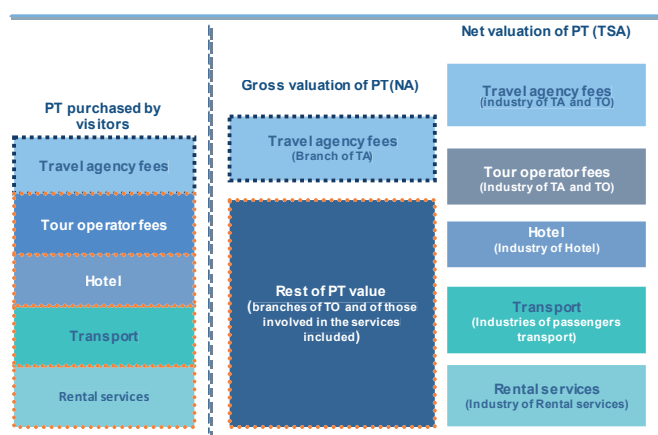
Net evaluation of Tour Operator services

From the TSA perspective, the net evaluation of Package Tours (PT) requires separately estimating the cost of the main services offered, in order to determine the share of production attributable solely to the organisational activity of Tour Operators (TO), which corresponds to the revenue associated with that activity, net of the commissions paid to Travel Agencies (TA) for any intermediary role in the sale of PT (Figure 4.1). This revenue, derived solely from the organisation's activity and with which the TO will still have to cover all other costs, is obtained by deducting the individual costs incurred in creating the PT from the PT's cost.

In National Accounts, services included in a PT are counted twice: once in the TO branch and again in the single branches of the service producers. This duplication, which also affects intermediate consumption, is offset when calculating added value.

In the case of the TSA, the estimate of TO production is net evaluated.

Figure 4.1 - Evaluation of a PT: NA vs TSA. Example of typical services



Source: Authors' processing

Second homes

“The ownership of a vacation home on one’s own account is unusual from a statistical perspective because it generates both a tourism characteristic service and an equivalent tourism consumption. In the System of National Accounts (SNA 1993 and 2008), a housing service on own account is associated with the ownership of a dwelling occupied by its owner and is treated as both a production activity and the output and consumption of a specific service. This situation covers both the principal dwelling and all other dwellings owned by a household for its own use. It covers, in particular, owner-occupied vacation homes” (United Nations Commission of the European Communities et al. 2010; European Commission et al. 2009).

From a TSA perspective, NA activity and product associated with such an issue are included globally within the corresponding tourism categories. Nevertheless, secondary dwellings are very important for tourism, and in some countries they

may represent a significant share of the overall accommodation service. In all these cases, compilers are encouraged “to develop a specific subcategory for owner-occupied secondary dwelling, both as a product and as an “industry” (United Nations Commission of the European Communities et al. 2010).

In Italy, the use of second homes for tourism is widespread among the resident population and can be considered a typical Italian way of meeting accommodation demand for domestic trips. T5 thus depicts this situation by reporting the imputed value of services for holiday purposes.

As already mentioned, in principle, T5 reflects the total output of tourism and non-tourism industries, regardless of the share directly demanded by tourists. The second-homes industry is an exception, as the industry’s total output is not gross production but rather a tourism component. In particular, the total production of the Italian NA branch has been weighted by the national share of second-home use for tourism purposes.

In this regard, the resort to two sources is needed, in particular:

- the final consumption expenses of households;
- the Census of Housing.

In accordance with European regulations, the national estimate of housing services is derived using a stratification method based on actual rents: information on the housing stock, divided into strata (Census), is combined with data on actual rents paid in each stratum (Households Budget Survey).

For the purposes of the TSA, the estimation of housing services for second homes used for tourism, a subset of the national total, follows the same method used in NA, with the addition of weighting the total value of the service by the rate of tourism use. In this regard, the use of a structure for the use of housing for tourism purposes, obtained by combining, on the one hand, the data of the Census of Housing and, on the other, the data of the Travel and Holidays survey, leads to the calculation of a specific statistical ratio necessary to extract the part of the TSA from the national data.

4.3 Estimation of tourism demand

On the demand side, the main measures of the tourism sector's contribution to the overall economy are tourism expenditure and tourism consumption.

In particular:

1. Inbound expenditure (T1);
2. Domestic expenditure (T2);
3. Outbound spending (T3);
4. Internal tourism consumption (T4).

4.3.1. *Inbound tourism expenditure*

The estimate of expenditure by non-resident visitors is derived from the Bank of Italy's Survey on International Tourism. It is the main source for estimating inbound expenditure; however, to compile T1 correctly in accordance with the required international standards and ensure its completion, it was essential to integrate it with other sources. The excessive aggregation of some expenditure items detected in this survey, or the lack thereof, necessitated recourse to Istat sources to complete the final estimate.

The expenditure on accommodation, as estimated by the Bank of Italy survey, has been adjusted to exclude the component for meals consumed in accommodation facilities, which, from the TSA's point of view, fall within the scope of the Catering product. To do this, the estimate of hotel turnover from the Quarterly Survey of the Turnover of Services was used, which provided data on hotel turnover by overnight and meal. On the other hand, the estimate derived from the Bank of Italy for services in inhabited dwellings used free of charge by relatives and friends has been eliminated and replaced with estimates of imputed rents derived from the NA.

The inbound expenditure for the transport service includes only expenditure related to the purchase of services rendered by resident carriers. In road transport, integrating the estimate derived from the Bank of Italy survey with the NA data enabled the separation of the share attributable to petrol for allocation to "Other product".

4.3.2 Domestic tourism expenditure

Italian T2, relating to domestic tourism expenditure, is the one that makes the most use of integrated sources in its compilation. Since domestic tourism expenditure is a component of internal tourism expenditure, together with inbound, part of the T2 compilation process involved taking the difference between inbound and internal tourism estimates, with the latter derived from National Accounts HCS data.

The contribution of THS to T2 compilation is very important, first of all for several pieces of non-monetary information provided (distinction between personal and business purposes of travel, which allowed a more accurate allocation of tourism expenditure between the competences of T2 and those of T4; number of trips and excursions, number of overnight stays, type of accommodation and prevalent transport, which provided the basis for the calculation of indicators and per capita expenditure), but also, in some cases, for the direct estimation of expenditure.

The estimation of domestic expenditure on air transport uses the SUTs of the National Accounts, in particular the data on final consumption, adjusted to exclude the share imputed to imports; this adjustment is used to exclude services provided by non-resident carriers.

4.3.3 Outbound tourism expenditure

Outbound tourism expenditure is expenditure incurred by resident visitors outside their country of reference.

T3 completes the estimation of demand but does not contribute to estimating tourism value added and, therefore, to determining the tourism sector's role within the whole economy. From T3's point of view, outbound expenditure refers entirely to services provided by non-resident companies and benefits only the economies of the countries visited.

Almost all of the outbound estimate is obtained from the Bank of Italy survey data. The additional use of National Accounts data has enabled the correct allocation of expenditure for products 7, 9, and 10, as well as the estimation of the "Other product".

As with inbound tourism expenditure, the transport estimate for outbound tourism expenditure, on the one hand, refers to the expenses of Italian visitors within the country visited, assuming that the carrier used is always resident in the country visited; on the other hand, to the expenses of Italian visitors for the Italy-abroad-Italy route, if the carrier is not resident. As with this Table, the road transport figure has undergone the usual treatment to estimate the petrol-related share.

Outbound and inbound tourism expenditures approximate the Travel item of the BoP compiled by the Bank of Italy, although substantial differences persist.

Recalling what has already been said about the visitor (TSA perspective) being a conceptual category within the broader category of traveller (BoP perspective), the main differences can be summarised as follows:

- TSA excludes cross-border, seasonal, and other short-term workers; students or long-term patients; nomads and refugees; transport carrier crews; embassy staff; armed forces;
- International transport, included in the TSA Tables, is not within the scope of the Travel item of the BoP but within that of International Passenger Transport Services.

For a detailed breakdown of the differences, please refer to the international manuals (UNSD and UNWTO 2010).

4.3.4 Other components and internal tourism consumption

Internal tourist consumption is largely derived from the sum of expenditure by resident and non-resident visitors and also includes “Other components”, whose most characteristic elements are business travel and second homes for tourism purposes. The estimate of business trips for the Italian TSA has been derived entirely from the perspective of the production units, primarily using National Accounts data on intermediate consumption, both domestic and imported, as provided by the Use Table.

In accordance with international recommendations, the intermediate consumption of Italian companies for accommodation, air transport, and travel agency services is entirely attributable to employees’ business travel expenses. For the Italian TSA, intermediate consumption for rail transport

was also included in the “Other components” column, as it was considered attributable to the same purpose. These were included in the total amount, with the understanding that the risk of overestimation from using accounting data without imputation was preferable to the possible underestimation that would result from using data insufficiently robust for this purpose.

However, expenses incurred during a business trip also include services (first and foremost, catering) and goods purchased for personal use. In practice, even catering is treated as intermediate consumption and, therefore, falls under the “Other components” if payment is made with the company credit card. In this case, there is, in fact, no disbursement by the employee, unlike with accommodation services and travel agencies. In practice, however, it is extremely difficult to make such a distinction. Excluding this case, the catering product in the “Other components” column of Table 4 should be zero, because it is considered part of all expenditure for final consumption. For reasons of data presentability to an external user unfamiliar with accounting concepts, it was decided to also include this product in the “Other components” column of T4. These expenses are estimated in T2, as they are considered final consumption. From the perspective of inbound tourism, business travel expenses are borne by the companies in the foreign visitor’s country of origin. Therefore, from the perspective of the country compiling the TSA, the related expense is entirely considered final and therefore pertains to Table 1.

For second homes, see the relevant Section above.

Finally, the “Other components” of the Italian TSA also include expenses attributable to general government and non-profit institutions serving households that support tourism consumption. Data on NPISH are available in the accounting matrices, with the usual product-level detail. For government data, the selected services for ITSA, based on the COFOG⁸ classification of public expenditure, are listed in Table 4.4.

8 The classification of public expenditure by function used in National Accounts refers to the COFOG, an international classification adopted as a standard by ESA95. The COFOG is divided into 3 levels of analysis: the first level consists of ten divisions, each of which is divided into groups, which in turn are divided into classes. The first six divisions cover expenditure on collective interventions and services; expenditure of an individual type is included in the remaining divisions.

Table 4.4 - Individual final consumption of general government expenditure by function class of the COFOG

Class (III Liv.)	Class Description	Tourist content of the class
07.2.4	Paramedical services	Thermal baths or seawater treatments
08.1.1	Leisure activities	Operation or support of recreational facilities such as beaches, parks, camping areas, furnished accommodation, and swimming pools
08.2.1	Cultural activities	Operation or support of structures for cultural purposes such as museums, art galleries, monuments, zoological and botanical gardens. Production, operation or support of cultural events (concerts, theatrical productions, art exhibitions)

Source: Authors' processing

4.4 Reconciliation of estimates from the supply and demand side

The reconciliation of tourism supply and demand estimates is presented in Table 6 of the TSA (T6), where the adjustment element ensures coherence across the estimates under different evaluation criteria⁹.

The tourism share of product, defined as the ratio of tourism consumption to total production, is a fundamental indicator for assessing the tourism sector as a whole within the reference economy. It also serves, in time series, as a tool for constructing confidence intervals to analyse discrepancies between supply – and demand-side estimates, thereby helping to verify the accuracy of tourism estimates.

The comparison in T6 does not end with the calculation of the tourism share of the product. It is also used to calculate the corresponding tourism share of output for each industry in T5. The approach used for the Italian TSA to determine the tourism share of the value of tourism industry output has taken into account both a product and an industrial perspective. For example, for the accommodation services industry, a specific assessment of its tourism production was carried out product by product, on the basis that, in the context of accommodation facilities, the use of other services that may be produced, in addition to the characteristic ones, has a tourism impact comparable to that of the main product.

⁹ The tourist offer is estimated at basic prices, whereas demand is expressed at purchase prices. In T6, there are special columns – imports, net taxes and trade margins – which allow production to be expressed at purchase prices, so that it can be related to final demand.

For second homes, on the other hand, the tourism share column in value production is the same, yet it is estimated in T5, where housing services have already been estimated only for the part intended for tourism demand.

4.5 The other Tables of the TSA for Italy: Table 7 and Table 10

4.5.1 Sources and estimation of Table 7 and Table 10d

Employment linked to tourism activities is undoubtedly a fundamental variable in an in-depth analysis of the tourism sector's contribution to a country's economic development.

Employment in the tourism sector is characterised by pronounced seasonality, high variability in working conditions, considerable flexibility, and a lack of formality in employment contracts. This makes employment in the tourism sector structurally very different from that in other sectors. In addition, these specific features require different employment measures, each of which can better represent one aspect than another.

Table 7 of the TSA (T7) is broken down by tourism industry, for each of which the number of establishments and three employment measures are required: number of jobs, hours worked and Full-Time Equivalent (FTE) jobs. The former provides an overall measure of the sector's employment size, but for comparison between the tourism sector and other sectors or the rest of the economy, more standardised measures of labour input, such as FTE, which neutralise the effects of part-time work, are more appropriate.

It is important to emphasise that, as with the production in T5, employment in T7 does not refer only to the share of production that directly served tourism consumption; at the same time, it excludes employment in other industries with secondary tourist outputs.

The compilation of T7 for the Italian TSA begins with an estimate of the number of establishments, based on the ASIA-Local Units Register.

The following information, drawn from the National Accounts data archives, was required to estimate the occupational variables:

- a. number of jobs at the 4-digit level of the classification of economic activities (class);
- b. FTE jobs at the 4-digit level of the classification of economic activities (class);
- c. hours worked at 2-digit level of the Classification of Economic Activities (division).

Table 10d is an extension of the T7 column on the number of establishments, disaggregated into 9 employee-size classes. For its compilation, data from the ASIA-Local Units Register were used, by economic activity and by employee-size classes.

4.5.2 Sources and estimates of the other sub-Tables for non-monetary data (Table 10a, 10b, 10c)

The remaining sub-Tables that make up Table 10 of a TSA (T10) – Tables 10a, 10b and 10c – collect information on tourism flows or accommodation facilities, and, taken together, derive from the surveys of the Bank of Italy and Istat.

Table 10a, compiled from the aforementioned surveys, details the number of trips and overnight stays by visitor type and tourism flow type.

Table 10b, compiled entirely from the Bank of Italy Survey, breaks down the total inbound trips and overnight stays in Table 10a by transport type. Compared with the international standard in Table 10b, the Italian one differs in some headings, providing less detail in some cases (e.g., for air and sea transport) and a different structure in others (e.g., for road transport).

With reference to Table 10c, which is entirely compiled from the Istat Survey on the Capacity of tourism accommodation and details the number of accommodation facilities by type of accommodation, it should be noted that, for Division 55 of the ISIC Rev.4 international classification, the number of accommodation facilities reported in Table 10c differs significantly from that indicated in Table 10d. This difference arises from the different sources underlying the compilation of the two sub-Tables. Table 10c is based on data derived from the two Istat surveys on the accommodation sector. These surveys are subject to the Regulation on tourism statistics, and their statistical unit is the local-KAU, i.e., the production unit identified on the basis of its geographical location.

5. Concluding remarks

The almost exclusive use of official sources in compiling the Italian TSA ensured the quality and reliability of the final estimates relative to those of the national accounts. In addition, the adopted approach laid the methodological and informational foundations for replicability. However, in some cases, the need to explore new sources to improve estimates has emerged, for example in the case of inbound tourism expenditure for the item “Other product”, which appears, in our opinion, to be underestimated. Direct data are not yet available, so indirect imputation methods are required.

The comparison Table of tourism supply and demand has been an excellent tool for verifying the consistency of the national accounts data from which the TSA is derived. This is a typical feature of satellite accounting, as it is an extension of the central framework of national accounts and a representation of events contained within them but not visible.

In Italy, the compilation of the TSA at the national level is now a consolidated practice. At the sub-national level, however, the available data do not allow the implementation of a regional TSA, and its compilation therefore still represents a challenge for our country, but also an important step towards the possibility of developing increasingly territorially contextualised analyses of the tourism sector, whose ability to influence the economic and social structures of the areas visited makes it valuable to have a territorial measurement of the tourism sector comparable with that produced in the territorial accounting schemes.

References

Alivernini, A., E. Breda, and E. Iannario. 2014. *Survey on international tourism in Italy (1997-2012)*. Questioni di Economia e Finanza, Occasional papers, N. 220/2014. Roma, Italy: Bank of Italy. <https://www.bancaditalia.it/pubblicazioni/qef/2014-0220/220.pdf>.

European Commission, International Monetary Fund - IMF, Organisation for Economic Co-operation and Development - OECD, United Nations, and World Bank. 2009. *System of National Accounts 2008*. Washington, D.C., U.S.: United Nations. <https://unstats.un.org/unsd/nationalaccount/docs/sna2008.pdf>.

Eurostat. 2014. *Methodological manual for tourism statistics. Ver. 3.1. Manuals and Guidelines*. Luxembourg: Publications Office of the European Union. <https://ec.europa.eu/eurostat/web/products-manuals-and-guidelines/-/ks-gq-14-013>.

Eurostat. 2013. *European System of Accounts. ESA 2010*. Luxembourg: Publications Office of the European Union. <https://ec.europa.eu/eurostat/web/products-manuals-and-guidelines/-/ks-02-13-269>.

Eurostat. 2009a. *Tourism Satellite Accounts in the European Union. Volume 3: Practical Guide for the Compilation of a TSA: Directory of Good Practices*. Methodologies and Working papers. Luxembourg: Office for Official Publications of the European Communities. <https://op.europa.eu/en/publication-detail/-/publication/15056e6c-4505-4967-8e12-95a1e82a2d17/language-en>.

Eurostat. 2009b. *Tourism Satellite Accounts in the European Union Volume 1: Report on the implementation of TSA in 27 EU Member States*. Methodologies and Working papers. Luxembourg: Office for Official Publications of the European Communities. <https://ec.europa.eu/eurostat/web/products-statistical-working-papers/-/ks-ra-09-021>.

Eurostat. 2002. *European Implementation Manual on Tourism Satellite Accounts (TSA)*. Luxembourg: Eurostat. https://ec.europa.eu/eurostat/documents/747990/748067/TSA_EIM_FINAL_VERSION.pdf/896f9dab-b9fa-45c1-b963-3028a73b71c6.

International Monetary Fund - IMF, Statistics Department. 2004. *Revision of the Balance of Payments Manual. Fifth Edition (Annotated Outline)*. Washington, D.C., U.S.: IMF. <https://www.imf.org/external/np/sta/bop/pdf/ao.pdf>.

Istituto Nazionale di Statistica - Istat. 2025a. *Annual national accounts. Years 2023-2024*. Statistics Flash. Roma, Italy: Istat. <https://www.istat.it/wp-content/uploads/2025/09/ANNUAL-NATIONAL-ACCOUNTS-Years-2023-2024.pdf>.

Istituto Nazionale di Statistica - Istat. 2025b. *Registro statistico delle imprese attive - Anno 2023*. Tavole di dati. Roma, Italia: Istat. <https://www.istat.it/tavole-di-dati/registro-statistico-delle-imprese-attive-anno-2023/>.

Istituto Nazionale di Statistica - Istat. 2025c. *Indagine CAPI Viaggi e vacanze*. Informazioni sulle rilevazioni. Roma, Italia: Istat. <https://www.istat.it/informazioni-sulla-rilevazione/viaggi-e-vacanze-anno-2014/>.

Istituto Nazionale di Statistica - Istat. 2025d. *Tourist nights spent increased in the fourth quarter, 2024 sets a new record year*. Statistiche Today. Roma, Italy: Istat. <https://www.istat.it/en/press-release/tourist-flows-fourth-quarter-2024/>.

Istituto Nazionale di Statistica - Istat. 2009. *Classificazione delle attività economiche. Ateco 2007, derivata dalla Nace Rev. 2*. Metodi e Norme, N. 40/2009. Roma, Italia: Istat. https://www.istat.it/it/files/2022/03/volume_integrale_ATECO2007.pdf.

Maresca, S., M. Anzalone, and I. Piscitelli. 2014. "Methodology for the production assessment of tourism industries". *Rivista di statistica ufficiale/Review of official statistics*, N.3/2014: 61-96. Roma, Italy: Istat. <https://www.istat.it/produzione-editoriale/rivista-di-statistica-ufficiale-32014/>.

Ministero delle Attività Produttive - MAP, Direzione Generale per il Turismo, Istituto Nazionale di Statistica - Istat, Ufficio Italiano dei Cambi - UIC, and Centro Internazionale di Studi sull'Economia Turistica - Ciset. 2003. *Towards the implementation of the Tourism Satellite Account in Italy. Final Report*. Roma, Italy: MAP.

United Nations Commission of the European Communities, Eurostat, World Tourism Organization - UNWTO, and Organisation for Economic Co-operation and Development - OECD. 2010. *Tourism Satellite Account: Recommended Methodological Framework 2008*. Studies in Methods, Series F N. 80/Rev.1. New York, NY, U.S.: United Nations Publication. https://unstats.un.org/unsd/publication/seriesf/seriesf_80rev1e.pdf.

United Nations Statistics Division - UNSD. 2008. *Central Product Classification (CPC) Version 2*. Series M: Miscellaneous Statistical Papers, N. 77, Ver. 2. New York, NY, U.S.: United Nations. <https://unstats.un.org/unsd/classifications/Family/Detail/1073>.

United Nations Statistics Division - UNSD, and World Tourism Organization (UNWTO). 2010. *International Recommendations for Tourism Statistics 2008*. Studies in Methods, Series M N. 83/Rev.1. New York, NY, U.S.: United Nations Publication. https://unstats.un.org/unsd/publication/seriesm/seriesm_83rev1e.pdf.

The *Rivista di statistica ufficiale* publishes peer-reviewed articles dealing with cross-cutting topics: the measurement and understanding of social, demographic, economic, territorial and environmental subjects; the development of information systems and indicators for decision support; the methodological, technological and institutional issues related to the production process of statistical information, relevant to achieving official statistics purposes.

The *Rivista di statistica ufficiale* aims at promoting interactions and exchanges between and among researchers, stakeholders, policy-makers and other users who refer to official and public statistics at different levels, to improve data quality and enhance trust.

The *Rivista di statistica ufficiale* was born in 1992 as a series of monographs titled “*Quaderni di Ricerca Istat*”. In 1999, the series was entrusted to an external publisher, changed its name to “*Quaderni di Ricerca - Rivista di Statistica Ufficiale*” and started being published on a four-monthly basis. The current name was assumed from the Issue N. 1/2006, when the Italian National Institute of Statistics – Istat returned to being its publisher.

La Rivista di statistica ufficiale pubblica articoli, valutati da esperti, che trattano argomenti trasversali: la misurazione e la comprensione di temi sociali, demografici, economici, territoriali e ambientali; lo sviluppo di sistemi informativi e di indicatori per il supporto alle decisioni; le questioni metodologiche, tecnologiche e istituzionali relative al processo di produzione dell'informazione statistica, rilevanti per raggiungere gli obiettivi della statistica ufficiale.

La Rivista di statistica ufficiale promuove sinergie e scambi tra ricercatori, stakeholder, policy-maker e altri utenti che fanno riferimento alla statistica ufficiale e pubblica a diversi livelli, al fine di migliorare la qualità dei dati e aumentare la fiducia.

La Rivista di statistica ufficiale nasce nel 1992 come serie di monografie dal titolo “Quaderni di Ricerca Istat”. Nel 1999 la collana viene affidata a un editore esterno, cambia nome in “Quaderni di Ricerca - Rivista di Statistica Ufficiale” e diventa quadrimestrale. Il nome attuale è stato scelto a partire dal numero 1/2006, quando l'Istituto Nazionale di Statistica - Istat è tornato a esserne l'editore.

e-ISSN 1972-4829

p-ISSN 1828-1982