

Article

Improving the Representativeness of a Simple Random Sample: An Optimization Model and Its Application to the Continuous Sample of Working Lives

Vicente Núñez-Antón ^{1,*},[†],[‡] , Juan Manuel Pérez-Salamero González ^{2,‡} ,
Marta Regúlez-Castillo ^{1,‡}  and Carlos Vidal-Meliá ^{3,4,5,‡} 

¹ Department of Econometrics and Statistics (A.E. III), Faculty of Economics and Business, University of the Basque Country UPV/EHU, 48015 Bilbao, Spain; marta.regulez@ehu.eus

² Department of Financial Economics and Actuarial Science, Faculty of Economics, University of Valencia, 46022 Valencia, Spain; juan.perez-salamero@uv.es

³ Department of Financial Economics and Actuarial Science, Faculty of Economics, University of Valencia, 46022 Valencia, Spain; carlos.vidal@uv.es

⁴ Instituto Complutense de Análisis Económico (ICAE), Complutense University of Madrid, 28223 Madrid, Spain

⁵ Centre of Excellence in Population Ageing Research (CEPAR), UNSW, 2033 Sydney, Australia

* Correspondence: vicente.nunezanton@ehu.eus; Tel.: +34-94-601-3749

† Current address: Avenida Lehendakari Aguirre 83, 48015 Bilbao, Spain.

‡ These authors contributed equally to this work.

Received: 30 June 2020; Accepted: 23 July 2020; Published: 25 July 2020

Abstract: This paper proposes an optimization model for selecting a larger subsample that improves the representativeness of a simple random sample previously obtained from a population larger than the population of interest. The problem formulation involves convex mixed-integer nonlinear programming (convex MINLP) and is, therefore, NP-hard. However, the solution is found by maximizing the size of the subsample taken from a stratified random sample with proportional allocation and restricting it to a p -value large enough to achieve a good fit to the population of interest using Pearson's chi-square goodness-of-fit test. The paper also applies the model to the Continuous Sample of Working Lives (CSWL), which is a set of anonymized microdata containing information on individuals from Spanish Social Security records and the results prove that it is possible to obtain a larger subsample from the CSWL that (far) better represents the pensioner population for each of the waves analyzed.

Keywords: chi-square test; continuous sample of working lives; optimization; p -value; subsampling

1. Introduction

In practice, the success of any statistical analysis usually depends on asking the right questions or defining the right problem to be analyzed, and this includes accurately defining the population that is going to be used as a source of information. The researcher needs to carefully and completely define this population before collecting the sample and give a description of the members to be included. If the sample that the researcher has to work on is drawn by simple random sampling from a population larger than the target population (i.e., a subset of the previous set), it might not be representative of the target population as far as the variables of interest are concerned. This situation can be improved by using a sample obtained by stratified random sampling. If it is possible to obtain a subsample that is more representative of the population of interest than the simple random sample from which it is to be extracted, then all efforts should be directed towards obtaining such a subsample with the aim of achieving results of the highest possible quality.

However, a smaller sample might not be good enough to identify any relevant characteristics that may be present in the population of interest, so it is desirable to have a larger subsample, although this still needs to be of a manageable size. It is therefore vital for the sample selected to be both representative enough for an appropriate analysis and representative of the target population with regard to its main characteristics. A number of papers deal with the problem of selecting representative samples, including, for example, Ramsey and Hewitt [1], Grafström and Schelin [2], Kruskall and Mosteller [3–6], and Omair [7], among others.

The aim of this paper is to develop an optimization model for selecting a larger subsample that improves the representativeness of a simple random sample previously obtained from a population larger than the population of interest. The researcher in this process does not have to have all the data on the population of interest, but must be able to classify the population into homogeneous groups or strata as they would if they were using a stratified random design. Simple random sampling can be vulnerable to sampling error because the randomness of the selection may result in a sample that does not reflect the makeup of the population. A subsample designed using systematic and stratified techniques will fit the population of interest better than the original simple random sample. This method is more efficient than simple random sampling because it ensures the adequate representation of elements across strata as far as the variables of interest are concerned.

The problem formulation involves convex mixed-integer nonlinear programming (convex MINLP) and is, therefore, NP-hard (see [8,9]). However, the optimization model we propose finds a global solution by solving a nonlinear programming problem with just one decision variable that is a real positive number. Thus, the subsample is selected by maximizing the so-called “constant of proportionality”—i.e., maximizing the size of the subsample taken from a stratified random sample with proportional allocation—and restricting it to a p -value large enough to achieve a good fit to the population of interest using Pearson’s chi-square goodness-of-fit test. Doing this also ensures that the subsample will be contained within the initial simple random sample as well as in the population of interest.

In this paper, we propose an enumeration algorithm for finding the optimal global solution to the problem. By means of a simulation, we analyze the performance of the subsample selection method and the efficiency of the algorithm in solving the problem depending on whether the original simple random sample had a “bad” fit or a “fine” fit to the population. As will be later described, when the procedure is applied to many cases, it usually finds the optimal global solution within a reasonable time. This resolution time depends on the size of the original simple random sample and how good a fit it is to the target population.

Finally, we show the usefulness of the proposed mathematical optimization model and how its resolution procedure works by applying it to a real case using the Continuous Sample of Working Lives (CSWL, Muestra Continua de Vidas Laborales) (in addition, referred to by several authors as the Continuous Working Life Sample (CWLS) or the Continuous Survey of Working Lives (CSWL)). The CSWL comprises matched anonymized social security, income tax, and census records for a simple random sample containing 4% of Spanish contributors, pensioners, and unemployment benefit recipients. Starting from 2004, an edition of the CSWL dataset has been released every year. The application process for obtaining the CSWL data are simple, and approved users are allowed to work with the data on their own computers (see [10]). The work done by Pérez-Salamero González, Regúlez-Castillo, and Vidal-Meliá [11,12] draws attention to the existence of a lack of representativeness of the CSWL with regard to the population with a pension income for the waves 2005–2013 and, using an ad-hoc approach as a preliminary exploratory analysis, suggest that there is room for improvement. Applying our model to waves 2014–2017 provides us with larger subsamples included in the CSWL that are much more representative of the retired population in terms of pension type, gender, and age.

The structure of the paper is as follows. Section 2 proposes a convex MINLP model designed to select subsamples and presents the mathematical formulation used to solve the problem. Section 3 presents the algorithm for finding the global solution and verifies its effectiveness by means of

a simulation study. Section 4 shows the real-life application of our methodology to the CSWL. For the sake of clarity, after we provide a brief description of the CSWL, we divide this section into two subsections. The first one presents the optimization model tailored to suit the CSWL, while the second analyzes the results and the main implications. The paper ends with our conclusions, possible directions for future research, and two appendices. Appendix A is the Nomenclature Chapter and in Appendix B we prove the convexity of the chi-square statistic function in R_{+}^k .

2. The Optimization Model for Improving the Representativeness of a Simple Random Sample

To solve the problem, we need to obtain from the initial simple random sample, which was drawn from a population larger than the population of interest, a larger subsample that is more representative of the target population according to a statistical goodness-of-fit criterion, after carrying out a poststratification process (see Figure 1).

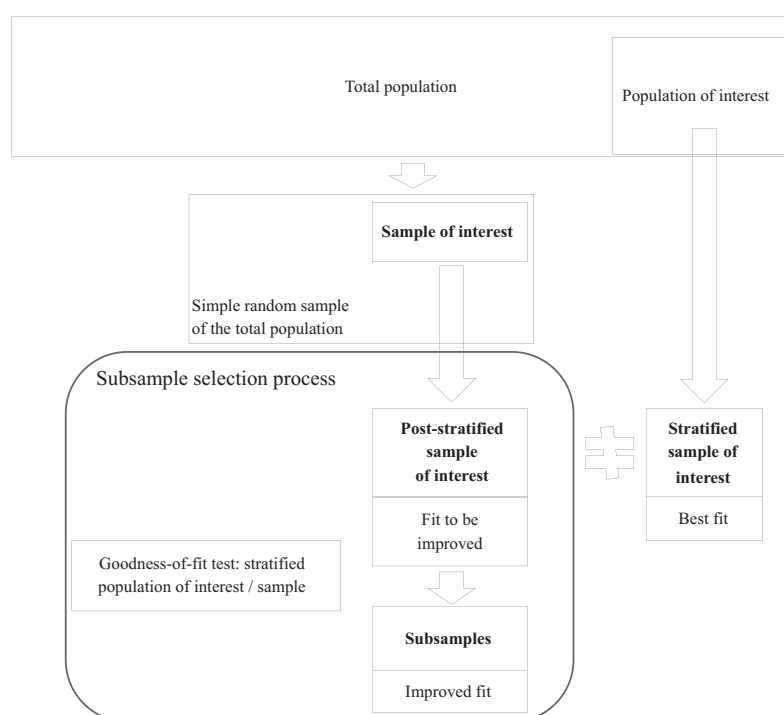


Figure 1. Graphical illustration of the problem.

The selection process has to take into account the following requirements:

1. The subsample must be as large as possible. We do not want to lose valuable information during the selection process, so we need to keep as many records as possible.
2. The subsample must form part of both the target population and the original simple random sample. This sounds an obvious requirement, but constraints need to be included so as to avoid outliers.
3. The elements to be included in each stratum of the subsample must take the form of a natural number (i.e., a non-negative integer).
4. The fit or representativeness with regard to the population under study must be improved. The optimization model should therefore include a goodness-of-fit test for the distribution by strata. It should also make it a requirement that the value of the statistic is smaller than the critical value given a predetermined significance level so as to avoid rejecting the null hypothesis

that the subsample has the same distribution as the population of interest. We use Pearson's goodness-of-fit test with the test statistic:

$$\chi^2 = \sum_{i=1}^k \frac{(O_i - E_i)^2}{E_i} \quad (1)$$

where the O_i are the observed values (those chosen from the simple random sample to build the subsample), and the E_i are the expected or theoretical values (those obtained from the distribution of the population of interest) and k is the number of strata for the variable of interest.

Equations (2)–(7) include the mathematical formulation for our subsample selection procedure aimed at improving the representativeness of a simple random sample and fulfilling the requirements listed above for the case of the univariate stratification (for multivariate stratification, the mathematical approach could be adapted using as many summation terms and sub-indices as stratifying variables to be considered).

$$\max_{n_i^{SUB}} \left\{ n^{SUB} = \sum_{i=1}^k n_i^{SUB} \right\}, \quad (2)$$

subject to constraints ($\forall i = 1, 2, \dots, k$):

$$\chi^2(n_1^{SUB}, \dots, n_k^{SUB}) = \sum_{i=1}^k \frac{(n_i^{SUB} - n_i^{EXP})^2}{n_i^{EXP}} \leq \chi_{(\alpha, r)}^2 \quad (3)$$

$$n_i^{EXP} = \frac{N_i}{N} \cdot n^{SUB} = \frac{N_i}{N} \cdot \sum_{i=1}^k n_i^{SUB} \quad (4)$$

$$0 \leq n_i^{SUB} \leq N_i \quad (5)$$

$$0 \leq n_i^{SUB} \leq n_i^{SRS} \quad (6)$$

$$n_i^{SUB} \in \mathbb{Z}, \quad (7)$$

where all of the terms included above are appropriately described in the Nomenclature Chapter in Appendix A. The first requirement for the subsample selection procedure involves the objective function (2), i.e., maximize the size of the subsample. Constraints (5) and (6) take into account the second requirement. The requirement to improve the goodness of fit is incorporated with constraints (3) and (4), using the chi-square goodness-of fit test once the significance level (α) has been chosen. Finally, constraints (5)–(7) incorporate the requirement that each stratum of the subsample must be a non-negative integer.

The objective function (2) is concave and linear, so it is continuous in a compact set, closed (the constraints are not strict), and bounded (all the decision variables are bounded below and above, by (5) and (6)). Therefore, if the set is not empty, i.e., if we establish a goodness-of-fit improvement criterion that can be satisfied using the data from the original sample, then a solution to the optimization problem will exist, given that the function is bounded (Weierstrass Theorem). We just need to apply a method to efficiently find it. Given constraints (3) and (7), this is a nonlinear integer programming problem. In addition, by considering the objective function and the constraints, it becomes a convex programming problem with decision variables bounded both above and below. The constraints, with the exception of (3), are linear inequalities, so they define convex sets. Without considering integer constraint (7), the function that calculates the value of the test statistic that leads to constraint (3) is

a convex function on R_+^k constrained with (5) and (6) because it can be shown that the associated Hessian matrix is positive semidefinite (see Appendix B).

The goodness-of-fit improvement constraint (3) can be rewritten equivalently as (8) by using a function that calculates the p -value using Pearson's chi-square test statistic and the number of strata. This results in the same convex set given by (3) for a fixed significance level equal to α . The value of α must be chosen in such a way that the problem is feasible, i.e., that a subsample that is a solution to the problem actually exists in the original simple random sample. The p -value is the probability that the test statistic will take on a value at least as extreme as the observed value, assuming that the null hypothesis is true, i.e., that the subsample has the same distribution as the population of interest. If the p -value is smaller than α , say 0.05, the null hypothesis is rejected. If it is greater than α , then the null hypothesis is not rejected:

$$pValue \left[\chi^2 \left(n_1^{SUB}, \dots, n_k^{SUB} \right); r \right] \geq \overline{pValue_{\min}},$$

where

$$\begin{aligned} pValue \left[\chi^2 \left(n_1^{SUB}, \dots, n_k^{SUB} \right); r \right] &= \\ &= 1 - \int_0^{\chi^2(n_1^{SUB}, \dots, n_k^{SUB})} \frac{\left(\frac{1}{2}\right)^{\frac{r}{2}}}{\Gamma\left[\frac{r}{2}\right]} x^{\left(\frac{r}{2}-1\right)} e^{-\frac{x}{2}} dx, \end{aligned} \quad (8)$$

r are the degrees of freedom and $\overline{pValue_{\min}}$ is the fixed significance level chosen to be equal to α .

Therefore, excluding the integer requirement for the variables, the opportunity set of the optimization problem is convex and thus can be framed in a wider class: convex mixed-integer nonlinear programming (convex MINLP).

2.1. The Optimization Model

Applying mathematical programming to sampling has been done before [13–22]. However, the aim of our optimization model is not to solve problems of optimum sample allocation in surveys like Neyman allocation [23]; cost-constrained optimal and precision-constrained optimal allocations seek to do. In our case, the population parameters are known, unlike in those cases in which they have to be estimated (e.g., population size, strata mean, etc.).

The optimization model for solving the problem set out in Equations (2)–(7) is based on stratified random sampling and uses proportional allocation to find a larger subsample from within the original simple random sample while improving representativeness, in line with work developed by authors such as Kontopantelis [24]. Proportional allocation in stratified sampling dates back to Bowley [25] and, given its simplicity, is very common in practice. It uses a sampling fraction in each stratum that is proportional to that of the total population. When no other information except stratum size (N_i) is available, allocating a given sample of size n in the different strata is proportional to their sizes. This implies that the sampling fractions are all equal and the same as the global sampling fraction, its value being the constant of proportionality $q = n/N$.

The problem of maximizing the size of the subsample will be solved by maximizing the constant of proportionality that depends on the number of elements in each stratum, after rewriting the constraints appropriately in terms of the new and only decision variable(q):

$$n_i^{SUB} = q \cdot N_i$$

The mathematical formulation of the optimization model in Equations (9)–(15) is the result of replacing the vector of decision variables, n_i^{SUB} , by q in the functions of the mathematical formulation of the problem in Equations (2)–(7).

Maximizing q is equivalent to maximizing the size of the subsample, even though what we are actually maximizing is $\hat{q}$, the adjusted constant of proportionality, given that we have to

consider the integer constraints for the number of units in each stratum. Constraint (12) guarantees constraints (5) and (7), while constraint (13) adapts constraint (6). Constraint (15) guarantees that q is positive and it is not larger than the constant of proportionality obtained by the ratio of the size of the simple random sample n^{SRS} with respect to the size of the target population N . Equations (9)–(15) include the mathematical formulation of the optimization model, maximizing the constant of proportionality and verifying the chi-square goodness-of-fit test follows:

$$\max_q \left\{ \hat{q}(q) = \frac{n^{SUB}(q)}{N} = \frac{\sum_{i=1}^k n_i^{SUB}(q)}{N} \right\}, \quad (9)$$

subject to constraints ($\forall i = 1, 2, \dots, k$):

$$pValue[\chi^2(q); r] \geq \overline{pValue_{\min}} \quad (10)$$

$$\chi^2(q) = \sum_{i=1}^k \frac{[n_i^{SUB}(q) - n_i^{EXP}(q)]^2}{n_i^{EXP}(q)} \quad (11)$$

$$n_i^{SUB}(q) = \text{Round}[q \cdot N_i] \quad (12)$$

$$0 \leq n_i^{SUB}(q) \leq n_i^{SRS} \quad (13)$$

$$n_i^{EXP}(q) = \frac{N_i}{N} \cdot n^{SUB}(q) = \frac{N_i}{N} \cdot \sum_{i=1}^k n_i^{SUB}(q) \quad (14)$$

$$0 \leq q \leq \frac{n^{SRS}}{N} \quad (15)$$

where all of the terms included above are appropriately described in the Nomenclature Chapter in Appendix A. The optimization model chosen to solve the problem moves from the MINLP framework in Equations (2)–(7), with a relatively high number of integer decision variables, to a non-differentiable nonlinear programming problem with only one non-negative real decision variable, the constant of proportionality (q), which means that the intermediate variables need to be integers and have non-differentiable functions. Given that this problem originates from the general problem shown in Equations (2)–(7), its mathematical properties guarantee that a solution exists as long as the set of possible solutions is not empty. This can be assured if the p -value, $\overline{pValue_{\min}}$, is chosen appropriately and is lower than that resulting from the stratified random sample contained in the original simple random sample. This solution gives a global maximum to the problem and provides the larger subsample using the data available in the original simple random sample, which is closer to the sample obtained by stratified random sampling with proportional allocation with knowledge of the distribution by strata of the target population.

Finally, it is important to mention that in some real problems it might be necessary to regroup those strata that do not reach the minimum size required for Pearson's chi-square goodness-of-fit test to have a good convergence to a χ^2 distribution. In such cases, the problem will always have an optimal and, in some extreme cases, a trivial solution. In the extreme case featuring a large size reduction, there might be just one stratum of regrouped expected, and observed values adding up to the same amount and these would provide the same theoretical and observed distributions, distorting the initial purpose of the test.

3. The Algorithm for Solving the Model and Some Simulation Results

In this section, we develop an efficient method for solving a specific convex MINLP using the optimization model set out in Section 2.1 (Equations (9)–(15)) and an algorithm with the following steps:

1. **Preliminary:** We verify that the null hypothesis of the chi-square goodness-of-fit test (H_0)—i.e., that the original simple random sample has the same distribution as the target population—is rejected. If not, we would not need to proceed. It is worth mentioning here that Pearson's chi-square depends on the size of the sample, so the impact of effect size on the rejection of the null hypothesis has to be taken into account (see [26–28]). Increasing emphasis has been placed on the use of effect size reporting when analyzing social science data.

To carry out this preliminary goodness-of-fit test, we have to:

- **Introduce the data (Input)** N_i from the target population, n_i^{SRS} from the simple random sample, and set the chosen $pValue_{min}$ that will be the significance level for the test.
 - **Compute the initial values (output):** size of the target population, N ; size of the original simple random sample, n^{SRS} ; expected values, n_i^{EXP} , given by (14); degrees of freedom, r ; and initial value of the constant of proportionality, $q_{start} = n^{SRS} / N$, $q = q_{start}$.
 - **Compute the test statistic and apply the decision rule as in (10) and (11).** The standardized effect size estimates also need to be reported. Cramer's V as interpreted by Cohen [29] is used for this purpose.
2. **Computing the observed values obtained by stratified sampling.** Using the value for the constant of proportionality, q , which in the first iteration of the algorithm is equal to the ratio between the size of the simple random sample and the size of the target population, $q = q_{start} = n^{SRS} / N$, we calculate the size of each stratum in the subsample using the nearest integer of $q \cdot N_i$ as in (12), $n_i^{SUB}(q) = Round[q \cdot N_i]$. The observed values obtained will be the same as those obtained using stratified random sampling from the target population with constant of proportionality q .
 3. **Fitting the observed values.** We compare the observed values obtained by stratified random sampling with proportional allocation from the target population, $n_i^{SUB}(q)$, which were found in the previous step, with those in the original simple random sample, n_i^{SRS} , using (13). If $n_i^{SUB}(q) > n_i^{SRS}$, we choose $n_i^{SUB}(q) = n_i^{SRS}$ for the subsample instead of the observed value obtained using $Round[q \cdot N_i]$ from the previous step. However, if $n_i^{SUB}(q) < n_i^{SRS}$, then we take the observed value for the subsample to be $n_i^{SUB}(q) = Round[q \cdot N_i]$, as obtained in the previous step.
 4. **Fitting the expected values.** After fitting the observed values, we calculate the total size of the subsample by adding up the number of units along the k strata, $\sum_{i=1}^k n_i^{SUB}(q) = n^{SUB}$. We then compute the expected values as $n_i^{EXP} = n^{SUB} \cdot N_i / N$ in order to obtain the same sum when adding up the expected values as we obtained with the observed values. If any stratum is found to be smaller than the minimum required by the test, which is usually 5, we will add it to the smallest nearest one until we reach the minimum number of elements. For Pearson's chi-square goodness-of-fit test to be valid, the sample size must be large enough to provide a minimum number of expected elements per category. Núñez-Antón, Pérez-Salamero González, Regúlez-Castillo, Ventura-Marco, and Vidal-Meliá [30] have developed functions for regrouping strata automatically no matter where they are located, thus enabling the goodness-of-fit test to be performed within an iterative procedure. The functions are written in Excel VBA (Visual Basic for Applications 7.1) and Mathematica, a registered trademark of Wolfram Research Inc. version 11.
 5. **Goodness-of-fit test.** Using (10) and (11), we now test the null hypothesis, H_0 —i.e., that the subsample has the same distribution as the target population—by comparing the fitted observed values with the fitted expected values obtained in steps 2 and 3 above. If the null hypothesis is not rejected, we can stop the algorithm because we have found the optimal $\hat{q}^*$, i.e., the distribution of the largest subsample contained in the original simple random sample that fits the population of interest. If the null hypothesis is rejected, we will proceed to the next step.

6. **New value of q .** To find the new value of q in order to start the process again, we now aim to obtain a reduction of q , requiring $q \geq 0$ that will provide the global optimal solution. The new value of q_j , q_j^{SUB} , used to start iteration j , is obtained by subtracting from the previous constant of proportionality, q_{j-1}^{SUB} , a step value (q_{step}) that is obtained in each case from the initial value of the constant of proportionality:

$$q_j^{SUB} = q_{j-1}^{SUB} - q_{step}; \quad q_{step} = 0.1/N$$

To shorten the time taken to find the solution, we incrementally reduce q as much as possible to the point where it does not reject the null hypothesis in the goodness-of-fit test. We then reverse the process in order to consider the immediately preceding value of q^{SUB} . From that value onwards, we are conducting a grid search in a finer set.

To analyze the performance of this algorithm as applied to solve the optimization model for the subsample selection procedure proposed in Section 2.1, we provide some results from a simulation study. This was carried out using MS Excel, a commonly-used non-specialist software, the advantages of which include the availability of pre-defined functions for calculating $\chi^2_{(\alpha,r)}$ and $pValue[\chi^2; r]$ in the goodness-of-fit test and the possibility of incorporating functions defined by the user in VBA.

We have generated 4000 populations and their corresponding simple random samples for the following scenarios: a “bad” fit of the sample to the population using the chi-square goodness-of-fit test and a “fine” fit with a better fit. The main characteristics of the simulations are:

- The number of strata is a random integer number between 2 and 20.
- The 4000 populations are generated in four blocks of 1000 as a function of the maximum size of the population strata—1000, 10,000, 100,000, and 1,000,000—given that the minimum size is 1 for all cases.
- The size of the simple random sample stratum is an integer number that results from rounding the product of the size of the corresponding stratum in the population by a percentage. This percentage is randomly selected from an interval that we have set to be [0%, 10%) for the “bad” fit and [3%, 5%) for the “fine” fit.
- For the 4000 simulations in each group, the optimization model is solved by means of our proposed algorithm with a given significance level, $\overline{pValue_{min}}$, of 5%, which is the most common significance level in practice. The optimization model was solved using MsExcel Professional Plus 2016, VBA 7.1, in a computer with an Intel Core i7-2600 Quad-Core Processor 3.4 GHz, 32 GB RAM and a Windows 7 Enterprise 64-bit system.

A summary of the simulation results can be seen in Table 1, which shows the number of global solutions obtained (row 2) and the average time taken to obtain the solution (row 4) for those cases in which an effect size (common practice when interpreting effect sizes is to use the benchmarks for “small”, “medium” and “large” effects suggested by Cohen [29]. Effect sizes may inform about practical significance, but they are not inherently meaningful.) (row 11) at least exists, even if it is small (row 12).

Towards the end of the table, we show the average values for Cramer’s V (row 9) and degrees of freedom (row 10) in those cases for which we have categorized the effect size (following Cohen’s methodology, a table has been constructed that categorizes the effect size as a function of the value of Cramer’s V and the degrees of freedom. This is available to any interested researchers upon request from the authors) (rows 11 to 14). These serve to provide general information about the simulated cases that have a global solution. As mentioned earlier, the proposed algorithm always finds the global solution.

The total average size (row 5) and the relative average size of the selected subsample with respect to the simple random sample (row 6) from which it is extracted are presented, as is the number of cases in which the subsample is greater than zero ($q > 0$) (row 8), even though at least one stratum

with zero units exists in the original simple random sample, which would prevent us from obtaining the subsample that best fits the population with a constant of proportionality equal to $\min\{n_i^{SRS}/N_i\}$. Finally, we include the size of the obtained subsample with respect to that resulting using as a constant of proportionality $q = \min\{n_i^{SRS}/N_i\}$ when it is not null (row 7).

Looking at the “bad” scenario, considering the 1000 simulated populations and simple random samples for each of the four cases according to maximum strata size, the null hypothesis in the chi-square goodness-of-fit test is rejected in 714, 938, 992, and 996, cases respectively (row 1). Given that the effect size is large in almost all cases (row 14), it could be stated that the test provides results to support the existence of statistically significant differences that are not exclusively due to sample size, i.e., the sample does not follow the same distribution as the population.

However, looking at the fine scenario for the first column of maximum strata size (1000), there is no rejection at all (row 1), whereas for the remaining three cases there is an increasing number of rejections of the null hypothesis but with a small effect size (row 12). Those rejections will therefore be mainly due to the size of the sample, although not completely, since some rejections do not even have this small effect size.

If we look at the simulation results for the “bad” scenario, we find that there are original simple random samples with null size strata, whereas the corresponding population strata are not null. However, the subsample solution that best fits the population is not the trivial one, $q^* = \min\{n_i^{SRS}/N_i\}$ because the size reached by the subsample was not 0% of the population (row 8).

For this group of simulations (“bad”), the subsamples obtained have a relative size (row 7) with respect to the smallest non-null size stratum that ranges from 1.32 times larger (32% larger) up to 4.64 times larger (364% larger), so the procedure gives solutions larger than the trivial subsample solution. As would be expected, the reduction in size between the original simple random sample and the subsample obtained (row 6) is greater in the “bad” scenario than in the “fine” one, given the better fit of the original simple random sample to the population in the latter group, and the time it takes to find the global solution is shorter for the “fine” cases than for the “bad” ones.

Table 1. Summary of the simulation results. Source: own. Two scenarios per strata size band: “bad” fit: [0%, 10%) and “fine” fit: [3%, 5%). $\overline{pValue}_{min} = 5\%$. Effect size (ES): small, medium, large. Simple random sample reject cases are verified cases rejecting the null hypothesis of the chi-square goodness-of fit, i.e., the original simple random sample does not have the same distribution as the target population.

		Cases by Maximum Strata Size							
		1000		10,000		100,000		1,000,000	
	Items	Bad	Fine	Bad	Fine	Bad	Fine	Bad	Fine
1	Simple random sample reject cases	714	-	938	437	992	901	996	989
2	Global solution cases with ES	714	-	936	434	977	783	977	840
3	Cases with regrouped stratum = 1	13	-	3	0	1	0	0	0
4	Average time seconds	3.34	-	23.26	4.19	50.98	11.21	101.15	13.89
5	Average subsample size	119.8	-	492.7	1696.8	3044.1	10,667.9	23,925.7	97,970.5
6	Relative average % $\frac{n^{SUB}}{n^{SRS}}$	65.81	-	41.01	93.49	28.56	88.92	22.9	84.09
7	Average $\frac{q^{SUB}}{Min\{n_i^{SRS}/N_i\}}$	4.64	-	3.39	1.20	2.02	1.12	1.32	1.06
8	$\{Min\{n_i^{SRS}\} = 0, n^{SUB} > 0\}$	306	-	107	15	21	2	7	0
9	Average Cramer's V	0.21	-	0.20	0.05	0.21	0.05	0.20	0.05
10	Average df	8.89	-	9.95	11.49	10.01	11.06	10.19	10.97
11	Type of ES	Large	-	Large	Small	Large	Small	Large	Small
12	Small cases	8	-	60	434	67	783	68	840
13	Medium cases	270	-	304	-	328	-	313	-
14	Large cases	436	-	572	-	582	-	596	-

In the specific cases or larger strata samples and, therefore, of the corresponding subsamples, the proposed algorithm takes more time to find the solution for the “bad” scenario, mainly because its solution implies taking larger steps towards verifying the test. Thus, the backward step will be a larger one, so that it actually starts at the preceding q^{SUB} value and then, using the same precision that the one used for smaller samples, start the forward search to find q^{SUB} using smaller steps. In this way, in the “bad” scenario and for strata samples up to 1,000,000, the time to find the solution is almost twice as much as the mean solution time of 101.15 s when compared to that of strata samples up to 100,000, which have a mean solution time close to 50.98 s. As for the “fine” scenario, this difference in mean solution times becomes smaller, going from 11.21 s for strata samples up to 100,000, when compared to the 13.89 s for larger strata samples, which represents a time increase of about 23.91%. This time scale increases as the size of the original simple random sample increases (row 4). Finally, for the “fine” scenario, the number of solutions obtained with just one regrouped stratum (row 3) is smaller than for the “bad” scenario, as is the number of cases with null-size strata in the original simple random sample that result in non-null-size optimal subsamples (row 8).

4. Applying the Model to the Continuous Sample of Working Lives (CSWL)

In this section, we apply the methodology developed in the previous sections to the Continuous Sample of Working Lives (CSWL), a set of microdata comprising anonymized information on individuals. This information is administrative data about people’s working lives and forms the basis of this sample taken from Spanish Social Security records. The sample reference population is defined as individuals who have had some connection with Social Security (through contributions or receiving some kind of pension or unemployment benefits) at any time during the year of reference. The total number of people involved during the year is larger than the number for any one particular date in that year, and one person can have several different simultaneous or successive relationships depending on their working situation during the year. Those who are outside the Social Security system (certain civil servants and those in the informal economy) are not included in this population. The first wave covers people who had an economic relationship of some sort with Social Security in 2004. However, each wave includes data covering the entire working and pension life of the people selected, starting in 1980. The sample is updated every year using information from the Social Security system, dating back to when computerized records began, and also from other administrative data sources, which record complementary information on individuals. The random sampling method is simple and uses no stratification. The sample provided by the INSS is 4% of the reference population and comprises around 1.2 million people. The population we want to study in this paper is the pensioner population categorized by age, gender, and type of pension for the period 2005 to 2017 (see [31]).

It seems difficult to hide the real importance of CSWL for several research fields. Data gathered from the CSWL have been widely used by academics to investigate a number of issues connected with the Spanish economy and relevant socioeconomic conditions. These include immigrants and immigration policy (see [32–38]), the Labor market (see [39–51]), the equity, sustainability, transparency and other aspects of the Public Pension System (see [52–66]), the impact of the economic crisis (see [67–70]), the usefulness of the MCVL (see [71–75]), the unemployment ([76–81]), disability and public health (see [82–89]), earnings, wealth and inequality (see [90–92]), retirement behavior/incentives (see [93–96]), and the gender gap (see [97,98]). The above reference list is not exhaustive but represents some of the most important published papers that have used the CSWL as the data source.

Despite the widespread use of the CSWL, very little attention has been given to the fact that administrative errors, misclassification problems and the type of sampling used might mean that the data selected are not representative of the population of interest. In studies on global ageing and the sustainability of the public pension system, the CSWL is one of the main datasets considered for research purposes in Spain. It is therefore important to know how representative it is of the pensioner

population (see [11,12]). After performing a poststratification process on the CSWL by type of pension, gender, and age, Pérez-Salamero González, Regúlez-Castillo, and Vidal-Meliá [12] revealed that the sample did not exactly fit the population of interest for waves 2005 to 2013 on the basis of the INSS (2006–2018) statistics report data (for brevity, we do not include this material here, but it is available upon request to the authors). In the same study, the authors show the gains in estimating pension expenditure for 2010 with a subsample distribution generated from the 2010 CSWL using an early stage of the procedure. These findings pointed out the necessity of further research towards a deeper analysis and development of the initial exploratory ad-hoc procedure. Once this has been done in this paper, we proceed to adapt and apply the proposed optimization model, together with the developed algorithm, to the waves 2005 to 2017 of the CSWL. We simultaneously consider the distribution by age, gender, and type of pension because this is a more general case than separately looking at the weight of each age cohort or the weight of one gender within a particular type of pension.

4.1. The Optimization Model Adapted to the CSWL

The optimization model considers the results obtained from adapting the mathematical approach laid out in Equations (9)–(15) and applies them to this real case involving the pensioner distribution in the CSWL. The procedure will take into account that the value of the test for the subsample to be selected is such that it does not reject the null hypothesis (that the subsample has the same distribution as the pensioner population by 31 December) against the alternative hypothesis (that the subsample does not have the same distribution as the pensioner population by 31 December). The procedure should therefore include a goodness-of-fit test on the distribution of the number of pensioners by age (18 cohorts covering 5-year intervals except for the last one, which represents 85 years and over), gender (male and female) and type of pension (permanent disability, retirement, widow(er)'s, orphan's, and family responsibilities (this is a special type of survivor benefit for family members included in the public Spanish Social Security System. It is not compatible with the beneficiary receiving other public pensions)), and this includes taking into account the associated p -values.

In order to guarantee that the subsample distribution fits the population of interest by pension type, gender, and age, constraint (10) is tailored specifically to this case, producing ten constraints that require that the null hypothesis not be rejected for each of the ten combinations of pension type and gender. The mathematical approach to this real application is detailed in Equations (16)–(22) as follows:

$$\max_q \left\{ \hat{q}(q) = \frac{n^{SUB}(q)}{N^{INSS}} \right\} \quad (16)$$

subject to constraints ($\forall i = 1, 2, \dots, 18; \forall j = 1, 2; \forall k = 1, \dots, 5$):

$$pValue_{j,k} \left[\chi^2_{j,k}(q); r_{j,k}(q) \right] \geq \overline{pValue}_{\min} \quad (17)$$

$$\chi^2_{j,k}(q) = \sum_{i \in I_{j,k}(q)} \frac{\left[\bar{n}_{i,j,k}^{SUB}(q) - \bar{n}_{i,j,k}^{EXP}(q) \right]^2}{\bar{n}_{i,j,k}^{EXP}(q)} \quad (18)$$

$$n_{i,j,k}^{SUB}(q) = \text{Round} \left[q \cdot N_{i,j,k}^{INSS} \right] \quad (19)$$

$$0 \leq n_{i,j,k}^{SUB}(q) \leq n_{i,j,k}^{CSWL} \quad (20)$$

$$n_{i,j,k}^{EXP}(q) = \frac{N_{i,j,k}^{INSS}}{N^{INSS}} \cdot n^{SUB}(q) \quad (21)$$

$$0 \leq q \leq \frac{n^{CSWL}}{N^{INSS}}, \quad (22)$$

where i is the index for the 18 cohorts into which the “age” variable has been categorized; j is the index corresponding to “gender” (male, female); and k is the index for the five types of pension benefit (permanent disability, retirement, widow(er)’s, orphan’s and family responsibilities), and:

$n^{SUB}(q) = \sum_{k=1}^5 \sum_{j=1}^2 \sum_{i=1}^{18} n_{i,j,k}^{SUB}(q)$: the size of the subsample.

$N_{i,j,k}^{INSS}$: the size of the stratum i, j, k in the target population. It is the number of pensioners in the population, with the sub-indices representing the corresponding groups by age, gender, and type of benefit. These data are obtained from INSS [99].

$N^{INSS} = \sum_{k=1}^5 \sum_{j=1}^2 \sum_{i=1}^{18} N_{i,j,k}^{INSS}$: the total number of beneficiaries. These data are obtained from INSS [99].

$n_{i,j,k}^{CSWL}$: the size of the stratum i, j, k in CSWL. It is the number of pensioners in CSWL, with the sub-indices representing the corresponding groups by age, gender, and type of benefit.

$n^{CSWL} = \sum_{k=1}^5 \sum_{j=1}^2 \sum_{i=1}^{18} n_{i,j,k}^{CSWL}$: the size of the CSWL.

$n_{i,j,k}^{EXP}(q) = \frac{N_{i,j,k}^{INSS}}{N^{INSS}} \cdot n^{SUB}(q)$: the expected size of stratum i, j, k in the subsample. This depends on the population relative frequency $N_{i,j,k}^{INSS} / N^{INSS}$ and the size of the subsample $n^{SUB}(q)$.

$n_{i,j,k}^{SUB}(q) = \text{Round}[q \cdot N_{i,j,k}^{INSS}]$: the size of the stratum i, j, k in the subsample (observed values).

Round: a function that returns the nearest integer to its argument.

$\bar{n}_{i,j,k}^{EXP}(q)$: the expected size of the regrouped stratum i, j, k in the subsample.

$\bar{n}_{i,j,k}^{SUB}(q)$: the size of the regrouped stratum i, j, k in the subsample (observed values).

$pValue_{j,k}[\chi_{j,k}^2(q); r_{j,k}(q)]$: a function that depends on the chi-square statistic and the degrees of freedom, both of which in the end also depend on the constant of proportionality and, above all, the values estimated and observed after regrouping (where applicable).

$\chi_{j,k}^2(q)$: the sample value for the chi-square statistic calculated in each iteration for each gender and type of pension (ten cases, i.e., five types of pension/2 genders). It evaluates the difference between the regrouped observed values and the expected values for cohort indices with five or more elements, avoiding those indices in which the regrouped cohort has no elements.

$r_{j,k}(q) = g(\bar{n}_{1,j,k}^{EXP}(q), \dots, \bar{n}_{18,j,k}^{EXP}(q)) - 1$: a function that returns the degrees of freedom in each iteration once the goodness-of-fit test is calculated for each type of pension and gender. It is equal to the expected number of regrouped cohorts minus 1, given that there are no parameters to estimate because the population distribution is already known.

$pValue_{min}$: a pre-established α level of statistical significance which is the criterion for the subsample to improve the goodness of fit to the population. This pre-established minimum p -value will be the same for all ten cases (five types of pension/2 genders) and has to be high in order to guarantee a better fit to the population than the value given by the CSWL.

$\bar{I}_{j,k}(q) = \{i \in I_{j,k}(q) \mid \bar{n}_{i,j,k}^{EXP}(q) \geq 5\}$: a set of indices for regrouped age cohorts that contain five or more elements, for each type of pension and gender.

$I_{j,k}(q) = \{1, 2, 3, \dots, 18\}$: a set of indices for age cohorts by type of pension and gender, which in all cases has 18 age cohorts.

5. Main Results

In this section, we provide the results of applying the above procedure to the CSWL with the aim of finding a larger subsample designed to improve the fit to the distribution of the pensioner population. Table 2 shows the global optimal solutions for 2005 to 2017, which range in size from 14.4% of the CSWL in 2013, associated with a $pValue_{min}$ of 0.95, to 99.39% in 2012, associated with a $pValue_{min}$ of 0.05. We have verified that all the solutions are global. What immediately attracts our attention in Figure 2 is the small size of the subsamples for 2013 and 2014 compared to all the other waves analyzed. Indeed, feasible solutions for the 2013 wave were found to range in size from 43.59%

of the CSWL, associated with a $\overline{pValue}_{min}$ of 0.05, to 14.4%, associated with a $\overline{pValue}_{min}$ of 0.95. Similar results are reported for the 2014 wave.

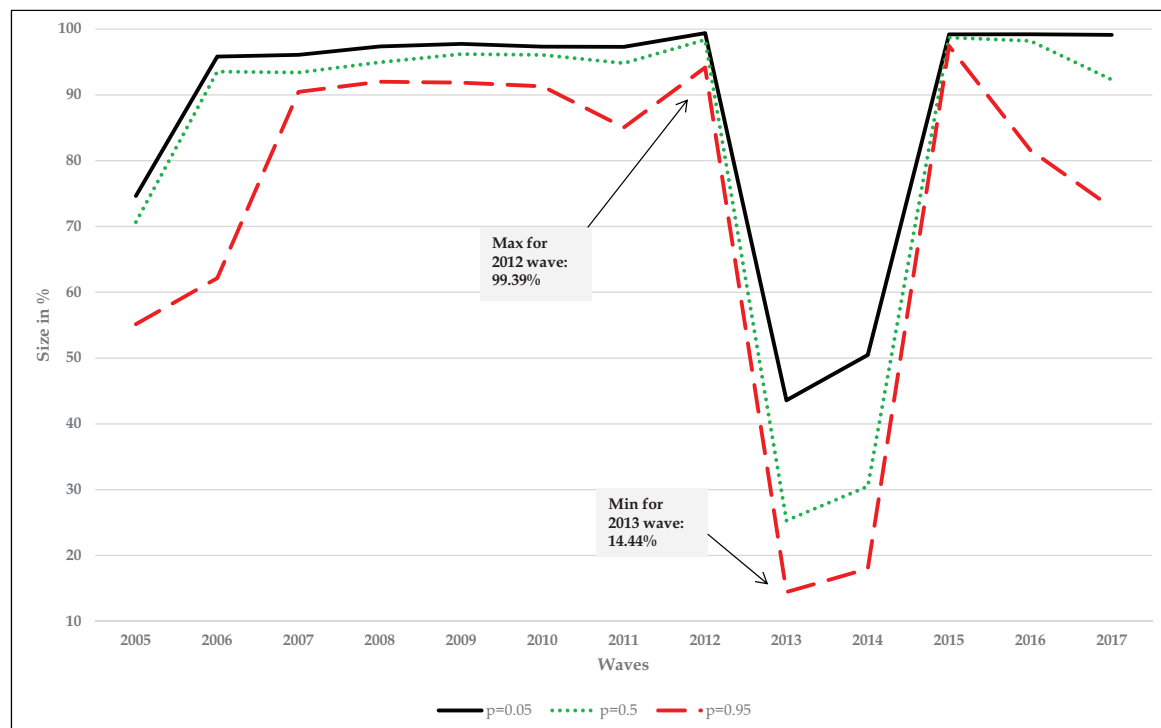


Figure 2. Size of the optimal subsamples for several p -values.

The explanation for this apparent anomaly in 2013 is that there are some cohorts (permanent disability males (0.21%) and females (1.67%), group 65–69 years) that are underrepresented in the subsample with respect to their real weight in the population (4%). For 2014, the explanation is the same, with some cohorts being underrepresented due to administrative errors. This would include the fact that pensioners over 65 with permanent disability benefits have not been reclassified, so that they are considered as retirement pension beneficiaries in the CSWL, but as disabled in the official population statistics. Since our procedure relies on maximizing the constant of proportionality that depends on the number of units in each stratum, the resulting percentages in the subsample for 2013 and 2014 are coherent.

For all the CSWL waves considered, the value of $\min\{n_{i,j,k}/N_{i,j,k}\}$ is 0.00%, and, therefore, the subsamples obtained by our proposed procedure are much larger than the one contained in the original CSWL that would have been obtained by stratified sampling using a constant of proportionality equal to $q = \min\{n_{i,j,k}/N_{i,j,k}\}$. The effect size obtained, following Cohen [29], can be classified as medium for the 2005 to 2011 waves with the exception of 2007, which is small. From 2012 to 2017, the size of the effect is negligible, so the rejection of the hypothesis that the CSWL has the same distribution as the population might be attributed to the large size of the sample.

Using Solver, a Microsoft Excel add-in program, with an IntelR CoreTM i7 – 2600 Processor (32 GB RAM, up to 3.40 GH) to solve the convex MINLP problem shown in Equations (16)–(22), the time needed to solve each of the 39 cases (13 years times three significance levels $\overline{pValue}_{min}$) ranged from 1.638 to 11.076 seconds (see Figure 3). This is a reasonable time for these types of problem with quite large dimensions, as would be the original problem of maximizing the size of the subsample instead of considering the constant of proportionality.



Figure 3. Time needed to obtain the subsamples for several p -values.

Using the 2005 to 2017 waves, Tables 3 and 4 show the p -values for the ten cases considered (five types of pension/2 genders) for the goodness-of-fit test for pensions in the subsample obtained using this procedure compared with those obtained for pensions in the CSWL. Overall, we find a lack of representativeness in the CSWL for total pensions in all the waves analyzed. Looking at the range of different pensions, the results seem to suggest that, for most of the waves analyzed, the CSWL does not fit the distribution of the population well in terms of pension type, gender and age for two types of pension benefit: permanent disability and widow(er)'s.

Once our proposed procedure is carried out, and after adjusting the distribution of total pensions with respect to the population using Pearson's goodness-of-fit test, the p -values obtained move towards or become equal to one. The procedure therefore provides subsamples with a better fit and many more observations than would be attained by a stratified random sample taken from the CSWL (for each of the years considered the distribution of pensions by pension type, gender, and age of the optimal-design larger subsample for the three significance levels considered is available upon request to the authors).

Table 2. Summary of results by $\overline{pValue}_{min}\%$. Source: own. D.I.U.: Dimension of the Integer Unrestricted problem. Number of combinations of integer values that the pensioner strata may take, within their bounds but without requiring the constraint that the null hypothesis of the statistical test is not rejected.

	n^{SUB}			$\hat{q} = \left(\frac{n^{SUB}}{n^{INSS}} \right) \%$			$\left(\frac{n^{SUB}}{n^{CSWL}} \right) \%$			Time (sec.)			D.I.U	Effect Size
Year	0.05	0.5	0.95	0.05	0.5	0.95	0.05	0.5	0.95	0.05	0.5	0.95		
2005	240,702	227,849	177,856	2.971	2.812	2.195	74.65	70.67	55.16	5.570	5.180	8.814	$7.188 \cdot 10^{339}$	Medium
2006	315,634	308,116	204,716	3.836	3.745	2.488	95.83	93.55	62.16	7.051	8.752	7.629	$3.564 \cdot 10^{339}$	Medium
2007	319,612	310,649	300,913	3.835	3.727	3.610	96.09	93.40	90.47	8.439	4.727	8.222	$7.216 \cdot 10^{335}$	Small
2008	329,204	321,031	311,082	3.887	3.790	3.673	97.37	94.95	92.01	8.939	7.738	6.864	$6.468 \cdot 10^{338}$	Medium
2009	335,665	330,298	315,426	3.897	3.835	3.662	97.76	96.20	91.87	8.549	10.359	8.487	$1.083 \cdot 10^{338}$	Medium
2010	339,831	335,482	318,778	3.885	3.835	3.644	97.33	96.08	91.30	7.613	7.191	5.585	$2.204 \cdot 10^{336}$	Medium
2011	343,317	334,590	300,129	3.871	3.772	3.384	97.30	94.82	85.06	8.860	8.767	6.848	$4.670 \cdot 10^{332}$	Medium
2012	358,046	354,395	339,223	3.975	3.935	3.766	99.39	98.38	94.17	1.638	7.910	7.519	$5.332 \cdot 10^{330}$	Negligible
2013	158,486	92,024	52,502	1.732	1.005	0.574	43.59	25.31	14.44	3.791	3.994	5.819	$4.939 \cdot 10^{327}$	Negligible
2014	186,109	112,455	66,322	2.005	1.212	0.715	50.48	30.50	17.99	6.646	3.307	5.648	$4.666 \cdot 10^{325}$	Negligible
2015	371,174	369,381	364,563	3.969	3.950	3.898	99.20	98.72	97.43	6.568	11.076	8.377	$4.588 \cdot 10^{326}$	Negligible
2016	376,017	372,166	309,241	3.973	3.932	3.267	99.22	98.20	81.60	9.313	10.967	9.719	$1.497 \cdot 10^{324}$	Negligible
2017	378,463	352,488	278,254	3.954	3.682	2.907	99.13	92.32	72.88	9.329	9.032	9.453	$1.012 \cdot 10^{321}$	Negligible

Table 3. Goodness-of-fit test (population/samples) p -values (CSWL and subsamples) M: Male; F: Female. Source: own.

							CSWL						
Type of Pension	2005	2006	2007	2008	2009	2010	2011	2012	2013	2014	2015	2016	2017
Perm. Disability M	0.000000	0.000000	0.261722	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
Perm. Disability F	0.000000	0.000000	0.022240	0.000000	0.000000	0.000000	0.000000	0.000000	0.187779	0.042312	0.000000	0.000000	0.000021
Retirement M	0.000000	0.000004	0.000001	0.000000	0.000000	0.000000	0.000000	0.005268	0.140715	0.013956	0.000743	0.001314	0.002751
Retirement F	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.001217	0.000120	0.000052	0.000011	0.000065	0.013070
Widower's M	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000
Widow's F	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000014	0.000006	0.000451	0.009670	0.008186	0.208145
Orphan's M	0.000000	0.005082	0.202030	0.261090	0.591337	0.462422	0.837630	0.944315	0.593409	0.370949	0.849506	0.000039	0.112380
Orphan's F	0.000000	0.118497	0.164789	0.141561	0.393848	0.561802	0.296694	0.117598	0.106684	0.285731	0.000101	0.000000	0.070073
Family Responsib. M	0.002755	0.111863	0.115631	0.396466	0.061782	0.428490	0.208140	0.323662	0.463862	0.327626	0.834403	0.830553	0.915809
Family Responsib. F	0.003573	0.249222	0.154051	0.021609	0.689061	0.841654	0.960454	0.333362	0.156821	0.659514	0.886466	0.344878	0.758455
Total Pensions	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000	0.000000

Subsample with p -value = $\overline{pValue}_{min}\%$ = 5%

Type of Pension	2005	2006	2007	2008	2009	2010	2011	2012	2013	2014	2015	2016	2017
Perm. Disability M	0.999062	0.997165	1.000000	1.000000	0.999995	0.999884	1.000000	1.000000	0.050022	0.050023	0.999079	0.999962	1.000000
Perm. Disability F	0.050167	1.000000	0.999935	1.000000	0.999449	0.999896	0.998295	0.986718	1.000000	0.95254	0.998444	0.999978	0.990695
Retirement M	1.000000	0.992407	0.999452	0.674773	0.105610	0.113748	0.050031	0.136429	0.999905	0.999937	0.433788	0.050327	0.050202
Retirement F	0.999948	0.999994	1.000000	0.734123	0.050021	0.050184	0.237496	0.536713	0.999268	0.999864	0.050727	0.068528	0.495790
Widower's M	1.000000	0.050289	0.050099	0.050117	0.094137	0.289216	0.361538	0.205785	1.000000	1.000000	0.90782	0.903889	0.926686
Widow's F	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	0.997186	1.000000	1.000000	0.886394	0.946287	0.999344
Orphan's M	1.000000	0.973781	0.974972	0.996229	0.997206	0.999660	0.999868	0.999343	1.000000	1.000000	0.998689	0.982116	0.968999
Orphan's F	1.000000	0.994200	0.992617	0.999344	0.996990	1.000000	0.999810	0.999604	1.000000	1.000000	1.000000	1.000000	0.999870
Family Responsib. M	0.485665	0.452965	0.870183	0.960915	0.834992	0.984285	0.739576	0.894736	1.000000	0.999996	0.984555	0.992232	0.987172
Family Responsib. F	0.592029	0.958494	0.915957	0.995480	0.997873	0.997951	0.994521	0.839826	0.999999	0.999998	0.999575	0.999999	0.999967
Total Pensions	1.000000	1.000000	1.000000	1.000000	0.999999	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000

Table 4. Goodness-of-fit test (population/samples) p -values (CSWL and subsamples) M: Male; F: Female. Source: own.

Subsample with $p\text{-value} = \overline{pValue}_{min}\% = 50\%$													
Type of Pension	2005	2006	2007	2008	2009	2010	2011	2012	2013	2014	2015	2016	2017
Perm. Disability M	0.999824	0.998681	1.000000	1.000000	1.000000	0.999995	1.000000	1.000000	0.500135	0.500169	0.999989	1.000000	1.000000
Perm. Disability F	0.500237	1.000000	1.000000	1.000000	0.999988	0.999996	0.999946	0.999053	1.000000	0.999968	0.999604	0.999998	1.000000
Retirement M	1.000000	0.999888	1.000000	1.000000	0.993192	0.500103	0.500086	0.742335	0.999991	0.999987	0.834558	0.500022	0.500002
Retirement F	0.999996	1.000000	1.000000	1.000000	0.974744	0.831733	0.999998	0.998129	0.999696	0.999979	0.500139	0.888871	0.999492
Widower's M	1.000000	0.500484	0.500491	0.500446	0.500167	0.704613	0.973377	0.500940	1.000000	1.000000	0.974067	0.994449	1.000000
Widow's F	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	0.999922	1.000000	1.000000
Orphan's M	1.000000	0.993038	0.990978	0.999867	0.999881	0.999983	1.000000	0.999930	1.000000	1.000000	0.999604	0.99941	1.000000
Orphan's F	1.000000	0.998723	0.998609	0.999976	0.998971	1.000000	0.999987	0.999933	1.000000	1.000000	1.000000	1.000000	1.000000
Family Responsib. M	0.596180	0.516585	0.926654	0.981715	0.865071	0.991481	0.815663	0.908726	1.000000	0.999994	0.986995	0.95072	0.996236
Family Responsib. F	0.730443	0.983661	0.965339	0.998284	0.999231	0.999351	0.997099	0.879598	1.000000	1.000000	0.999852	1.000000	1.000000
Total Pensions	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000
Subsample with $p\text{-value} = \overline{pValue}_{min}\% = 95\%$													
Type of Pension	2005	2006	2007	2008	2009	2010	2011	2012	2013	2014	2015	2016	2017
Perm. Disability M	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	0.975219	0.950021	1.000000	1.000000	1.000000
Perm. Disability F	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	0.999985	1.000000	1.000000
Retirement M	1.000000	1.000000	1.000000	1.000000	1.000000	0.950005	0.999992	0.996452	0.999993	0.999999	0.950024	0.950063	0.950000
Retirement F	1.000000	1.000000	1.000000	0.999996	1.000000	0.999897	0.999986	0.999985	0.999986	0.999988	0.962697	1.000000	0.999905
Widower's M	1.000000	1.000000	0.950055	0.950193	0.999535	0.999988	1.000000	0.999689	1.000000	1.000000	0.999129	1.000000	1.000000
Widow's F	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000
Orphan's M	1.000000	1.000000	0.995183	0.999997	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	0.999970	1.000000	1.000000
Orphan's F	1.000000	1.000000	0.999779	1.000000	0.999919	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000
Family Responsib. M	0.950041	0.950010	0.962938	0.992756	0.950023	0.999119	0.950023	0.950006	0.999987	0.998926	0.989983	0.999760	0.999937
Family Responsib. F	0.987717	0.999995	0.984731	0.999267	0.99988	0.999992	0.999793	0.950267	0.999999	1.000000	0.999985	1.000000	1.000000
Total Pensions	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000	1.000000

6. Discussion

This paper contains our proposal for an optimization model for improving the representativeness of an already available simple random sample obtained from a larger population than the population of interest, and the development of a novel methodology for selecting large subsamples can be considered the main contribution of this research. To our knowledge, neither the criterion used to select the subsamples nor the optimization model has been comprehensively explored or addressed in the literature.

The optimization model chosen to solve the problem moves from an MINLP framework with a relatively high number of integer decision variables to a non-differentiable nonlinear programming (NLP) problem with only one non-negative real decision variable: the constant of proportionality. We develop a method to efficiently solve a specific convex MINLP using the proposed optimization model by means of an algorithm that we have proved always finds the global solution. Using a simulation study, the procedure has been shown to work well in different scenarios as regards the goodness-of-fit of the simple random sample to the target population with an efficient use of time.

Another important contribution of this research is the real application of the procedure to the CSWL, a dataset widely used by a broad range of social science researchers comprising a simple random sample obtained from Spanish Social Security records. This dataset has become a baseline for researchers as it provides invaluable information about working lives and enables in-depth studies to be made of many aspects of the Spanish pension system that were previously overlooked. The methodology developed in this paper is applied to all the waves available 2005–2017 of the CSWL at the time of writing (may 2019). We find that, overall, the CSWL lacks representativeness when all pensions in all the waves analyzed are considered. Looking at the different types of pension, the results seem to suggest that most of the CSWL waves analyzed do not fit the distribution of the population well in terms of pension type, gender, and age for two types of pension benefit: permanent disability and widow(er)'s.

The application of the adapted optimization model to the 2005–2017 waves of the CSWL shows that larger subsamples can be obtained that will satisfy the chi-square goodness-of-fit test with associated *p*-values close to one. It can, therefore, be concluded that, for all the waves considered, it is possible to select large subsamples from the CSWL that better represent the pensioner population than the CSWL's own dataset, with a better fit to the distribution of the population's pensioners by type of pension, age and gender. In addition, last but not least, with this procedure, the users can choose between the desired goodness-of-fit and the size of the subsample they want, thereby allowing them a certain user-selection criterion to adapt the procedure to the research to be conducted.

From an applied perspective, since the CSWL has been widely used by researchers to investigate various issues in connection with the Spanish economy and its socioeconomic conditions, but without testing its representativeness with respect to the population of interest, the real example developed in this paper could and should be extended to other groups of interest such as contributors, recipients of unemployment benefits, immigrants, and/or the native population.

Finally, the model has been implemented using MS Excel, so as a future line of research we would like to use other optimization software and include in it the functions we have developed for regrouping strata automatically, and then compare the results.

Author Contributions: Data curation, J.M.P.S.G.; formal analysis, V.N.A., J.M.P.S.G., M.R.C. and C.V.M.; investigation, V.N.A., J.M.P.S.G., M.R.C. and C.V.M.; methodology, V.N.A., J.M.P.S.G., M.R.C. and C.V.M.; software, J.M.P.S.G.; writing—original draft, V.N.A., J.M.P.S.G., M.R.C. and C.V.M.; writing—review and editing, V.N.A., J.M.P.S.G., M.R.C. and C.V.M. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by the Ministerio de Economía y Competitividad (Spain) and the Basque Government for projects ECO2015-65826-P, IT 793-13 and IT1336-19, respectively, and Ministerio de Economía y Competitividad, Agencia Estatal de Investigación (AEI), Fondo Europeo de Desarrollo Regional (FEDER), the Department of Education of the Basque Government (UPV/EHU Econometrics Research Group), Universidad del País Vasco UPV/EHU and Generalidad Valenciana (Valencian Government) under research grants MTM2016-74931-P (AEI/FEDER, UE), IT-642-13, IT1359-19, UFI11/03, and AICO/2019/075.

Acknowledgments: The authors are grateful for the comments received at the Seventh International Conference MAF 2016, the 1st Workshop on Pensions and Insurance and the XXVIIth ASEPUMA Workshop—XVth International Meeting, plus all those made by several colleagues to previous versions of this article. The authors thank Peter Hall for his English support. The authors wish to thank the editor and the anonymous referees for providing thoughtful comments and suggestions which have led to substantial improvement in the presentation of the material in this paper.

Conflicts of Interest: The authors declare no conflict of interest.

Abbreviations

The following abbreviations are used in this manuscript:

MDPI	Multidisciplinary Digital Publishing Institute
DOAJ	Directory of Open Access Journals
TLA	Three Letter Acronym
LD	Linear Dichroism

Appendix A. Nomenclature Chapter

For the terms included in Equations (2)–(7), we have that:

n_i^{SUB} : the size of stratum i in the subsample (observed values).

n^{SUB} : the size of the subsample.

n_i^{EXP} : the expected size of stratum i in the subsample. This depends on the population relative frequency $\frac{N_i}{N}$ and the size of the subsample n^{SUB} .

N_i : the size of the stratum i in the target population.

N : the size of the target population.

n_i^{SRS} : the size of stratum i in the simple random sample.

n^{SRS} : the size of the simple random sample.

$\chi^2(n_1^{SUB}, \dots, n_k^{SUB})$: the chi-square goodness-of-fit test statistic.

$\chi^2_{(\alpha, r)}$: the tabulated value of the chi-square distribution with r degrees of freedom and a significance level equal to α .

k : the number of strata for the variable of interest.

$r = k - 1$: the degrees of freedom equal to the number of strata minus 1, given that in this case there are no parameters to be estimated because the population distribution is known.

Z : the set of integer numbers.

For the terms included in Equations (10)–(15), we have that: $n_i^{SUB}(q)$: $\text{Round}[q \cdot N_i]$: the size of stratum i in the subsample (observed values).

Round : function that rounds its argument to the nearest integer.

$n^{SUB}(q) = \sum_{i=1}^k n_i^{SUB}(q)$: the size of the subsample.

$n_i^{EXP}(q) = \frac{N_i}{N} \cdot n^{SUB}(q)$: the expected size of stratum i in the subsample. This depends on the population relative frequency $\frac{N_i}{N}$ and the size of the subsample $n^{SUB}(q)$.

k : the number of strata for the variable of interest.

N_i : the size of the stratum i in the target population.

N : the size of the target population.

n_i^{SRS} : the size of stratum i in the simple random sample.

$\chi^2(q)$: the chi-square goodness-of-fit test statistic.

$r = k - 1$: the degrees of freedom equal to the number of strata minus 1, given that in this case there are no parameters to be estimated because the population distribution is known.

$p\text{Value}[\chi^2(q); r]$: a function that calculates the p -value given the value of the test statistic, $\chi^2(q)$ and the degrees of freedom, r , as in (8).

$p\text{Value}_{\min}$: a pre-established α level of statistical significance which is the criterion for the subsample to improve the goodness of fit to the population.

Appendix B. Proof of the Convexity of the Chi-Square Statistic Function in R_+^k

In Pearson's chi-square goodness-of-fit test statistic, the expected values depend on known population frequencies, on the size of the sample and, therefore, on the observed values or absolute frequencies that add up to the sample size. This allows us to write the statistic as a function of the observed values. Considering the observed and expected values of a variable as the corresponding units of a given stratum from a stratified sample, the mathematical expression for the statistic is:

$$\chi^2(n_1^{SUB}, \dots, n_k^{SUB}) = \sum_{i=1}^k \frac{(n_i^{SUB} - n_i^{EXP})^2}{n_i^{EXP}} = \sum_{i=1}^k \frac{\left[n_i^{SUB} - \left(\frac{N_i}{N}\right) \sum_{j=1}^k n_j^{SUB}\right]^2}{\left(\frac{N_i}{N}\right) \sum_{j=1}^k n_j^{SUB}} \quad (A1)$$

This statistic is a function of k variables and it can be written as the sum of k functions. That is,

$$f_i(n_1^{SUB}, \dots, n_k^{SUB}) = \frac{\left[n_i^{SUB} - \left(\frac{N_i}{N}\right) \sum_{j=1}^k n_j^{SUB}\right]^2}{\left(\frac{N_i}{N}\right) \sum_{j=1}^k n_j^{SUB}} \quad (A2)$$

These functions are the ratio of a second-degree polynomial in $(n_1^{SUB}, \dots, n_k^{SUB})$. In particular, the numerator is a positive semidefinite quadratic form, and the denominator is a linear function of the same variables. The coefficients $\frac{N_i}{N}$ are the relative frequencies, which are known from the population and, in addition, they are constant. Given that the denominator is not null for nonempty samples, if every term in the sum (i.e., f_i) is a convex function with respect to $(n_1^{SUB}, \dots, n_k^{SUB}) \in R_+^k$, the test statistic itself will be also a convex function. That is,

$$\begin{aligned} f_i(n_1^{SUB}, \dots, n_k^{SUB}) &= \frac{\left[n_i^{SUB} - \left(\frac{N_i}{N}\right) \sum_{j=1}^k n_j^{SUB}\right]^2}{\left(\frac{N_i}{N}\right) \sum_{j=1}^k n_j^{SUB}} \\ &= \frac{(n_i^{SUB})^2 + \left[\frac{N_i}{N} \sum_{j=1}^k n_j^{SUB}\right]^2 - 2 \cdot n_i^{SUB} \cdot \left[\frac{N_i}{N} \sum_{j=1}^k n_j^{SUB}\right]}{\left(\frac{N_i}{N}\right) \sum_{j=1}^k n_j^{SUB}} \\ &= \frac{(n_i^{SUB})^2}{\left(\frac{N_i}{N}\right) \sum_{j=1}^k n_j^{SUB}} + \frac{N_i}{N} \sum_{j=1}^k n_j^{SUB} - 2 \cdot n_i^{SUB} \end{aligned}$$

Since the last two terms in the expression above are linear functions of $(n_1^{SUB}, \dots, n_k^{SUB})$ and, thus, convex functions, the convexity of the function $f_i(n_1^{SUB}, \dots, n_k^{SUB})$ depends on the convexity of the first term given by:

$$g_i(n_1^{SUB}, \dots, n_k^{SUB}) = \frac{(n_i^{SUB})^2}{\left(\frac{N_i}{N}\right) \sum_{j=1}^k n_j^{SUB}} \quad (A3)$$

Therefore, the function will be convex in R_+^k if we have that:

$$g_i[\lambda x + (1 - \lambda)y] \leq \lambda g_i(x) + (1 - \lambda)g_i(y), \quad \forall x, y \in R_+^k, \quad \forall \lambda \in [0, 1] \quad (A4)$$

$$g_i[\lambda x + (1 - \lambda)y] - \lambda g_i(x) - (1 - \lambda)g_i(y) \leq 0, \quad \forall x, y \in R_+^k, \quad \forall \lambda \in [0, 1] \quad (A5)$$

Figure A1 graphically illustrates the non-convexity of $\chi^2(n_1^{SUB}, \dots, n_k^{SUB})$ in all R^2 .

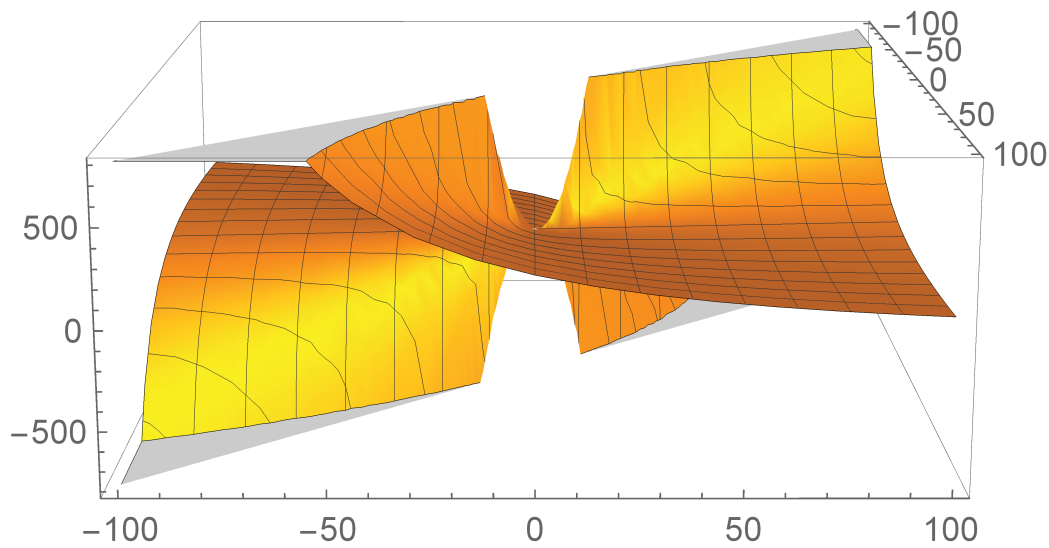


Figure A1. Non-convexity of the $\chi^2(n_1^{SUB}, \dots, n_k^{SUB})$ in R^2 .

Proof. We have that:

$$\begin{aligned} g_i[\lambda x + (1 - \lambda)y] &= \frac{[\lambda x_i + (1 - \lambda)y_i]^2}{\frac{N_i}{N} \sum_{j=1}^k [\lambda x_j + (1 - \lambda)y_j]} \\ &= \frac{\lambda^2 x_i^2 + 2\lambda(1 - \lambda)x_i y_i + (1 - \lambda)^2 y_i^2}{\frac{N_i}{N} [\lambda (\sum_{j=1}^k x_j) + (1 - \lambda) \sum_{j=1}^k y_j]} \end{aligned}$$

In addition, we also have that:

$$\begin{aligned} \lambda g_i(x) + (1 - \lambda)g_i(y) &= \lambda \frac{x_i^2}{\frac{N_i}{N} \sum_{j=1}^k x_j} + (1 - \lambda) \frac{y_i^2}{\frac{N_i}{N} \sum_{j=1}^k y_j} \\ &= \frac{\lambda x_i^2 (\sum_{j=1}^k y_j) + (1 - \lambda)y_i^2 (\sum_{j=1}^k x_j)}{\frac{N_i}{N} (\sum_{j=1}^k x_j) (\sum_{j=1}^k y_j)}, \end{aligned}$$

so that:

$$g_i[\lambda x + (1 - \lambda)y] - \lambda g_i(x) - (1 - \lambda)g_i(y) = \tag{A6}$$

$$= \frac{\lambda^2 x_i^2 + 2\lambda(1 - \lambda)x_i y_i + (1 - \lambda)^2 y_i^2}{\frac{N_i}{N} [\lambda (\sum_{j=1}^k x_j) + (1 - \lambda) \sum_{j=1}^k y_j]} - \frac{\lambda x_i^2 (\sum_{j=1}^k y_j) + (1 - \lambda)y_i^2 (\sum_{j=1}^k x_j)}{\frac{N_i}{N} (\sum_{j=1}^k x_j) (\sum_{j=1}^k y_j)} \tag{A7}$$

$$= \frac{[\lambda^2 x_i^2 + 2\lambda(1 - \lambda)x_i y_i + (1 - \lambda)^2 y_i^2] (\sum_{j=1}^k x_j) (\sum_{j=1}^k y_j)}{\frac{N_i}{N} [\lambda (\sum_{j=1}^k x_j) + (1 - \lambda) \sum_{j=1}^k y_j] (\sum_{j=1}^k x_j) (\sum_{j=1}^k y_j)} - \tag{A8}$$

$$- \frac{[\lambda x_i^2 (\sum_{j=1}^k y_j) + (1 - \lambda)y_i^2 (\sum_{j=1}^k x_j)] [\lambda (\sum_{j=1}^k x_j) + (1 - \lambda) (\sum_{j=1}^k y_j)]}{\frac{N_i}{N} [\lambda (\sum_{j=1}^k x_j) + (1 - \lambda) \sum_{j=1}^k y_j] (\sum_{j=1}^k x_j) (\sum_{j=1}^k y_j)} \tag{A9}$$

Since the denominator is positive for nonempty subsamples, $\sum_{j=1}^k x_j > 0$, $\sum_{j=1}^k y_j > 0$, $\lambda \in [0, 1]$, $\frac{N_i}{N} > 0$, the sign of the ratio depends solely on the sign of the numerator. Therefore, we have that:

$$\lambda^2 x_i^2 \left(\sum_{j=1}^k x_j \right) \left(\sum_{j=1}^k y_j \right) + 2\lambda(1-\lambda)x_i y_i \left(\sum_{j=1}^k x_j \right) \left(\sum_{j=1}^k y_j \right) + \quad (\text{A10})$$

$$+ (1-\lambda)^2 y_i^2 \left(\sum_{j=1}^k x_j \right) \left(\sum_{j=1}^k y_j \right) - \lambda^2 x_i^2 \left(\sum_{j=1}^k x_j \right) \left(\sum_{j=1}^k y_j \right) - \quad (\text{A11})$$

$$- \lambda(1-\lambda)x_i^2 \left(\sum_{j=1}^k y_j \right)^2 - \lambda(1-\lambda)y_i^2 \left(\sum_{j=1}^k x_j \right)^2 - (1-\lambda)^2 y_i^2 \left(\sum_{j=1}^k x_j \right) \left(\sum_{j=1}^k y_j \right) = \quad (\text{A12})$$

$$= 2\lambda(1-\lambda)x_i y_i \left(\sum_{j=1}^k x_j \right) \left(\sum_{j=1}^k y_j \right) - \lambda(1-\lambda)x_i^2 \left(\sum_{j=1}^k y_j \right)^2 - \quad (\text{A13})$$

$$- \lambda(1-\lambda)y_i^2 \left(\sum_{j=1}^k x_j \right)^2 = -\lambda(1-\lambda) \left[x_i \left(\sum_{j=1}^k y_j \right) - y_i \left(\sum_{j=1}^k x_j \right) \right]^2 \leq 0, \quad (\text{A14})$$

because $\lambda \in [0, 1]$.

Therefore, we conclude that, given that the function $\chi^2(n_1^{SUB}, \dots, n_k^{SUB})$ is the sum of convex functions, it is convex with respect to $(n_1^{SUB}, \dots, n_k^{SUB}) \in R_+^k$. $\square$

Figure A2 draws the function $\chi^2(n_1^{SUB}, n_2^{SUB})$ for non-negative values of (n_1^{SUB}, n_2^{SUB}) . Looking at Figure A2, we notice the convexity of the function for non-negative strata size, as we already proved above.

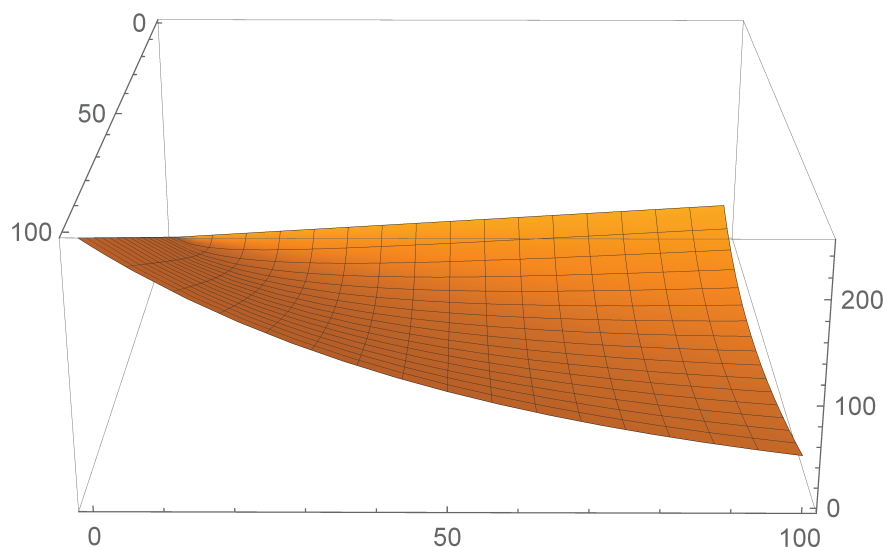


Figure A2. $\chi^2(n_1^{SUB}, n_2^{SUB})$ with $n_1^{SUB}, n_2^{SUB} \geq 0$; $N = N_1 + N_2 = 100 + 250$.

References

1. Ramsey, C.A.; Hewitt, A.D. A methodology for assessing sample representativeness. *Environ. Forensics* **2005**, *6*, 71–75.
2. Gr'afstrom, A.; Schelin, L. How to select representative samples. *Scand. J. Stat.* **2014**, *41*, 277–290.
3. Kruskall, W.; Mosteller, F. Representative sampling, I. *Int. Stat. Rev.* **1979**, *47*, 13–24.

4. Kruskall, W.; Mosteller, F. Representative sampling, II: Scientific literature, excluding statistics. *Int. Stat. Rev.* **1979**, *47*, 111–127.
5. Kruskall, W.; Mosteller, F. Representative sampling, III: The current statistical literature. *Int. Stat. Rev.* **1979**, *47*, 245–265.
6. Kruskall, W.; Mosteller, F. Representative sampling, IV: The history of the concept in statistics, 1895–1939. *Int. Stat. Rev.* **1980**, *48*, 169–195.
7. Omair, A. Sample size estimation and sampling techniques for selecting a representative sample. *J. Health Spec.* **2014**, *2*, 142–147.
8. Bonami, P.; Kilinç, M.; Linderoth, J. Algorithms and software for convex mixed integer nonlinear programs. In *Mixed Integer Nonlinear Programming. The IMA Volumes in Mathematics and its Applications*; Lee, J., Leyferr, S., Eds.; Springer: New York, NY, USA, 2012; Volume 154; pp. 1–39.
9. D’Ambrosio, C.; Lodi, A. Mixed integer nonlinear programming tools: An updated practical overview. *Ann. Oper. Res.* **2013**, *204*, 301–320.
10. MESS: Documentación Muestra Continua de Vidas Laborales: MCVL Guía. Madrid: Secretaría de Estado de la Seguridad Social. Ministerio de Trabajo, Migraciones y Seguridad Social. Available online: <http://www.seg-social.es/> (accessed on 12 March 2020).
11. Pérez-Salamero González, J.M.; Regúlez-Castillo, M.; Vidal-Meliá, C. Análisis de la representatividad de la MCVL: El caso de las prestaciones del sistema público de pensiones. *Hacienda Pública Esp.* **2016**, *217*, 67–130.
12. Pérez-Salamero González, J.M.; Regúlez-Castillo, M.; Vidal-Meliá, C. The continuous sample of working lives: Improving its representativeness. *SERIEs* **2017**, *8*, 43–95.
13. Cochran, W.G. *Sampling Techniques*; Wiley: New York, NY, USA, 1977.
14. Särndal, C.E.; Swensson, B.; Wretman, J. Model Assisted Survey Sampling. *Springer Series in Statistics*; Springer: New York, NY, USA, 1992.
15. Valliant, R.; Gentle, J.E. An application of mathematical programming to sample allocation. *Comput. Stat. Data An.* **1997**, *25*, 337–360.
16. Baillargeon, S.; Rivest, L.P. A general algorithm for univariate stratification. *Int. Stat. Rev.* **2009**, *77*, 331–344.
17. Díaz-García, J.A.; Ramos-Quiroga, R. Optimum allocation in multivariable stratified random sampling: Stochastic matrix mathematical programming. *Stat. Neerl.* **2012**, *66*, 492–511.
18. Gupta, N.; Sana Ifthekar, S.; Bari, A. Fuzzy goal programming approach to solve nonlinear bi-level programming problem in stratified double sampling design in the presence of non-response. *Int. J. Sci. Eng. Res.* **2012**, *3*, 1–9.
19. Valliant, R.; Dever, J.; Kreuter, F. *Practical Tools for Designing and Weighting Survey Samples*; Springer: New York, NY, USA, 2013.
20. Díaz-García, J.A.; Ramos-Quiroga, R. Optimum allocation in multivariable stratified random sampling: A modified Prékopa’s approach. *J. Math. Mod. Algorithms* **2014**, *13*, 315–330.
21. Gupta, N.; Ali, I.; Bari, A. An optimal chance constraint multivariate stratified sampling design using auxiliary information. *J. Math. Mod. Algorithms Oper. Res.* **2014**, *13*, 341–352.
22. De Moura Brito, J.A.; Do Nascimento Silva, P.L.; Silva Semaan, G.; Maculan, N. Integer programming formulations applied to optimal allocation in stratified sampling. *Surv. Methodol.* **2015**, *41*, 427–442.
23. Neyman, J. On the two different aspects of the representative method: The method of representative sampling and the method of purposive sampling. *J. R. Stat. Soc.* **1934**, *97*, 558–625.
24. Kontopantelis, E. A greedy algorithm for representative sampling: Repsample in Stata. *J. Stat. Softw.* **2013**, *56*, 1–18.
25. Bowley, A.L. Measurement of precision attained in sampling. *B. Int. Statist. Inst.* **1926**, *22*, 6–62.
26. Berkson, J. Some difficulties of interpretation encountered in the application of the chi-square test. *J. Am. Stat. Assoc.* **1938**, *33*, 526–536.
27. Wang, C. *Sense and Nonsense of Statistical Inference: Controversy, Misuse and Subtlety*; Marcel Dekker: New York, NY, USA, 1993.
28. Lin, M.; Lucas, H.C.; Shmieli, G. Research commentary: Too big to fail: large samples and the *p*-value problem. *Inform. Syst. Res.* **2013**, *24*, 906–917.
29. Cohen, J. *Statistical Power Analysis for the Behavioral Sciences*; Erlbaum: Hillsdale, NJ, USA, 1988.
30. Núñez-Antón, V.; Pérez-Salamero González, J.M.; Regúlez-Castillo, M.; Ventura-Marco, M.; Vidal-Meliá, C. Automatic regrouping of strata in the goodness-of-fit chi-square test. *SORT* **2019**, *43*, 113–142.

31. DGOSS: Muestra Continua de vidas Laborales, 2005–2017. Madrid: Secretaría de Estado de la Seguridad Social. Ministerio de Trabajo, Migraciones y Seguridad Social. Available online: <http://www.seg-social.es/wps/portal/wss/internet/EstadisticasPresupuestosEstudios/Estadisticas/> (accessed on 12 March 2020).
32. Agliari, E.; Barra, A.; Contucci, P.; Sandell, R.; Vernia, C. A stochastic approach for quantifying immigrant integration: The Spanish test case. *New J. Phys.* **2014**, *16*, doi: 10.1088/1367-2630/16/10/103034.
33. Barra, A.; Contucci, P.; Sandell, R.; Vernia, C. An analysis of a large dataset on immigrant integration in Spain. The statistical mechanics perspective on social action. *Sci. Rep.* **2014**, *4*, 4174.
34. Carrasco, R.; García Pérez, J.I. Employment dynamics of immigrants versus natives: evidence from the boom-bust period in Spain, 2000–2011. *Econ. Inq.* **2015**, *53*, 1038–1060.
35. De Pedraza, P.; Villacampa González, A.; Muñoz de Bustillo Llorente, R. Immigrants' employment situations and decent work determinants in the Spanish labour market. *Int. J. Humanit. Soc. Sci.* **2012**, *2*, 1–19.
36. Gómez Tello, A.; Nicolini, R. Immigration and productivity: A Spanish tale. *J. Prod. Anal.* **2017**, *47*, 167–183.
37. González, L.; Ortega, F. How do very open economies adjust to large immigration flows? Evidence from Spanish regions. *Labour Econ.* **2011**, *18*, 57–70.
38. Solé, M.; Díaz Serrano, L.; Rodríguez, M. Disparities in work, risk and health between immigrants and native-born Spaniards. *Soc. Sci. Med.* **2013**, *76*, 179–187.
39. Alonso Domínguez, A. Labor transitions of Spanish workers: A flexicurity approach. *Rev. Int. Org.* **2012**, *9*, 121–143.
40. Álvarez de Toledo, P.; Núñez, F.; Usabiaga, C. An empirical analysis of the matching process in Andalusian public employment agencies. *Hacienda Pública Esp.* **2011**, *198*, 67–102.
41. Álvarez de Toledo, P.; Núñez, F.; Usabiaga, C. An empirical approach on labour segmentation. Applications with individual duration data. *Econ. Model.* **2014**, *36*, 252–267.
42. Álvarez de Toledo, P.; Núñez, F.; Usabiaga, C. ¿Quién se empareja con quién en el mercado laboral español? Un análisis clúster basado en la muestra continua de vidas laborales. *Investigación Económica* **2017**, *76*, 3–182.
43. Álvarez de Toledo, P.; Núñez, F.; Usabiaga, C. Análisis “cluster” de los flujos laborales andaluces. *Rev. Estud. Reg.* **2013**, *97*, 195–221.
44. Cueto, B.; Rodríguez, V. Sheltered employment centres and labour market integration of people with disabilities: A quasi-experimental evaluation using Spanish data. In *Disadvantaged Workers. Empirical Evidence and Labour Policies*; Malo, M., Sciulli, D., Eds.; Springer: New York, NY, USA, 2014; pp. 65–91.
45. García Pérez, J.I.; Marinescu, I.; Castelló, J.V. Can fixed-term contracts put low skilled youth in a better career path? Evidence from Spain. *Econ. J.* **2019**, *129*, 1693–1730.
46. García Pérez, J.I.; Osuna, V. Dual labour markets and the tenure distribution: Reducing severance pay or introducing a single contract? *Labour Econ.* **2014**, *29*, 1–13.
47. García Pérez, J.I.; Rebollo Sanz Y. The use of permanent contracts across Spanish regions: Do regional wage subsidies work? *Investig. Econ.* **2009**, *33*, 97–130.
48. Garda, P. Essays on the Macroeconomics of Labor Markets. Ph.D. Dissertation, Universitat Pompeu Fabra, Barcelona, Spain, 2013.
49. Úbeda, M.; Cabasés, M.Á.; Sabaté, M.; Strecker, T. The Deterioration of the Spanish Youth Labour Market (1985–2015): An Interdisciplinary Case Study. *YoUnG* **2020**, doi: 10.1177/1103308820914838.
50. Vall Castelló, J. Promoting employment of disabled women in Spain: evaluating a policy. *Labour Econ.* **2012**, *19*, 82–91.
51. Vall Castelló, J. What happens to the employment of disabled individuals when all financial disincentives to work are abolished? *Health Econ.* **2017**, *26*, 158–174.
52. Alonso, F.; Devesa, J.E.; Devesa, M.; Domínguez, I.; Encinas, B.; Meneu, R.; Nagore, A. Towards an adequate and sustainable replacement rate in defined benefit pension systems: The case of Spain. *Int. Soc. Secur. Rev.* **2018**, *71*, 51–70.
53. Boado-Penas, M.C.; Valdés-Prieto, S.; Vidal-Meliá, C. An actuarial balance sheet for pay-as-you-go finance: Solvency indicators for Spain and Sweden. *Fisc. Stud.* **2008**, *29*, 89–134.
54. Conde Ruiz, J.I.; González, C.I. Reforma de pensiones 2011 en España. *Hacienda Pública Esp.* **2013**, *204*, 9–44.
55. Conde-Ruiz, J.I.; González, C.I. From Bismarck to Beveridge: The other pension reform in Spain. *SERIES* **2016**, *7*, 461–490.

56. Devesa, J.E.; Devesa, M.; Domínguez, I.; Encinas, B.; Meneu, R.; Nagore, A. Equidad y sostenibilidad como objetivos ante la reforma del sistema contributivo de pensiones de jubilación. *Hacienda Pública Esp.* **2012**, *201*, 9–38.
57. García García, M.; Nave Pineda, J.M. Impacto en las prestaciones de jubilación de la reforma del sistema público de pensiones español. *Hacienda Pública Esp.* **2018**, *224*, 113–137.
58. Moral Arce, I.; Patxot, C.; Souto, G. La sostenibilidad del sistema de pensiones. Una aproximación a partir de la CSWL. *Revista de Economía Aplicada* **2008**, *XVI*, 29–66.
59. Muñoz de Bustillo, R.; De Pedraza, P.; Antón, J.I.; Rivas, L.A. Working life and retirement pensions in Spain: The simulated impact of a parametric reform. *Int. Soc. Secur. Rev.* **2011**, *64*, 73–93.
60. Patxot, C.; Souto, G.; Villanueva, J. Fostering the contributory nature of the Spanish retirement pension system: An arithmetic micro-simulation exercise using the MCVL. *Presup. Gasto Público* **2009**, *57*, 7–32.
61. Peinado Martínez, P. Pension System's reform in Spain: A dynamic analysis of the effects on welfare. Ph.D. Dissertation. Universidad del País Vasco UPV/EHU, Bilbao, Spain, 2011.
62. Peinado Martínez, P. A dynamic gender analysis of Spain's pension reforms of 2011. *Fem. Econ.* **2014**, *20*, 163–190.
63. Peinado Martínez, P.; Serrano Pérez, F. A dynamic analysis of the effect of social security reform on Spanish widow pensioners. *Panoeconomicus* **2011**, *58*, 759–771.
64. Sánchez Martín, A.R.; Sánchez Marcos, V. Demographic change and pension reform in Spain: An assessment in a two-earner OLG model. *Fisc. Stud.* **2010**, *31*, 405–452.
65. Vidal-Meliá, C. An assessment of the 2011 Spanish pension reform using the Swedish system as a benchmark. *J. Pension Econ. Financ.* **2014**, *13*, 297–333.
66. Vidal-Meliá, C.; Boado Penas, M.C.; Settergren, O. Automatic balance mechanisms in pay-as-you-go pension systems. *Geneva Pap. Risk Insur. Issues Pract.* **2009**, *34*, 287–317.
67. Amuedo Dorantes, C.; Borra, C. On the differential impact of the recent economic downturn on work safety by nativity: the Spanish experience. *IZA J. Dev. Migr.* **2013**, *2*, 1–26.
68. Anghel, B.; Basso, H.; Bover, O.; Casado, J.M.; Hospido, L.; Izquierdo, M.; Kataryniuk, I.A.; Lacuesta, A.; Montero, J.M.; Vozmediano, E. Income, consumption and wealth inequality in Spain. *SERIEs* **2018**, *9*, 351–357.
69. Antón, J.I.; Muñoz, R. Public-private sector wage differentials in Spain. An updated picture in the midst of the great recession. *Investigación Económica* **2015**, *LXXIV*, 115–157.
70. Dudel, C.; López Gómez, M.A.; Benavides, F.G.; Myrskylä, M. The length of working life in Spain: Levels, recent trends, and the impact of the financial crisis. *Eur. J. Popul.* **2018**, *34*, 769–791.
71. Arranz, J.M.; García-Serrano, C. Are the MCVL tax data useful? Ideas for mining. *Hacienda Pública Esp.* **2011**, *199*, 151–186.
72. Arranz, J.M.; García-Serrano, C.; Hernanz, V. How do we pursue “labormetrics”? An application using the MCVL. *Estadística Española* **2013**, *55*, 231–254.
73. De la Roca, J.; Puga, D. Learning by working in big cities. *Rev. Econ. Stud.* **2017**, *84*, 106–142.
74. López, M.A.; Benavides, F.G.; Alonso, J.; Espallargues, M.; Durán, X.; Martínez, J.M. The value of using administrative data in public health research: the continuous working life sample. *Gac. Sanit.* **2014**, *28*, 334–337.
75. Pérez-Salamero González, J.M. La MCVL como fuente generadora de datos para el estudio del sistema de pensiones. Ph.D. Dissertation. Universidad de Valencia, Valencia, Spain, 2015.
76. Arranz, J.M.; García-Serrano, C. Duration and recurrence in unemployment benefits. *J. Labor Res.* **2014**, *35*, 271–295.
77. Arranz, J.M.; García-Serrano, C. Duration of joblessness and long-term unemployment: Is duration as long as official statistics say? In *Disadvantaged Workers. Empirical Evidence and Labour Policies*; Malo, M., Sciulli, D., Eds.; Springer: New York, NY, USA, 2014; pp. 297–320.
78. Arranz, J.M.; García-Serrano, C. The interplay of the unemployment compensation system, fixed-term contracts and rehiring: The case of Spain. *Int. J. Manpower* **2014**, *35*, 1236–1259.
79. Bentolila, S.; García-Pérez, J.I.; Jansen, M. Are the Spanish long-term unemployed employable? *SERIEs* **2017**, *8*, 1–41.
80. Nagore García, A. Gender differences in unemployment dynamics and initial wages over the business cycle. *J. Labor Res.* **2017**, *38*, 228–260.

81. Rebollo-Sanz, Y. Unemployment insurance and job turnover in Spain. *Labour Econ.* **2012**, *19*, 403–426.
82. Benavides F.G.; Durán, X.; Gimeno, D.; Vanroelen, C.; Martínez, J.L. Labour market trajectories and early retirement due to permanent disability: a study based on 14972 new cases in Spain. *Eur. J. Public Health.* **2015**, *25*, 673–677.
83. Carrillo-Castrillo, J.A.; Guadix, J.; Rubio-Romero, J.C.; Onieva, L. Estimation of the Relative Risks of Musculoskeletal Injuries in the Andalusian Manufacturing Sector. *Int. J. Ind. Ergonom.* **2016**, *52*, 69–77.
84. Castañer-Garriga, A.; Pérez-Salamero González, J.M.; Vidal-Meliá, C. Evaluación de las tarifas de las pensiones de accidentes de trabajo y enfermedades profesionales (2011–2015). *Rev. Innovar J.* **2017**, *27*, 153–167.
85. Durán, X.; Vanroelend, C.; Deboosere, P.; Benavides, F.G. Social security status and mortality in Belgian and Spanish male workers. *Gac. Sanit.* **2016**, *30*, 293–295.
86. Jiménez-Martín, S.; Juanmartí Mestres, A.; Vall Castelló, J. Hiring subsidies for people with a disability: Do They Work? *Eur. J. Health Econ.* **2019**, *20*, 669–689.
87. López Gómez, M.A.; Serra, L.; Delclos, G.L.; Benavides, F.G. Employment history indicators and mortality in a nested case-control study from the Spanish WORKing life social security (WORKss) cohort. *PLoS ONE* **2017**, *12*, E0178486.
88. López Gómez, M.A.; Durán, X.; Zaballa, E.; Sanchez-Niubo, A.; Delclos, G.L.; Benavides, G.L. Cohort profile: The Spanish WORKing life Social security (WORKss) cohort Study. *BMJ Open* **2016**, *6*, doi: 10.1136/bmjopen-2015-008555.
89. López, M.A.; Durán, X.; Alonso, J.; Martínez, J.M.; Espallargues, M.; Benavides, F.G. Estimating the burden of disease due to permanent disability in Spain during the period 2009–2012. *Rev. Esp. Salud Public.* **2014**, *88*, 349–358.
90. Bonhomme, S.; Hospido, L. Earnings inequality in Spain: new Evidence using tax data. *Appl. Econ.* **2013**, *45*, 4212–4225.
91. Bonhomme, S.; Hospido, L. The cycle of earnings inequality: evidence from Spanish social security data. *Econ. J.* **2017**, *127*, 1244–1278.
92. Marie, O.; Vall Castelló, J. Measuring the (income) effect of disability insurance generosity on labour market participation. *J. Public Econ.* **2012**, *96*, 198–210.
93. Cairó Blanco, I. An empirical analysis of retirement behaviour in Spain: Partial versus full retirement. *SERIES* **2010**, *1*, 325–356.
94. García-Gómez, P.; Jiménez-Martín, S.; Castelló, J.V. Health, disability, and pathways into retirement in Spain. In *Social Security Programs and Retirement around the World*; Wise, D.A., Ed.; University of Chicago Press: Chicago, 2012; pp. 127–174.
95. García Pérez, J.I.; Jiménez Martín, S.; Sánchez Martín, A.R. Retirement incentives, individual heterogeneity and labor transitions of employed and unemployed workers. *Labour Econ.* **2013**, *20*, 106–120.
96. Vegas Sánchez, R.; Argimón, I.; Botella, M.; González, C. Old age pensions and retirement in Spain. *SERIES* **2013**, *4*, 273–307.
97. Cebrián, I.; Moreno, G. Labour market intermittency and its effect on gender wage gap in Spain. *Rev. Interv. Econ.* **2013**, *47*, doi: 10.4000/interventionseconomiques.1950.
98. Cebrián, I.; Moreno, G. The effects of gender differences in career interruptions on the gender wage gap in Spain. *Fem. Econ.* **2015**, *21*, 1–27.
99. INSS: Informes estadísticos, 2005–2017. Madrid: Instituto Nacional de la Seguridad Social. Secretaría de Estado de la Seguridad Social. Ministerio de Trabajo, Migraciones y Seguridad Social. Available online: <http://www.mitramiss.gob.es/es/estadisticas/> (accessed on 12 March 2020).

