



Income Distribution, Household Heterogeneity And Consumption Insurance In The UK: A Mixture Model Approach

Gabriele Amorosi (University of Hull, Business School, UK)
Amanda Gosling (University of Kent, School of Economics)
Miguel Leon-Ledesma (University of Kent, School of Economics)

Paper prepared for the 34th IARIW General Conference

Dresden, Germany, August 21-27, 2016

Session 2G: Research Forum 1 on Income Distribution Issues

Time: Monday, August 22, 2016 [Afternoon]

Income distribution, household heterogeneity and consumption insurance in the UK. A mixture model approach.

Gabriele Amorosi[§], Amanda Gosling[†] and Miguel Leon-Ledesma[†]

[§]University of Hull, Business School; [†]University of Kent, School of Economics.

Abstract

This paper uses data from the FES to examine the changing joint distribution of income and consumption from 1968 to 1999. The analysis is given structure by the assumption that households are of three possible types. Within each type there is no *permanent* heterogeneity but some may have higher current income due to idiosyncratic shocks. This simple idea allows us to estimate the mean income of each group and its variance. This variance will be the variance of the temporary shock. With the further assumption that macro shocks are the same for each group of households, we can then predict the distribution of consumption and the relationship between income and consumption according to the life-cycle model (where current consumption is only a function of permanent income). These predictions are then compared against the data and against the predictions of a “no smoothing” extreme where households simply consume out of current income. Our results are as follows. Firstly, the income estimates illustrate quite starkly the idea of polarisation, the middle income groups are shown to be declining in mass and in relative income. Secondly, neither extreme model fits the data exactly, but the fit of the life-cycle model with 3 types of heterogeneity seems to improve over time; this gives further support to the idea that part of the increase in inequality is due to risk – not to permanent differences – and that capital markets appear to be improving households’ insurance mechanism against it.

1 Introduction

This paper looks at the evolution of the household distribution of income and consumption. In common with much other work, it looks at the questions of the degree to which households are able to insure themselves against risk, how this has changed over time and whether the increase in income inequality is a result of increases in permanent inequality or increases in risk. We diverge from the rest of the literature, however, as we examine cross-sectional data and *infer* the relative sizes of permanent and temporary shocks. We assume that households are of 3 possible types with different permanent means and vulnerability to temporary shocks. These differences can be estimated by a mixture model, where the estimation procedure fits the overall distribution of income by setting it as a function of 3 different distributions with unknown proportions.

The first part of the empirical analysis consists of the estimation of income densities through a model of mixed normal distributions; the model provides an estimate of three normal densities, which seems to fit the data substantially well and track the overall profile of the empirical distribution of income. These results illustrate that polarization (see for example Goos and Manning (2007) and Autor et al. (2006)) may have occurred both through an increase in gap between the middle classes and the rich and a decline in their relative proportion. The results also show that there has been both an increase in within and between variances. None of these findings are particularly new, of course, but can be used to validate our methodology.

Next, we compare the actual distribution of consumption, as described by the empirical kernel density estimate, with two hypothetical benchmark distributions appositely constructed by mixture of normal modelling: (i) the distribution of consumption that

would emerge in the case of no-insurance, which is the same as the distribution of income previously estimated, and (ii) the distribution of consumption in the case of (full) insurance, assuming that the variance of income is the same for the different sub-household groups and consumption is equal to permanent income. Then, we estimate the expected level of consumption for each group by means of conditional probabilities – given the income group j , $E(c|Pr(g_j) = 1)$ – and compare each group's expected consumption with the average group's mean income. In this way, we can compare how income and consumption trends evolve over time. Last we predict the expected level of consumption for each income level and compare it again to data and the “no smoothing” extreme.

The paper will proceed as follows: the next section provides an overview of some relevant literature; section three describes and analyses the data used; section four explains the methodology we have used in the modelling of income and consumption; section five illustrates the empirical analysis and its findings; section six concludes.

2 Background literature

2.1 Risk insurance and consumption smoothing

The basic idea of consumption insurance can be regarded as the cross-sectional counterpart of the permanent income hypothesis; if complete insurance is achieved, consumption of the individual units (be them agents, families or countries) should not vary in response to idiosyncratic income shocks. Common shocks, on the contrary, are supposed to affect all the underlying population in the same way, thereby leaving the overall distributional proportions unchanged.

One of the seminal papers addressing the full-insurance hypothesis is Mace (1991), which tests risk-sharing with complete markets and assesses how household consumption responds to changes in aggregate consumption and/or to other idiosyncratic shocks – e.g. individual income, employment status, etc. The empirical results seem to be highly consistent with the assumption of risk-sharing, even without complete markets. Cochrane (1991), which represents another landmark in the literature, provides a further cross-sectional consumption-insurance test, based on the fundamental proposition that, in the presence of a full-insurance mechanism, cross-sectional consumption changes should be independent of variables exogenous to households (idiosyncratic ones); his empirical analysis seems to fail to reject the underlying theory.

Amongst the empirical literature, an early example of the use of repeated-cross-section data for the study of the distributions of consumption and income is Cutler and Katz (1992).

More recently, a great bulk of the literature has focused on the study of the relationship between income and consumption by looking at the validity of the permanent income hypothesis (PIH) and at how permanent and transitory shocks to income affect the level of consumption; in other words, the focus here is on how consumption is smoothed over time and what type of shock is more relevant. Deaton and Paxons (1994) look at how the PHI can explain the patterns of increasing consumption inequality over time within age cohorts in Great Britain, Taiwan and the US. Blundell and Preston (1998) are an interesting example of how the literature has studied the distribution of income and consumption; the emphasis here is being placed on the distinction between the concepts of permanent and transitory income uncertainty in the assessment of the evolution of consumption. Using data from the Family Expenditure Survey (FES), they study how

permanent and transitory income uncertainty affects the change in consumption inequality within British cohorts. Another notable study about the links between the types of (risk) insurance and the way income shocks translate into consumption inequality is Blundell, Pistaferri and Preston (2008); building on the work of Blundell and Preston (1998), the authors set up a framework whereby the change in consumption can be specified by a process where both permanent and transitory shocks can play a role; they find evidence that the persistence of shocks is important for the determination of consumption inequality and that, if temporary shocks are fully insurable, permanent shocks are insurable only to a partial extent.¹ This study has also spurred the work by other scholars, like Kaplan and Violante (2010), who compare the empirical estimates of the insurance coefficients obtained by Blundell and Preston (1998) with the standard incomplete market (SIM) model, in order to assess the degree of self-insurance as explained by the available data. Japelli and Pistaferri (2010a), in a recent work based on the Bank of Italy's survey on consumption, income and wealth, carry out a highly informative and extensive statistical analysis of different consumption and income inequality measures about Italy, both at the individual and at the household level. Assessing the link between consumption and income trends, they find that income inequality is usually higher than consumption inequality and tends to grow at a faster pace – this is also reflected at different cohort levels. Since this income variation is mainly attributable to non-permanent income shocks, they argue that this is an evidence of the PIH and speculate an important role could have been played by deeper financial markets.

2.2 Financial market deepening and risk-sharing.

¹ The insurance to permanent shocks applies particularly to people near retirement age.

The full consumption insurance hypothesis is based on very strong assumptions, as it can only exist in a competitive equilibrium environment where financial markets are complete or where are institutions implementing optimal an allocation of resources. Indeed, as consumption insurance is based on the efficient allocation of resources from savers to users, the degree of efficiency of financial markets is a key factor for the implementation of the full risk-insurance mechanism. This explains also why models of incomplete markets, where agents are affected by credit constraints and the insurance mechanism works only partially, have been set up and used in the literature. As a consequence, it is widely accepted that the degree of financial market efficiency is very important factor.

Ventura (2008) provides an empirical analysis of risk-sharing opportunities in Latin America and the Caribbean; his preliminary results show that the financial development of an economy is crucial in determining the amount of risk sharing.

The positive effects of financial integration (through the enhancement of the potential for risk sharing) on an aggregate consumption level have been extensively advocated and documented in the international finance literature. Demyanyk and Volosovych (2008) report substantial welfare gains – both real and potential – when comparing permanent consumption amongst EU's 25 member-countries.² Fratzscher and Imbs (2009) introduce asset choice in a model aimed at the testing of international consumption risk sharing. They find evidence of the role played by financial constraints in hampering down the sharing of risk when foreign direct investment and bank loans

² Potential welfare gains here refer to those gains achievable by further increasing the degree of risk sharing within the EU.

are used, with idiosyncratic consumption growth becoming more subject to idiosyncratic income changes.³

A more analytical assessment of the effects of financial market integration on the joint evolution of income and consumption – indeed on consumption smoothing – is Japelli and Pistaferri (2011); the two authors, by means of the same data used in Japelli and Pistaferri (2010), assess the determinants of the divergence between the trends in the variance of income and in the variance of consumption in Italy following the implementation of the European Monetary Union (EMU). They find that the sharp increase in income inequality during the period considered has not been mirrored by a similar increase in consumption inequality. However, they conclude that the EMU has not changed the ability of consumers to smooth income shocks and that the diverging trends between income and consumption inequality are primarily due to the increase in transitory inequality.

3 The data

All the income and consumption data used in our analysis come from the UK Family Expenditures Surveys (FES). The FES was carried out by the UK's Office for National Statistics (ONS) and has been in use since 1957 and, probably representing the main source of information for official, scholarly and private statistics about British households' income. It was replaced in 2001, when the FES and the National Food Survey (NFS) have been combined into a new survey, the Expenditure and Food Survey

³ As the main sources of transaction costs in international investment, they mention the following factors: different tax treatments, intermediation fees, liquidity premia across countries or asset classes and, also, informational frictions.

(EFS).⁴ The FES consists of approximately 6,500 households each year, but the dataset is not a longitudinal one, as the number and consistency of household changes each time.

The type of income concept considered in the FES – which I also use in this paper – is ‘*gross normal weekly household income*’, which represents the sum of weekly incomes from employment, self-employment investments, occupational pensions, imputed income for owner-occupiers and most social security benefits across household members.⁵

The consumption variable relevant to our research refers to total expenditures (plus imputed values), including both durables and non-durables expenditures. The inclusion of durables in consumption data has some important implications for the resulting data analysis. As explained by Bar-Ilan and Blinder (1988), the expenditures for durable goods are essentially more volatile than the expenditures for nondurable goods and services, with the average stock of durables (and total expenditures) not being necessarily proportional to permanent income and following the growth rate implied by the PIH. On the contrary, changes in permanent income are likely to give rise to large changes in durable expenditures. It is, therefore, important to keep this aspect in mind when trying to interpret the relevant data. As also illustrated by Bertola, Guiso and Pistaferri (2005), durables expenditures fluctuate substantially around their optimal amount – consistent with a proportion of nondurables – however keeping a relatively stable trend. This implies that the variability of consumption is likely to be exacerbated by the inclusion of durable amongst the data described.

⁴ Further to income and consumption, the FES collects a large amount of information on socio-economic characteristics of the households, like composition, size, social class, occupation and age of the head of household. Data is collected throughout the year to cover seasonal variations in expenditures (See the ONS website for relevant publications and information material: <http://www.statistics.gov.uk>).

⁵ See: FES User Guide (different issues).

In order to make the analysis comparable across households, we equivalize the values of income and consumption by means of the equivalence scale McClements developed for the econometric analysis of the 1971 and 1972 waves of FES.⁶

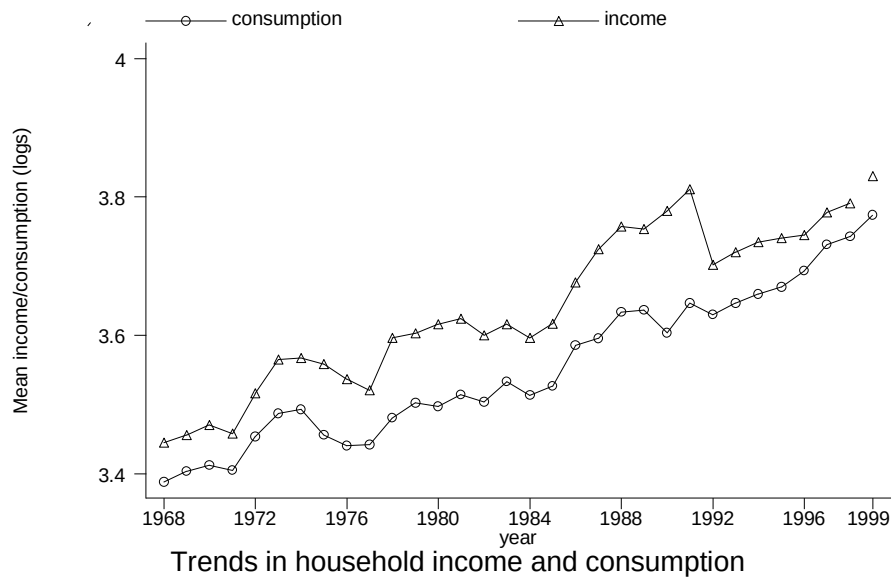
Another important step in the handling of the survey data, in order to work with real values of consumption and income, is to deflate both sets of equivalised data on a monthly basis; this is a necessary measure in order to work with time-consistent income and consumption values, due to the relatively high inflationary process affecting the UK economy during the 1970s and, to some extent, the early 1980s. Therefore, using the CPI index provided by the National statistical office we have constructed a series of monthly deflators that we have applied to the equivalized monthly consumption and monthly income.

The size of our yearly samples is the largest possible and we make no distinction between employed, self-employed and unemployed.

A common (but surely effective) way to assess the relationship between household income and household consumption is by looking at how the trends in the mean values of (the logs of) the two variables evolve relative to each other. Graph 1, which portrays the joint evolution of the two underlying variables in the period 1968-1999, shows how both income and consumption follow a fairly similar path; after an initial upwards trend, they both fall in the mid 70s and start to rise again since the late 70s, further accelerating during the first half of the 80s. A sudden fall occurs around 1991 – probably due to the ongoing recession affecting the UK – followed by a somewhat steady rise during the remaining period considered.

⁶ In an attempt to improve the equivalence scales in use at the time, McClements developed a scale that takes into account the effects of the number of children and the ages of the children on the living standards of the household. Unlike the OECD scales, the McClements scale equivalises the household income to the reference unit of an adult couple. See: Goodman, Johnson and Webb (1997).

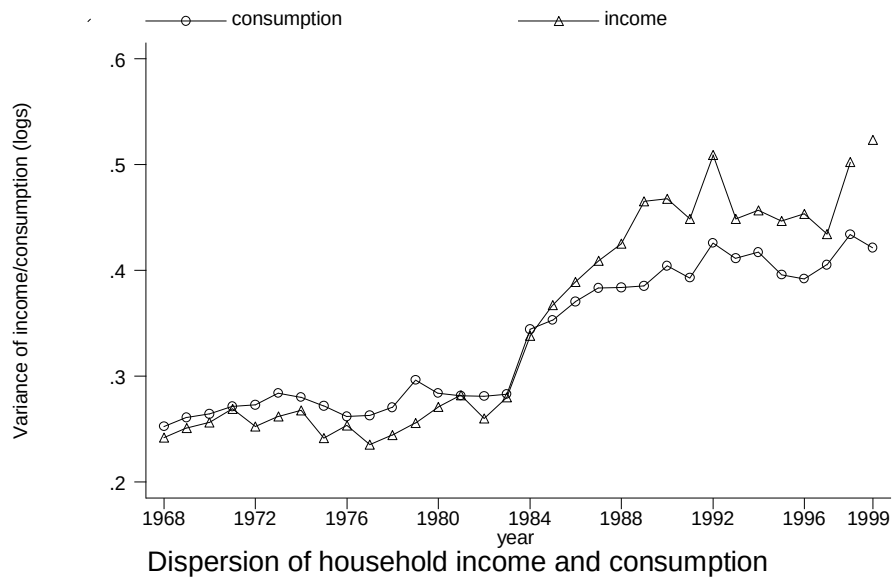
Graph 1



Looking at the relative trend, income seems to be growing at a slightly quicker pace than consumption during the 70s and the 80s. However, if one considers the recession that has affected the UK in the beginning of the 90s, it is possible to notice that the situation seems to be reversed; the decline in consumption looks far less pronounced than the decline in income. The opposite is true for the post-1992 recovery, as in this case income seems to pick up more slowly than consumption.

Like their mean values, the variances of household income and household consumption are also exhibiting a similar pattern; nevertheless, in this case, as pictured in graph 2, it is noticeable how the magnitude of the difference between the two variables remains somehow small during the first half of whole period examined – this is particularly true for the years before 1984, when the difference between the two variances tends to become almost negligible. After 1984, though, there seems to be some signs of a widening of the underlying gap between income and consumption that begins to widen and fluctuates in terms of magnitude.

Graph 2



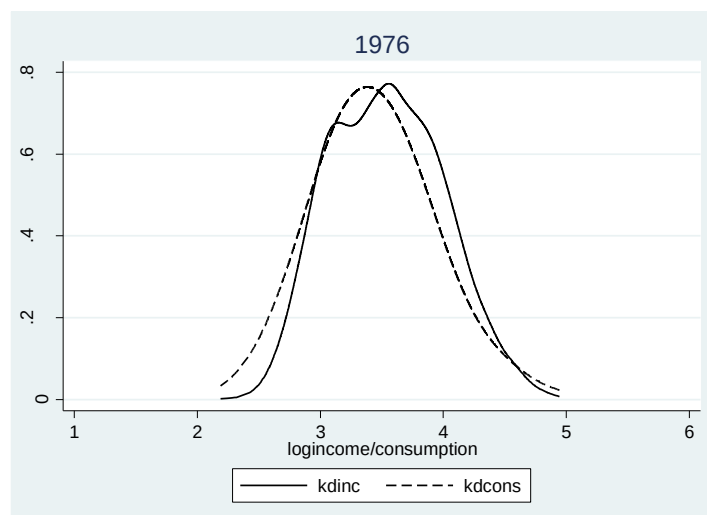
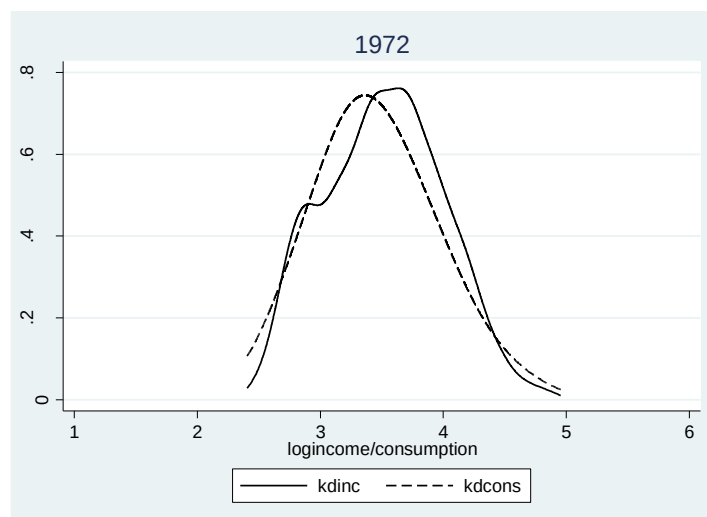
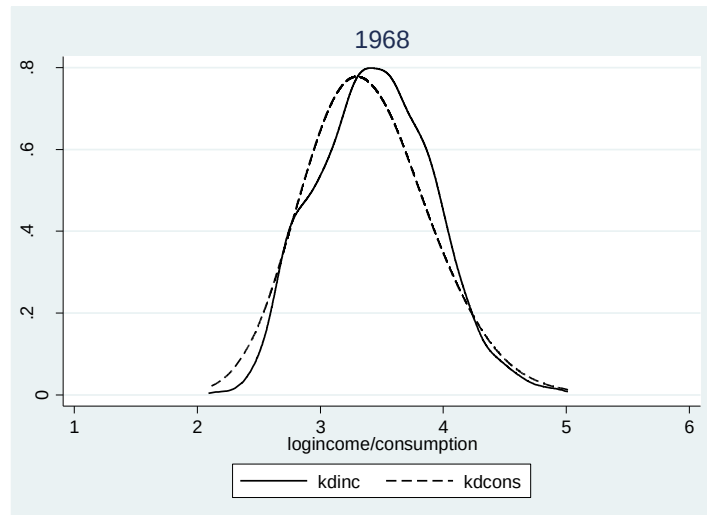
As far as the trend is concerned, the variance of consumption remains greater than the variance of income from 1968 to 1984; during this period the gap has remained relatively narrow, then to close up in 1984. In the subsequent period, however, income variability increases more than consumption for a few years, after which both values remain substantially stable until the late 90s. This is a sign that factors such as the labour market reforms, the increase in the skill-premium, and the financial liberalization process boosted in the mid-1980s may have played an effective role in shaping the relative distribution of income and consumption.

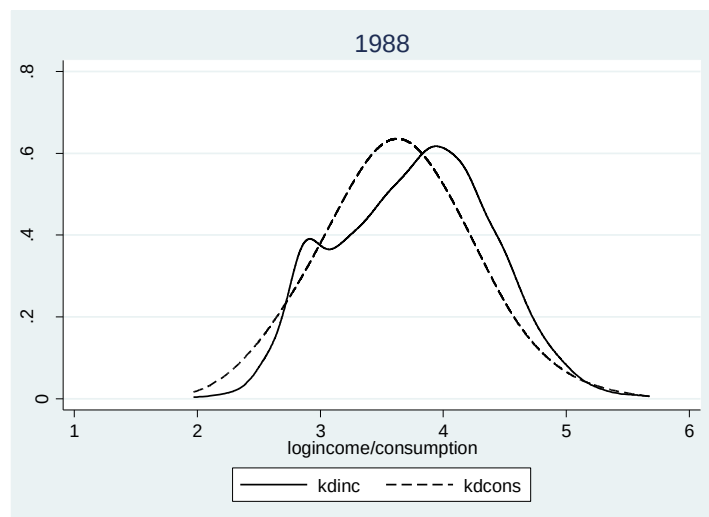
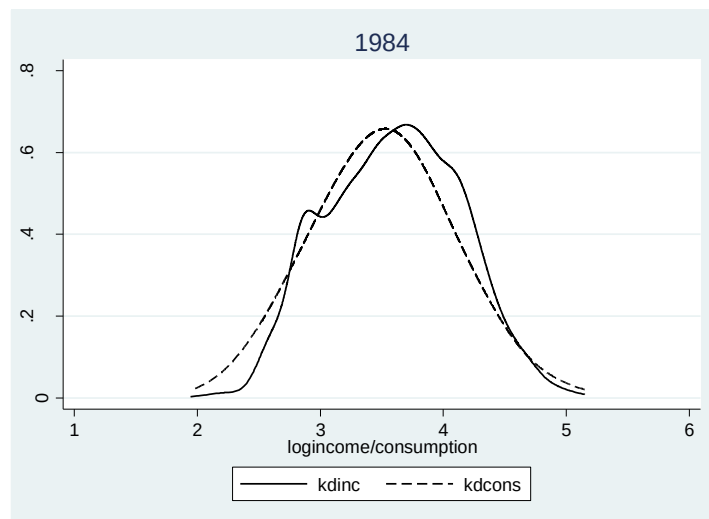
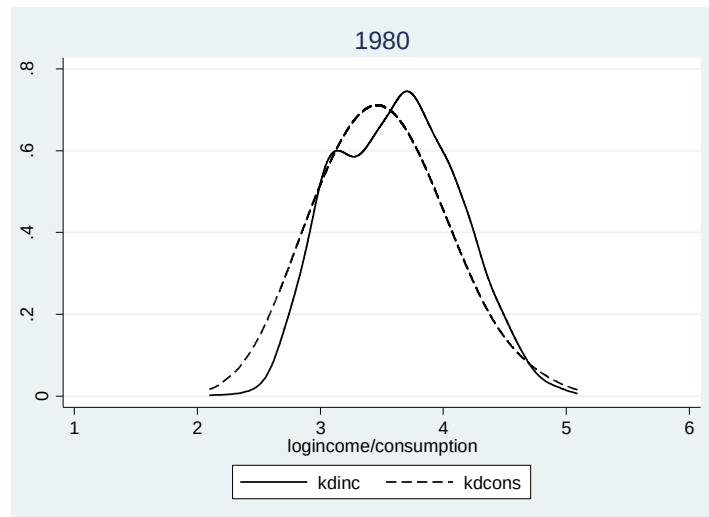
A more informative overview of the joint evolution of the distributions of income and consumption can be obtained by graph 3, which shows the empirical kernel estimates of the densities of income and consumption in selected years. The underlying data evidence two major features. First, the density of income is clearly developing a multimodal profile over time, while the same does not seem to apply to consumption, whose density profile seems to be rather unimodal, with a tendency towards a light polarization in the very end of our sample period; this discrepancy can be interpreted as

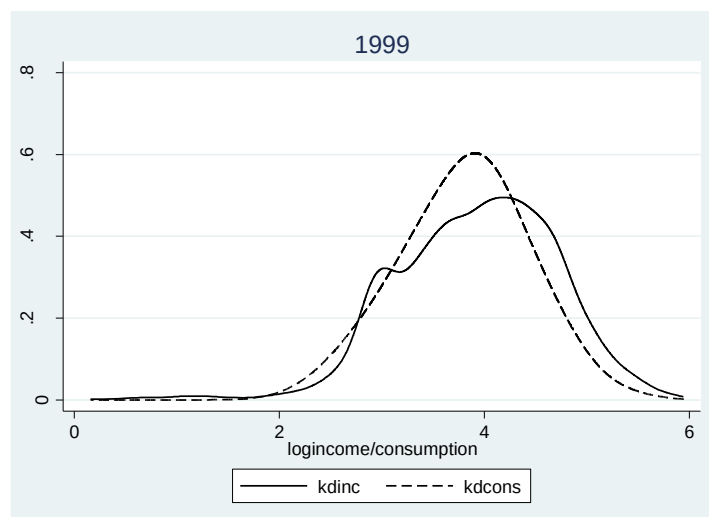
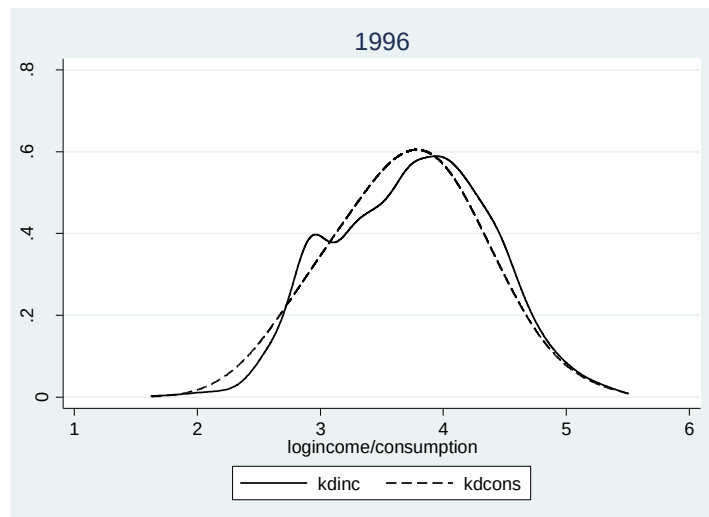
a consequence of the presence of some forms of insurance mechanism that allow the smoothing of consumption across different income groups of households. Second, the spread of both distributions proves to be increasing over time, with the density of income looking wider than the density of consumption; this is another sign of the likely presence of consumption insurance at large.

Graph 3

Empirical density of income and consumption







4 Methodology

If one looks at the trends (over time) in the mean or in the variance of the income of all the households that form our samples and compare them with graph 3, depicting the kernel density estimation (as illustrated in section 4), it is possible to realize how uninformative the formers are, in terms of distributional features. In fact, they simply describe the overall trend of a population that is regarded as a homogeneous group (the same reasoning can, by analogy, be applied to consumption). Yet, the mere use of the kernel density estimate, as depicted in the underlying diagrams, is certainly an improvement in the description of the main attributes across the whole distribution of income (and consumption), but does not provide a well defining characterization/quantification of the degree and scope of compositional diversity. This calls for the use of a methodology that is better equipped to address the heterogeneity issue when analyzing the distribution of income (and consumption).

The first step is to fit the data by a combination of normal distributions; this allows us to model the existing heterogeneity in terms of a given set of sub-population groups and obtain estimates of the parameters that give information of each normal distribution's income process therewith involved. The second step is the “construction” a counterfactual distribution of consumption that reflect the hypothesis of perfect smoothing; in order to do so, we constrain the distribution of consumption to have the same mean as the mean of permanent income and a constant common variance, with the number of means and variances depending on the number of sub-distributions estimated by the mixture model.⁷

⁷ Here the mean of permanent income is proxied by the mean of current income, under the assumption that the aggregate idiosyncratic component of income is zero (as will be explained later in this section).

Then, we estimate the predicted consumption level for each group of households, conditional of the actual level of income, and compare the trends in relative expected consumption levels with the average income of each group.

4.1 Mixture models (an overview)

Mixture models (also known as “finite” mixture models, due to the fact that the number of mixed distributions to be usually estimated in the underlying model is finite) are becoming an important tool for the analysis of economic data, as they are a very useful and flexible analytical method for the interpretation of complex distributions; in particular, mixture models can be relatively suitable to the description of multimodal, skewed, fat-tailed and other types of density curves.

Lindsay (1995), in describing the mixture model, states: “*The simplest and most natural derivation of the mixture model arises when one sample forms a population that consists of several homogeneous subpopulations,.....*”. McLachlan and Peel (2000) underline the usefulness of mixture distributions in modelling heterogeneity in different statistical contexts, as they provide a more direct role in the analysis of data and better model description for many distributions.

Mixture models are considered semi-parametric for two main reasons: firstly, even though each component of the mixture is a parametric model, the distribution of mixing proportions is “model-free”; secondly, there is no restriction that the components of the mixture belong to the same family of distributions.

Equation (1) below is a general specification of a mixture model:

$$(1) \quad f(x) = \sum_{j=1}^m \pi_j g(x; \theta_j) ,$$

where $f(x)$ is the mixture density function, $g(x; \theta_j)$ are the density functions of each i th component's distribution, with θ_j being the parameters vector of each j th component's distribution (e.g. the mean and the standard deviation), and π_j is the mixing proportion relative to the j th group's distribution – the summation $\sum \pi_j$ is obviously equal to one.

Amongst the different methods that are used to estimate the parameters of finite mixture models the maximum likelihood method is probably the most convenient one, as the estimates therewith obtained tend to converge to the true parameters values even under very general conditions. Specifically, if one considers a variable of sample of size n , $x_1, x_2, \dots, x_n$, consisting of m mixture components, the maximum likelihood method searches for the parameters values that maximize the likelihood function evaluated at the observations.

The generalized maximum likelihood function of a mixture can be specified as follows,

$$(2) \quad L(\theta, \pi) = \prod_{i=1}^n f(x_i) = \prod_{i=1}^n \left\{ \sum_{j=1}^m \pi_j g(x_i; \theta_j) \right\},$$

where θ is the vector of parameters of the density functions and π is the vector of mixing proportions as explained above.

The advantage of using such an approach is absolutely evident, as it makes it possible to overcome both the typical problems of pure parametric estimates, which do not capture all the distribution-wide features of the relevant population.

As illustrated by Celeux (2007), mixture models are versatile and parsimonious models with many available algorithms to estimate the mixture parameters, can be address special questions in the model-based clustering context in a proper way and can be compared and assessed in an objective way.

Another advantage of mixture models is that they combine much of the flexibility of nonparametric methods with the analytical advantages of parametric estimation.

As far as the application of mixture models to the study of economic phenomena is concerned, the literature is not very abundant of them and in most of the cases they refer to relatively recent exercises. Amongst these, there are some examples of the application of this methodology to the analysis of income distribution.

Pittau and Zelli (2005) uses a mixture model of normal distributions in order to test the presence of polarization vs. convergence in EU regions in the 1980s and 1990s and describe the evolution of the shape of the distribution of regional income; their main evidence shows a certain degree of convergence is the reduction of the number of components in fitting the empirical distribution.

In Flachaire and Nuñez (2007) one can find an analysis of conditional household income distributions using lognormal mixtures. Using FES data, they propose a conditional model by specifying the mixing probabilities as a particular set of functions of individual characteristics (e.g. working status). This allows the characterization of distinct homogeneous subpopulations: if one assumes that an individual's belonging to a specific subpopulation can be explained by his individual characteristics, then the probability of belonging to a given subpopulation may be varying among different individuals.

Bartošová and Forbelská (2010) apply the mixture model approach to the modelling of the distribution of household income in the Czech Republic, during the period 2005-2007, and asses its evolution using EU-SILC data.

4.2 Income modeling

Mixture models can potentially be applied to any type of mixed distribution and their choice should underlie the characteristics of the relevant population. However, a convenient way to model income is by using a mixture of normal distributions. Not only this choice reflects the great flexibility of normal distributions, whose combinations, as Marron and Wand (1992) point out, usually produces very close approximation of any type of density; it is also consistent with the consolidated notion that income is bound to follow a distribution that is approximated very closely by the lognormal distribution.

In analytical terms, in order to compute the ML estimation of the mixture of normals model, we need to consider we a model specification where we substitute the formula of the standard normal distribution into (2), thus we obtain the type of likelihood function to be used in our model that is illustrated below in equation (3);

$$(3) \quad L(\mu, \sigma^2, \pi) = \prod_{i=1}^n \left\{ \sum_{j=1}^m \pi_j \left[\frac{1}{\sqrt{2\sigma_j}} \times \exp \left(-\frac{1}{2} \left(\frac{(x_j - \mu_j)^2}{\sigma_j^2} \right) \right) \right] \right\},$$

where μ is the vector of subgroup j 's means and σ^2 is the vector of subgroup j 's variances.

The Maximum Likelihood (ML) estimate of the parameters of a mixture model is commonly computed by implementing the “expectation-maximization” (EM) algorithm, which is an efficient iterative procedure in the presence of missing or hidden data (in the

case of mixture model the missing values are the Z_i s that indicate the group membership for the i -th observation: Z_{ij} s indicates that the observation x_i belongs to the j -th group).⁸

A non-trivial problem we have encountered while implementing the estimation procedure is the large sampling variance, which can be reflected in a log-likelihood function relatively flat, at least near the maximum, or in the absence of a global maximum; this implies that different values of distribution parameters to be estimated are nearly equally likely to be fit the data. Therefore, we need to run the ML estimation a large number of times in order to obtain the highest ML value out of many.

In order to fit the model to the data, we have run the ML procedure with all the necessary iterations, in order to define the parameters and the proportions of the normal distributions derived from the mixture model we have implemented; it turns out that the number of groups/proportions composing the distribution of household income is equal to three. Therefore, our combination of normal distributions can be specified as follows:

$$(4) \quad f^{mix}(y_i) = \pi_{1,i}g(y_i | \theta_1) + \pi_{2,i}g(y_i | \theta_2) + \pi_{3,i}g(y_i | \theta_3).$$

As far as the estimation procedure is concerned, it is necessary to impose that the parameters θ_j estimated for the first year are identified with the same j -th group of households also in the following years.

Moreover, one major issue following the first estimation of the mixture model is represented by the large variability of the group proportions π_i , from one year to

⁸ Each iteration of the EM algorithm consists of two processes: The “expectation step” (E-step), and the “maximization step” (M-step). In the E-step, the missing data are estimated given the observed data and current estimate of the model parameters. This is achieved using the conditional expectation, explaining the choice of terminology. In the M-step, the likelihood function is maximized under the assumption that the missing data are known. The estimate of the missing data from the E-step is used in place of the actual missing data. Convergence is assured because the algorithm is guaranteed to increase the likelihood at each iteration. For more detailed information see: Dempster, Laird and Rubin (1977); and Aitkin and Rubin (1985).

another, which can be extremely large at times (especially until the mid-1970s).⁹ Therefore, in order to overcome this problem and obtain a more reasonable profile of the π_j parameters over time, we smooth the estimated values by using a five years moving average process MA(5). In this way, our proportions become:

$$(5) \quad \tilde{\pi}_{j,t} = \frac{1}{5} \sum_{m=t}^{t+4} \pi_{j,m},$$

where $j = 1, 2$ and 3 . This allows a better identification of the model.

We can now re-estimate the parameters of the mixed income densities, by constraining the maximum likelihood estimation of our mixture of normals to fit the proportion parameters $\tilde{\pi}_j$ as constructed in (5) for each year; thus we obtain a “smoothed” mixture model that becomes:

$$(6) \quad \tilde{f}^{mix}(y_i) = \tilde{\pi}_{1,t} g(y_i | \tilde{\theta}_1) + \tilde{\pi}_{2,t} g(y_i | \tilde{\theta}_2) + \tilde{\pi}_{3,t} g(y_i | \tilde{\theta}_3),$$

where $\tilde{\theta}$ denotes the new vector of income distribution parameters to be derived by means of ML estimation.

The methodology we have followed to compute the density of both the smoothed and unsmoothed mixture models is a typical the empirical kernel density estimation procedure, as outlined in equation (7) below:

$$(7) \quad \hat{f}(y_i) = \frac{\sum_i K\{(y_i - y)/n\}}{nh},$$

⁹ This is obviously due to noisiness of the survey data, which most probably suffer from a large sampling variation.

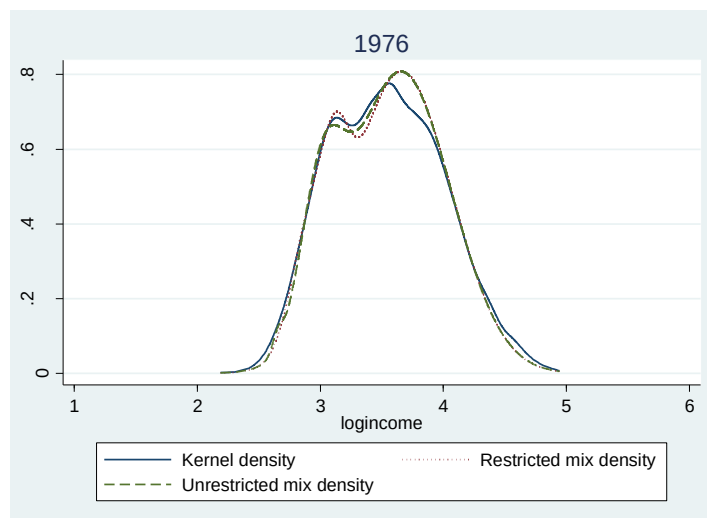
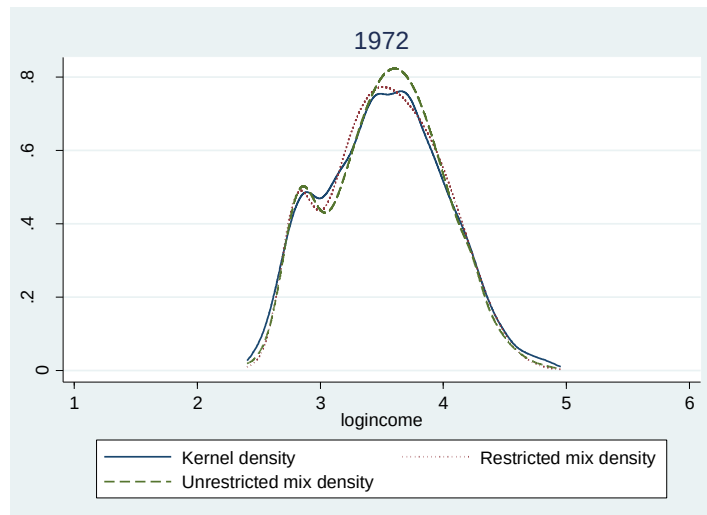
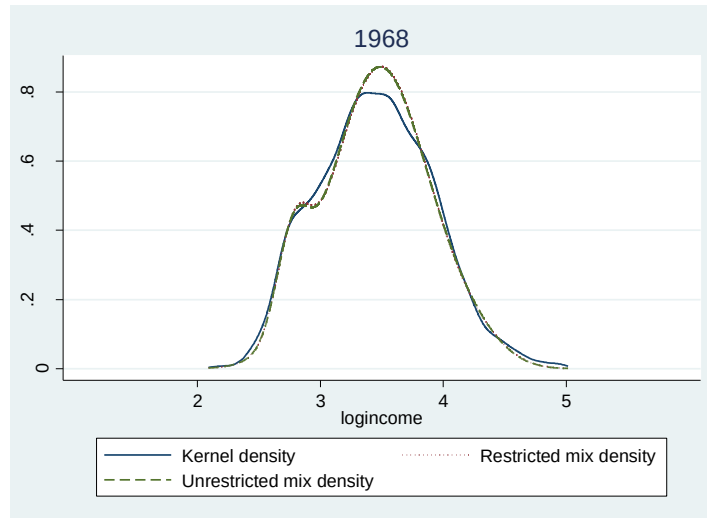
where K is the kernel function h is the smoothing parameter or band-width, it is possible to assess the suitability of the mixture models themselves in describing the actual data.

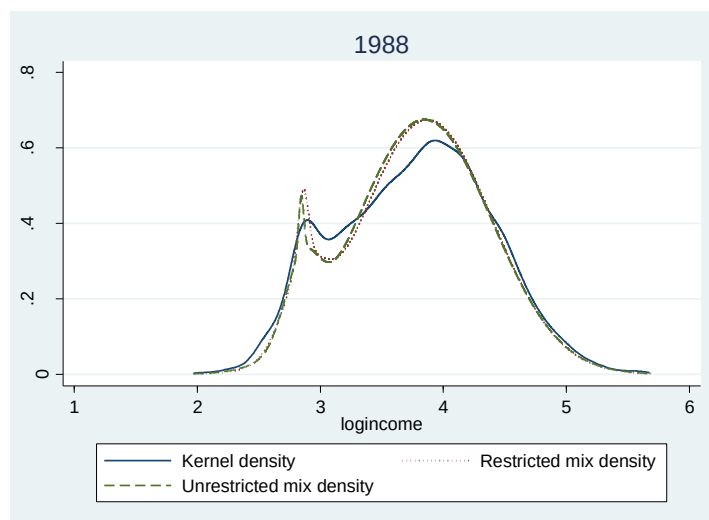
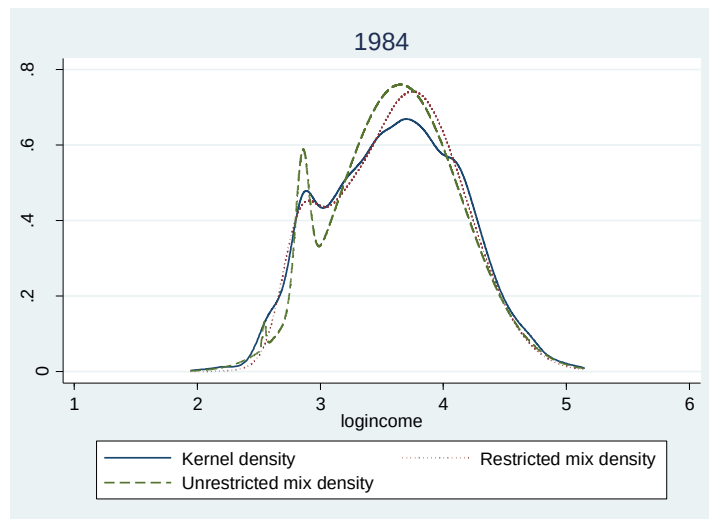
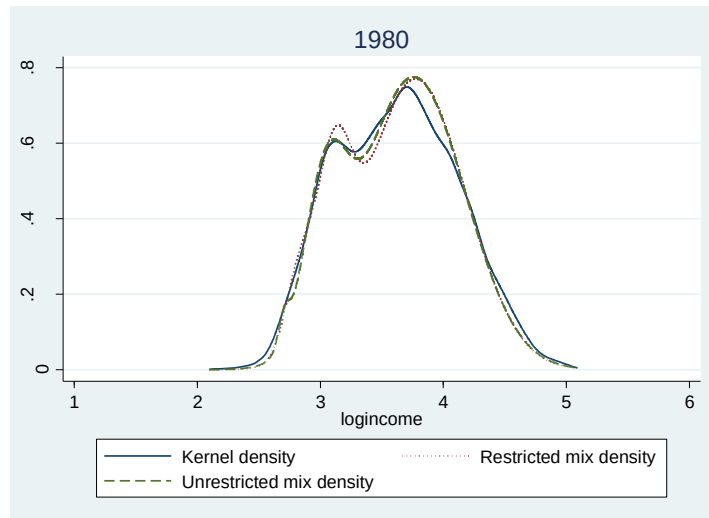
As is well known, an important step in the construction of the density profile is the choice of band-width h . In this case, we have used a rule of thumb and set $h = \text{s.d.}/7$ (this has proven to be an efficient choice for our estimation procedure).

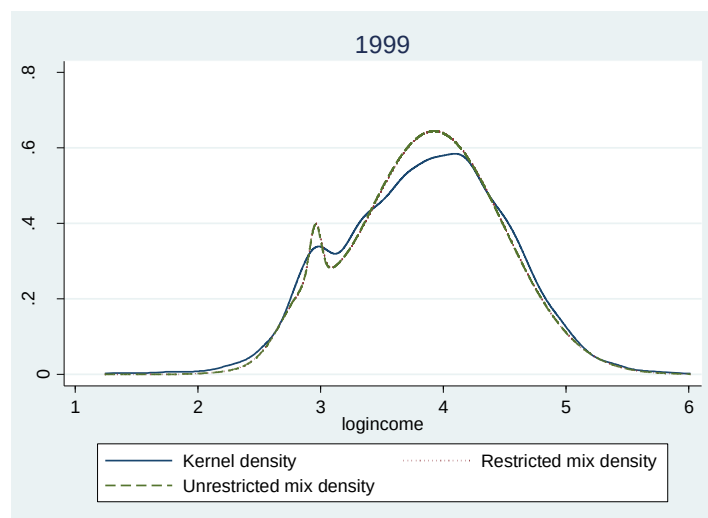
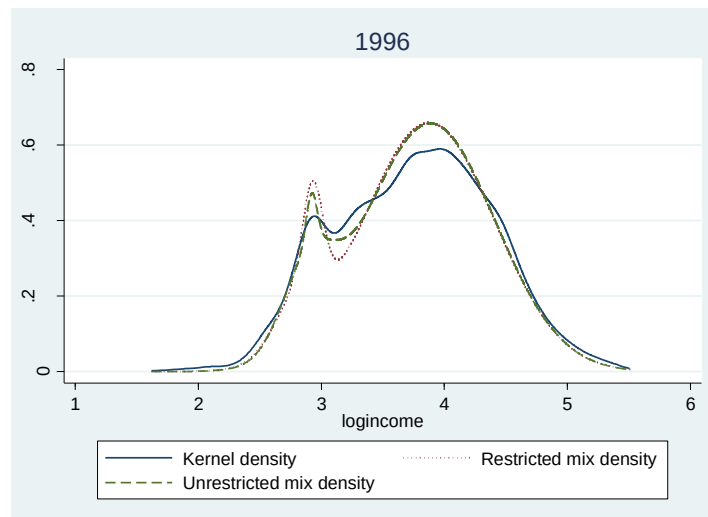
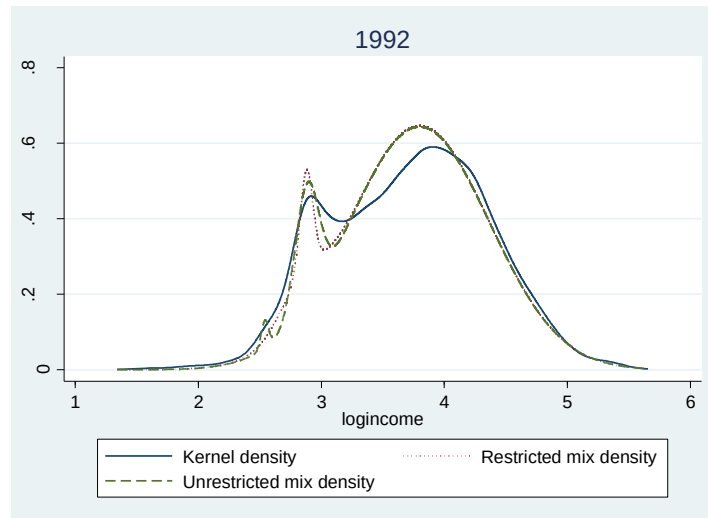
As shown in graph 4, the mixture of normals model we have used seems to track the main features of the actual income – e.g. the different modes – in a substantially close way, while the restricted mixture model differs from the unrestricted (with unsmoothed proportions) in a negligible way. Overall, the mixture of normal distribution seems to be a considerably suitable statistical procedure to model the distribution of UK's household income as reported in the FES survey. Indeed, the mixed normal distributions reflect the different modes detected by the empirical one. Moreover the discrepancy between the restricted (with smoothed proportions) and unrestricted (original proportions) density is minimal.

Graph 4

Income density-Actual and mixture model



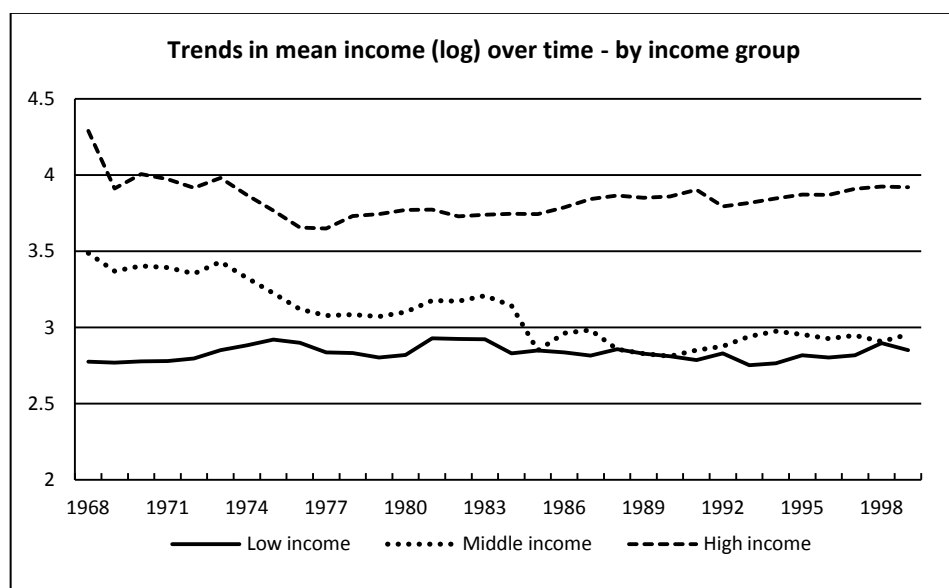




Some further analysis on household income is provided by graph 5 below, which shows the evolution of the mean income of the three income groups over time, once the mixture model has been restricted to have the smoothed proportions computed by the MA(5) in equation (6). If the means of the mid-income and the top-income groups have similar trends, especially until the mid-80s and after the early 90s (with a certain degree of divergence in between), the mean income level of the bottom-income group looks substantially stable over time.

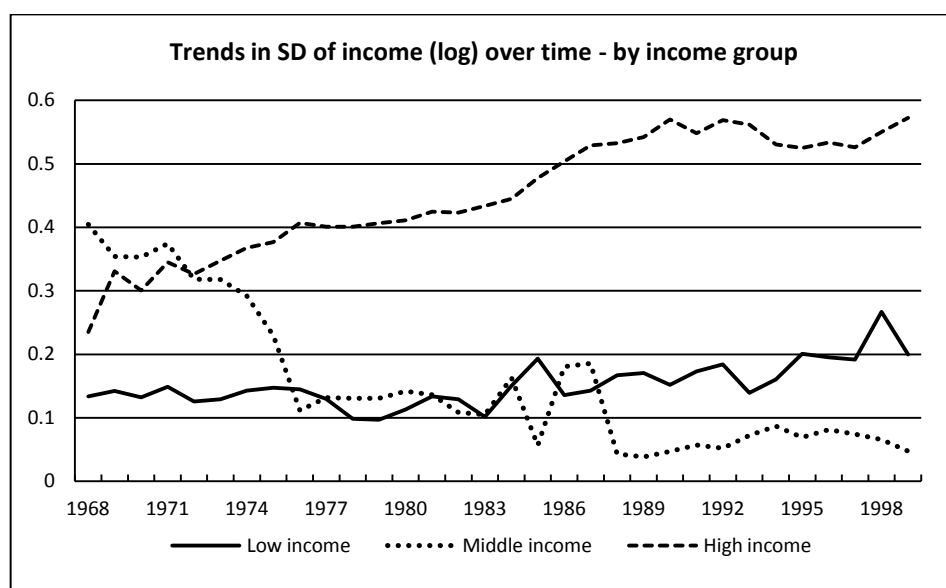
According to what is therein depicted, it is possible to distinguish between two distinct stages: (i) cross-group income inequality decreases until the late 1970s, principally due to the decrease in the income level of the richer and the middle groups; (ii) after that point in time, cross-group inequality starts to rise, mainly because of the widening of the gap between mid-income and top income households. At the same time, the mean incomes of the mid and the bottom groups converge to very similar values, something also in line with the estimation of the proportions of the two underlying groups by our mixture model.

Graph 5



As far as the trend in within group variation is concerned, as highlighted in graph 6, one can observe a significant and constant rise in the standard deviation of the higher income group, which after the initial period represents the vast majority of the households. This is mirrored by the sharp fall in the level of inequality within the so called middle-income group. The inequality of the low-income households remains substantially constant during the first half of the entire time period, and then it increases in the second half, but only to a moderate extent.

Graph 6



4.3 Consumption smoothing

The purpose of the second part of our empirical analysis is to assess the degree of consumption smoothing taking place across households across the relevant time period and is based on a simple assumption that can be summarized as follows; (i) if shocks to income are fully reflected in shocks to consumption, then there is no insurance, as consumption is proportional to income (meaning that agents/households will be credit

constrained and unable to insure in financial markets) (ii) if shocks to income and consumption are uncorrelated, this means that consumption depends only on permanent income and on consumption-specific shocks, but not on idiosyncratic shocks to income or group-specific shocks (meaning that agents/households will have full access to “insurance” markets).

After having defined the income process in the previous section, we need to define the model underlying consumption behavior. A common way to model consumption is by using the PIH framework. It is, therefore, necessary to first recall the income process as characterized by the PIH, which is as follows:

$$(8) \quad y_{it} = y_{it}^P + v_{it},$$

where y_{it}^P is the permanent component that follows a martingale process as specified below;

$$(9) \quad y_{it}^P = y_{it-1}^P + u_{it},$$

while v_{it} is an independently and identically distributed (i.i.d.) transitory income component.¹⁰ Like in Pischke (1995), we regard the income process set out in (9) above as a random walk, where u_{it} is the innovation component common to all households and v_{it} is an idiosyncratic shock that is uncorrelated across households; in this context, we can also assume that the underlying households fully adapt consumption to (permanent) aggregate shocks and adjust to household-specific income shock v_{it} by a factor α , as specified below;

¹⁰ See, among the others, Japelli and Pistaferri (2010a).

$$(10) \quad c_{it} = y_{it}^P + \alpha v_{it},$$

thus, α is a factor that determines the extent of risk-insurance.¹¹ From the perspective of our research, we have full insurance when α is equal to 1 and no insurance when α is equal to 0.

As the main purpose of this paper is to assess the degree risk insurance achieved by consumers, we need to compare the density of actual consumption with the densities of consumption relative to the cases of full smoothing and no smoothing. In this regard, the estimation of the density of actual consumption is straightforward, as we can derive it from the survey data, while the density of consumption with no-smoothing is the equivalent of the density of income as estimated by the mixture model specified by equation (6). Thus it is now necessary to estimate a density that reflects the case of perfect smoothing – indeed, assuming that a partial or full insurance mechanism is in place.

In terms of empirical analysis, the income that we observe can be specified as:

$$(11) \quad y_{it} = y_{it}^P + v_{it},$$

where y_i^P is consistent with the (permanent) aggregate income component and v_i is the idiosyncratic shock component.

For this purpose, we define consumption as:

¹¹ In Pischke α is equal to an annuity of value: $\frac{r}{1+r}$.

$$(12) \quad c_{it} = y_{it}^P + \varepsilon_{it},$$

where ε_i can be considered a generic idiosyncratic (household-specific) shock, substantially equivalent to the component αv in equation (10).

This means that, under the hypothesis of no insurance and/or full consumption smoothing, we have that:

$$\text{cor}(\nu_i, \varepsilon_i) = 0;$$

This is the situation that we need to simulate, by modelling consumption according to this assumption.

Under the opposite assumption of full insurance conditions, implying no consumption smoothing, it follows that:

$$\text{cor}(\nu_i, \varepsilon_i) = 1,$$

which is the case of consumption distribution fully replicating income distribution.

Using the same analytical framework of the mixture model of normal distributions we have used to fit income data (with its three income groups), we construct a model of consumption that reproduce an analogous mixture of normal distributions under the hypothesis that consumption reflects equation (12).

Given the mixture of normals modeling framework, it is necessary to assume that consumption is normally distributed. This implies that the mixture of normals setting underlying full consumption smoothing can be specified as follows:

$$(13) \quad \tilde{f}(c) = \sum_{j=1}^m \frac{1}{\bar{\sigma}^C} \phi_m \left(\frac{c_i - \bar{y}_j^P}{\bar{\sigma}^C} \right),$$

where $\bar{y}_j^P$ and $\bar{\sigma}^C$ are the actual mean of permanent household income of group j a common variance to all the groups, respectively.

Since the expected aggregate idiosyncratic shock to income ν_i in equation (8) can be regarded to be zero, as outlined in Mace (1991), we approximate the permanent mean income level to be equal to the mean income of the empirical distribution. In other words we assume that:

$$\bar{y}_j = \bar{y}_j^P,$$

with $\bar{y}_j$, the sample mean income of household group j , being used as a proxy for mean permanent household income.

As far as the $\bar{\sigma}^C$ is concerned, recalling that consumption is defined in equation (12), it is derived from the following variance:

$$\text{var}(c_i) = \text{var}(y_i^P) + \text{var}(\varepsilon_i),$$

where $\text{var}(c_i)$ is measured from consumption data and $\text{var}(y_i^P)$ is estimated with the log likelihood function; thus, $\text{var}(\varepsilon_i) = \text{var}(c_i) - \text{var}(y_i^P)$, as a consequence. In this way, we can implement the ML estimation of the mixture model by constraining the mean and the variance of the underlying distributions to reflect the mean (permanent) incomes of each of the three groups (each of proportion $\tilde{\pi}_1$, $\tilde{\pi}_2$ and $\tilde{\pi}_3$), and the variance as

specified above. This leads to the estimate of the density of the counterfactual distribution of full consumption smoothing as follows:

$$(14) \quad \tilde{f}(c_i) = \sum_{j=1}^3 \tilde{\pi}_j f(c_i; \tilde{\theta}_j),$$

with $\tilde{\theta}_j = (y_j^P; \text{var}(\bar{\varepsilon}_j))$.

As far as the distribution of consumption with no smoothing is concerned, the procedure outlined above is repeated by allowing the variance of consumption to differ by household income group. This leads to a mixture of normal model specification as follows;

$$(15) \quad \hat{f}(c) = \sum_{j=1}^m \frac{1}{\bar{\sigma}_j^C} \phi_m \left(\frac{c_i - \bar{y}_j^P}{\bar{\sigma}_j^C} \right),$$

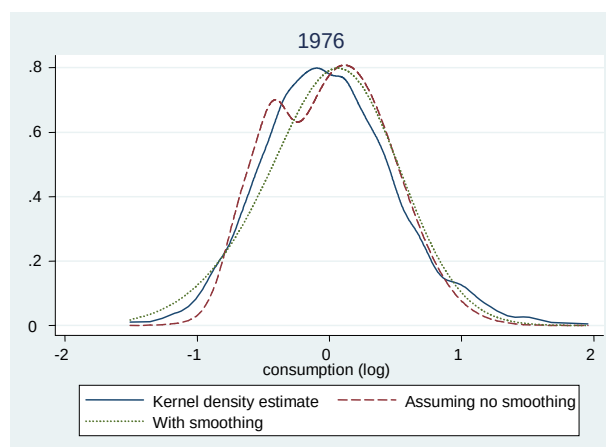
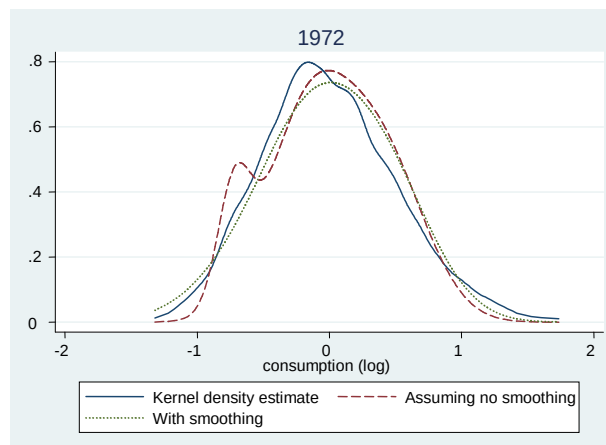
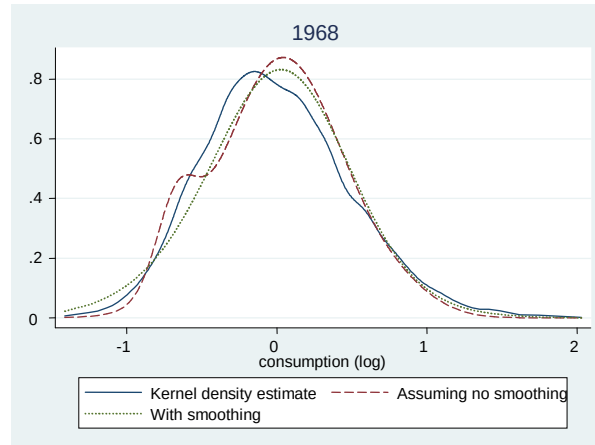
where the standard deviation of consumption is group-specific.

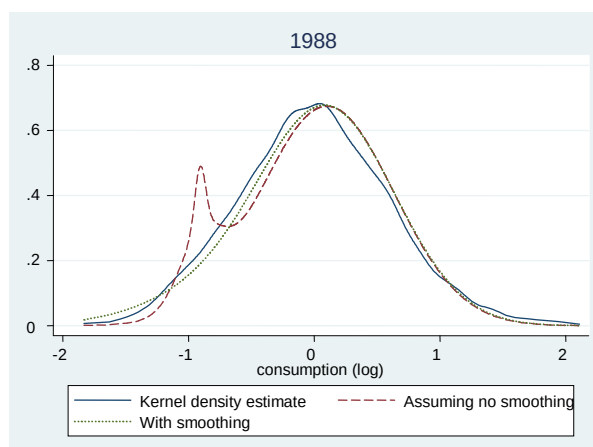
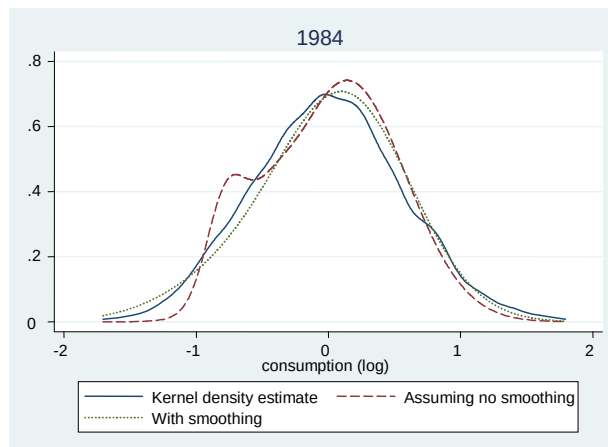
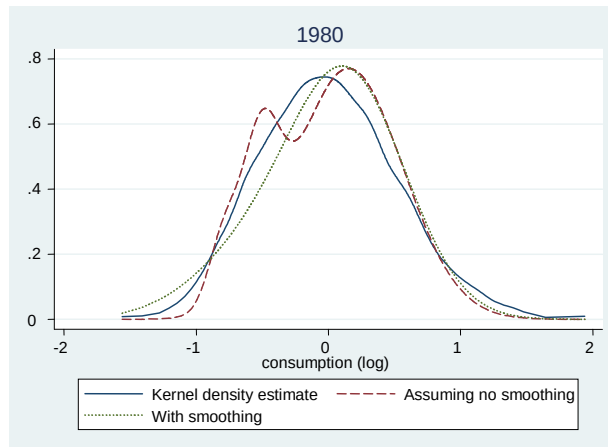
Graph 7 below provides a visual overview of how the density of actual household consumption compares with the two extreme-case distributions estimated in (6), which reflects perfect smoothing by households, and (15), which we have created assuming households do not smooth consumption at all. The first important indication thereby provided is the close fit with which the density of actual consumption approximates the density of consumption in the case of no smoothing; even though the profile of the former curve shows some signs of light bimodality by the end of the time period analyzed, the two curves become increasingly similar. This is a very meaningful result,

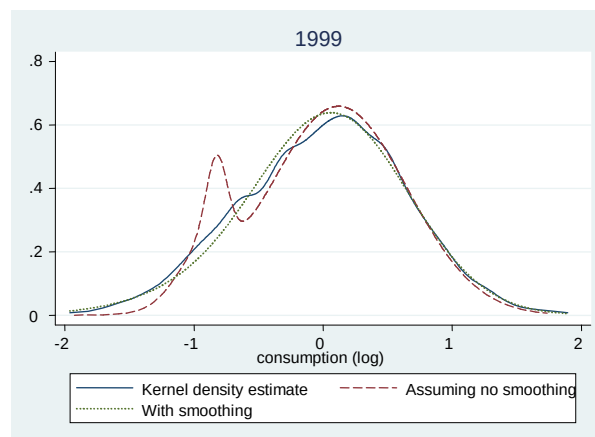
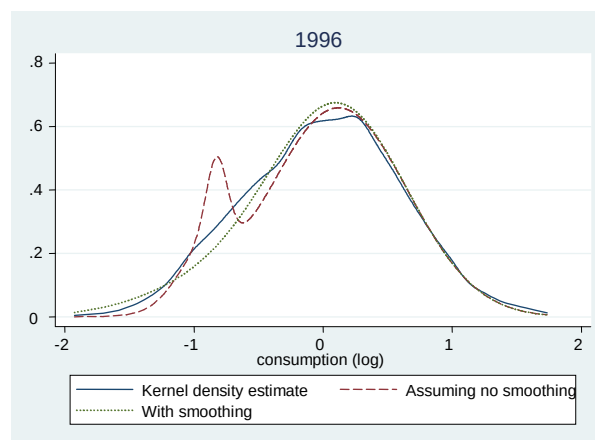
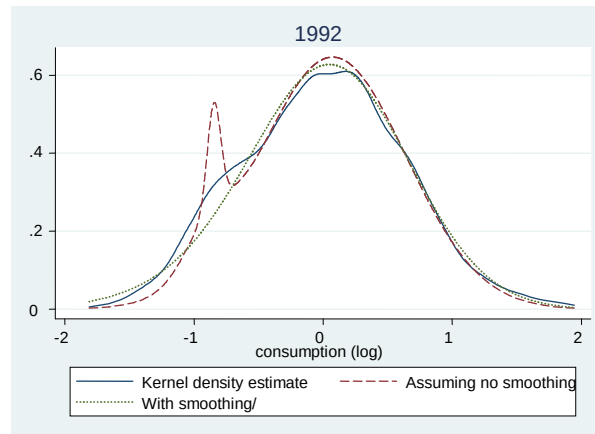
which underlies how insurance mechanisms may effectively be in place and lead to a sizable level of smoothing.

Graph 7

Density of Consumption-Actual, smoothing and no-smoothing

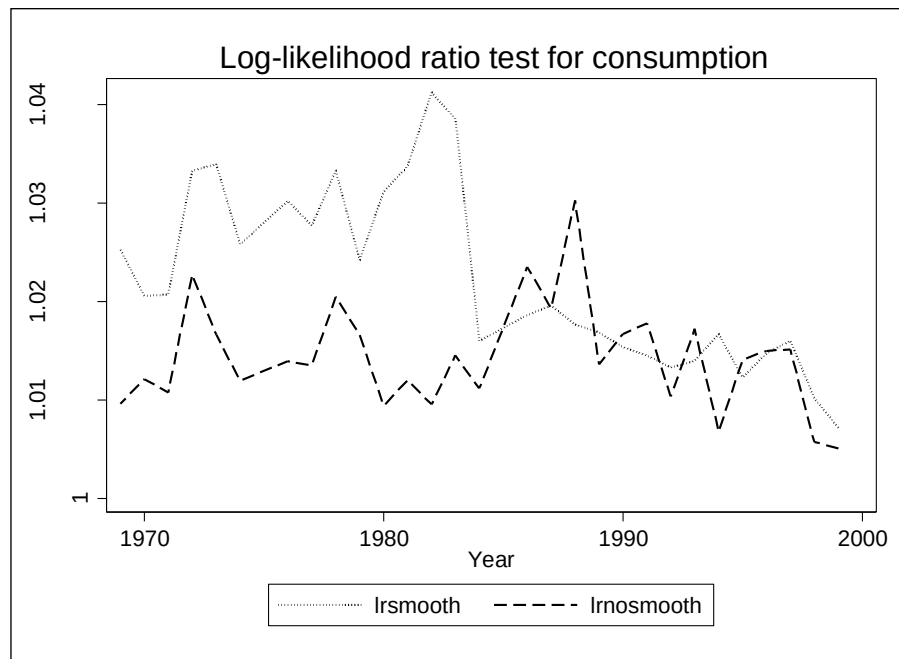






Further to the graphical analysis set out above, it is now useful to formally assess how well the mixture models of consumption we have built under the hypotheses of smoothing and no-smoothing fit the data relative to the mixture model based on the actual data. As we have used the MLE technique to estimate the parameters of the underlying distributions of consumption, the corresponding likelihood value, which is at the heart of the MLE parameter estimation method, can therefore be also employed to assess how the models constructed under the hypotheses of full-smoothing and no-smoothing fit the distribution of the data set.

Graph 8



Thus, we compute the log-likelihood value (L) for each of the three models of consumption relevant to our analysis: (i) L_C , which is the log-likelihood value of the mixture of (three) normal distributions model that makes the modelling of the actual

consumption data consistent with the modelling of our income data;¹² (ii) L_N , which is the log-likelihood value of the model related to the distribution of consumption that underlies no smoothing under the mixture of normals approach; and (iii) L_S , which is the log-likelihood value of the model related to the distribution of consumption that underlies perfect smoothing under the mixture of normals approach. Then, we compute two ratios, L_N/L_C (*lrnosmooth*) and L_S/L_C (*lrsmooth*), which represent the ratios of the log-likelihood values of the non-smoothed and smoothed mixed distribution models of consumption relative to the mixture distribution model of consumption data constructed with the mixture the three normal model, as specified above. Ratios above one would reflect a better fit of the relevant model.

The trend in the likelihood ratio test is depicted in graph 8 above. Even though the magnitude of the test is not very sizable, both ratios show a slightly improved fit when the two mixture models under the smoothing and no smoothing assumptions are considered; this is particularly true for the distribution of consumption under the assumption of full smoothing during the first part of the time period concerned. However, the fit seems to get worse in both models as the years approach the end of the period.

4.4 Expected consumption

The third and last part of our analysis is concerned with the evaluation of the expected level of consumption for each group G_j , obviously assuming consumption to follow a mixture of normal distribution with the same number of groups (and same proportions) as estimated in equation (6). This allows us to predict the level of consumption, conditional to the level of income, and compare the evolution of both variables over

¹² The only restriction applied to the distribution model L_C is based on is the fitting of consumption data according to the three mixtures of normal distributions, analogous to income data.

time. Such a comparison is also useful in order to infer the degree of cross-group consumption insurance, as we can compare how the differences between the two underlying variables change over time from group to group.

In order to do so, we make use of probability theory and specify our expected consumption equation in the following way:

$$(16) \quad E(c_i | G_j) = \frac{1}{n} \frac{1}{P(G=j)} \sum_{i=1}^n c_i P(G=j | y_i),$$

where $P(G=j)$ is simply $\tilde{\pi}_j$, and $P(G=j | y_i) = f(y_i | G_j)P(G=j)/f(y_i)$; the denominator in the right hand side of this expression is just the empirical density function of income, whilst $f(y_i | G_j)$, which represents the normal density of the income of the households belonging to group j , can be easily computed by considering a normal standard density whose means and variances are those of the mixture of normals model estimated by (6):

$$(17) \quad \tilde{f}(y_i | G_j) = \phi\left(\frac{y_i - \tilde{\mu}_j^Y}{\tilde{\sigma}_j}\right) \frac{1}{\tilde{\sigma}_j}.$$

This allows us to compare mean consumption with mean income $E(y_i | G_j)$, which we have already obtained from (6).¹³

The joint evolution of mean household consumption and mean income level, by income group, is shown in the graphs 9.a to 9.c below. Graph 9.a summarizes these trends for

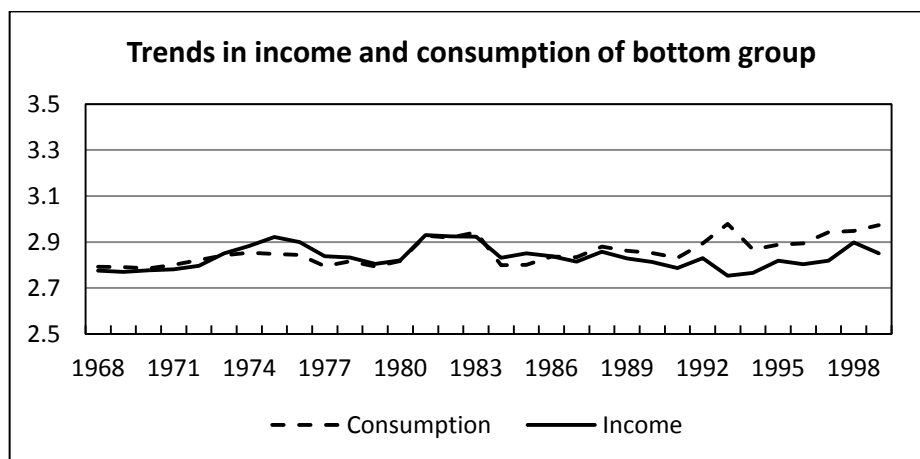
¹³ As $E(y_i | G_j)$ is just the mean household income $\tilde{\mu}_j^Y$ belonging to the vector of parameter $\tilde{\theta}_j$ estimated by the mixture of normals model we have employed.

the group of the poorest households, portrays a substantially similar and stable trend during the whole period, with a consumption starting to trace an increasing positive gap over income from the late 80s; this is a significant fact, likely to be related to the increase in the level of risk insurance by the less affluent households. The trends relative to the middle income households are shown in graph 8.b, where income and consumption evidence, with some ups and downs in between, a marked decrease. Consumption is almost constantly under the level of income, with the exception of the second half of the 90s. Graph 9.c displays income and consumption paths for the households that are better off. In this case it is possible to notice something opposite (relative to the other two groups) to what evidenced in the other two graphs; during the beginning of the time span under consideration, the level of consumption is above the level of income. Then, after the two variables have both stopped their decline, there is a slight, but constant, reversal of their trend, but consumption remains positioned below income.

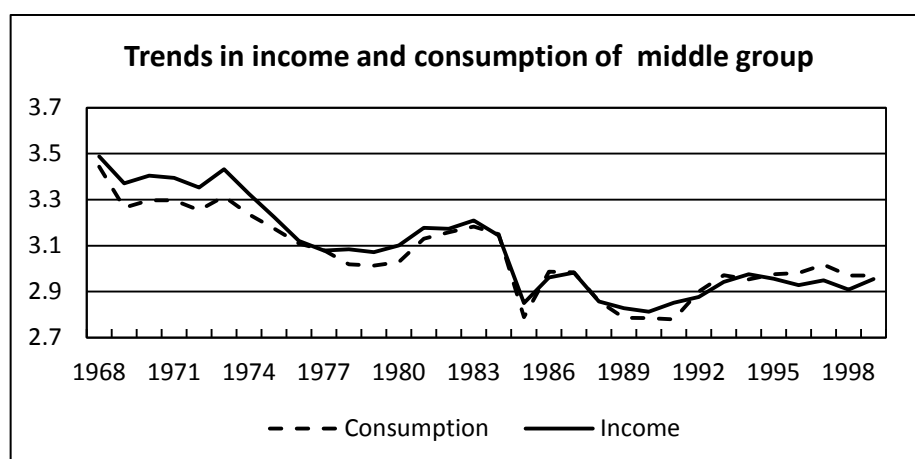
Graph 9

Trends in expected income and consumption by group

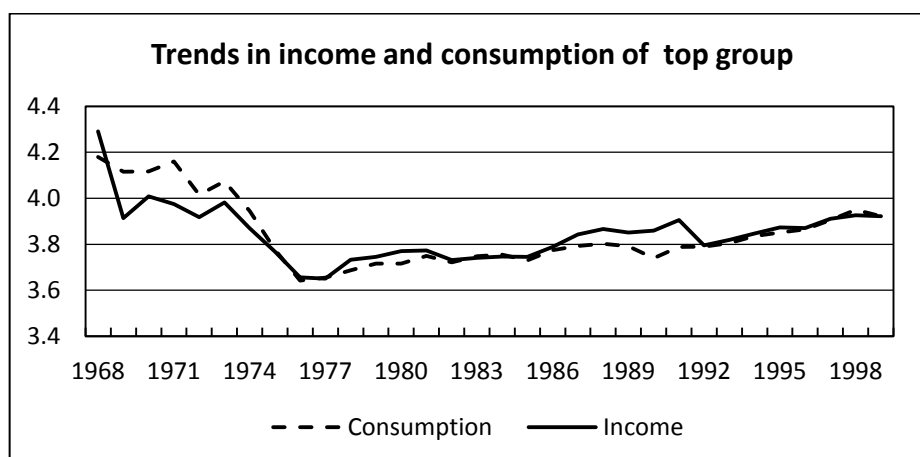
(a)



(b)



(c)



4.5 Predicted consumption given income

The final part of the empirical analysis of the paper aims to evaluate how well our mixture model is capable to predict the level of consumption at each level of income.

We can define the predicted level of consumption as follows:

$$(18) \quad E(\tilde{c} | y) = \sum_{j=1}^3 E(c | G_j) P(G_j | y),$$

where $P(G_1 | y)$ is the same as explained in the previous sub-section for equation (16), while

$E(c | G_j)$, is simply the average permanent income of group j , $\bar{y}_j^P$, as estimated by our model in (6). The level of predicted consumption thereby obtained is then compared with the level of consumption predicted by actual data and the actual level of household income. The former is estimated by the empirical distribution as:

$$(19) \quad E(c | y) = \frac{1}{nh} \sum_{i=1}^n c_i k(y_i),$$

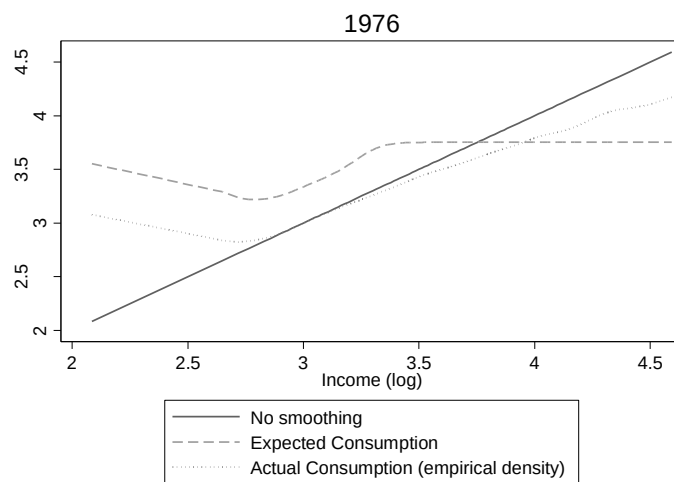
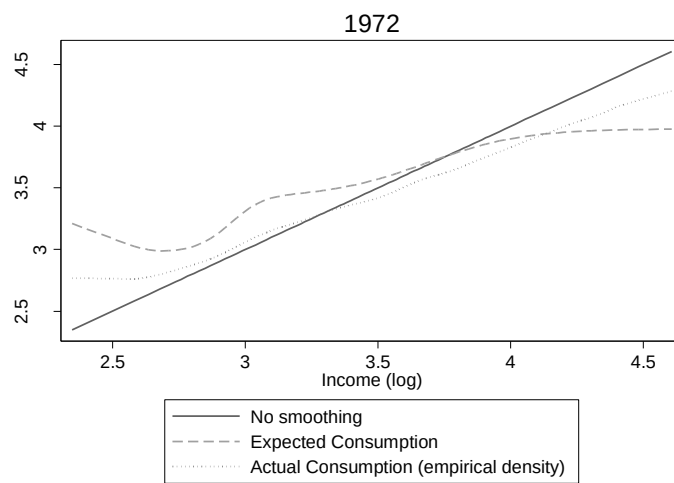
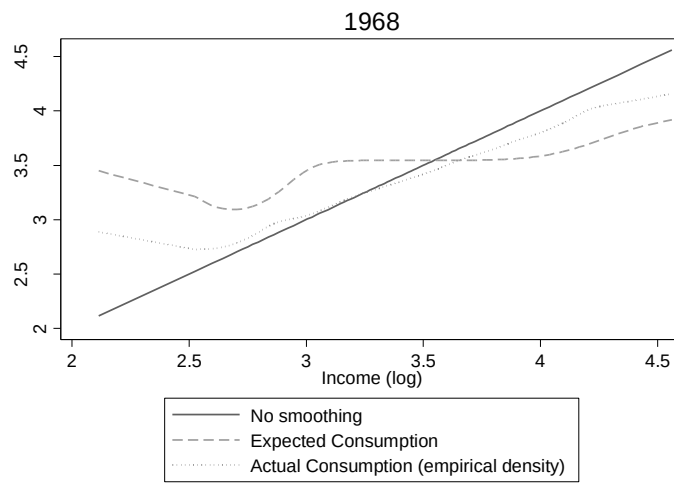
where each level of consumption of household i is weighted by its kernel density value $k(y_i)$. The latter is simply the empirical income level.

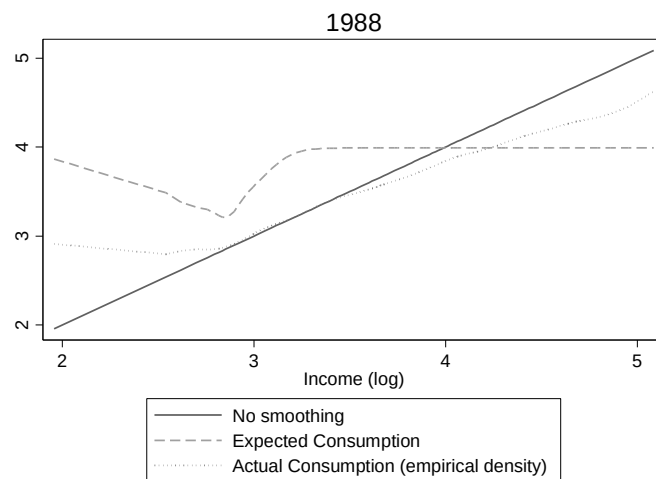
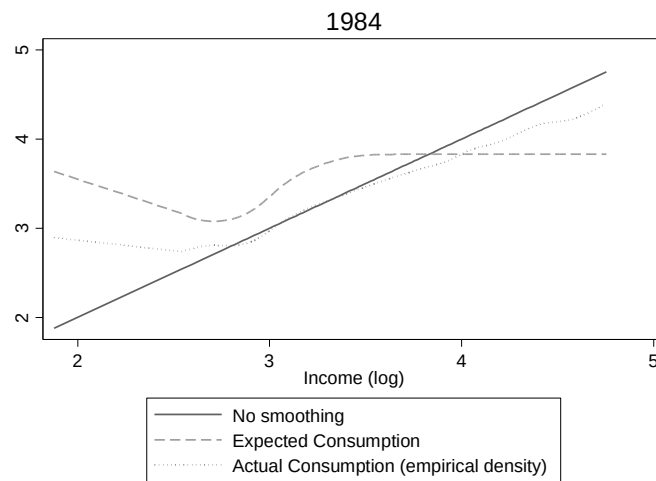
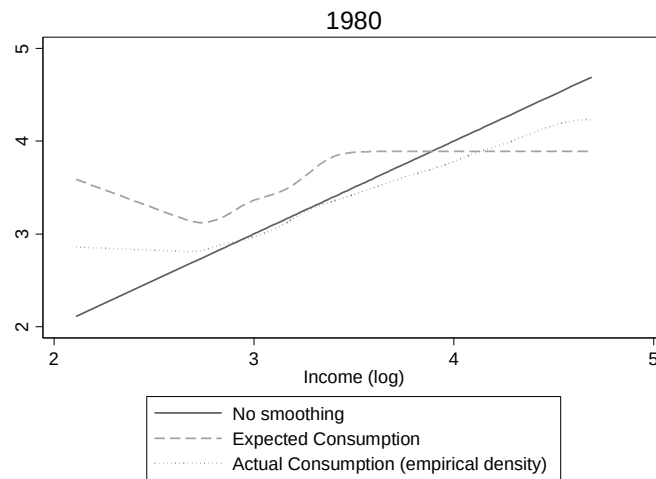
Graph 8 below plots the two types of predicted consumption against the level of income – obviously, the solid line represents the case of zero smoothing, with income = consumption.

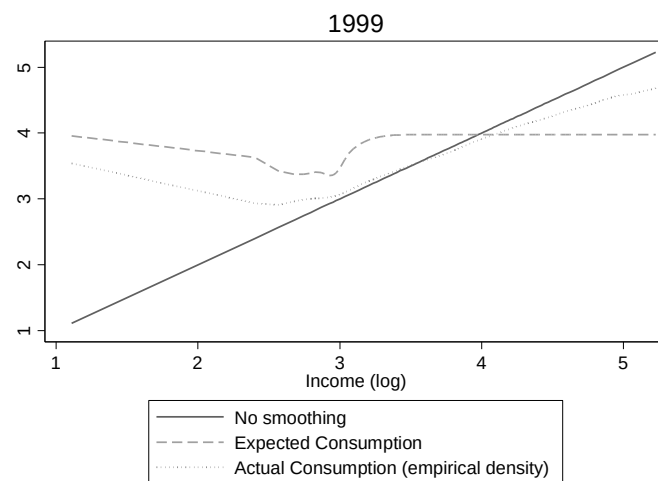
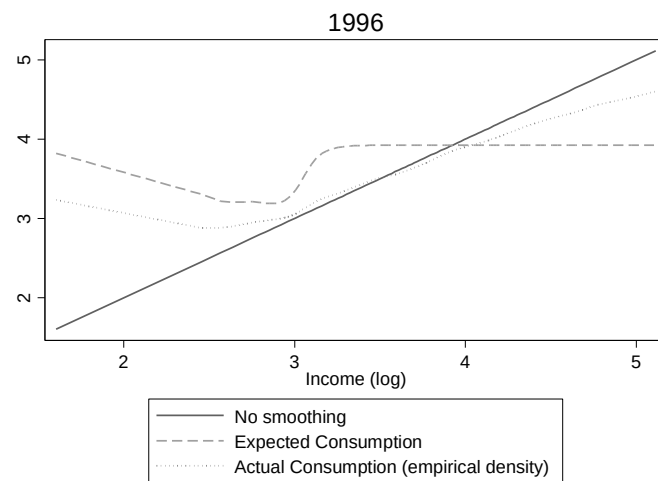
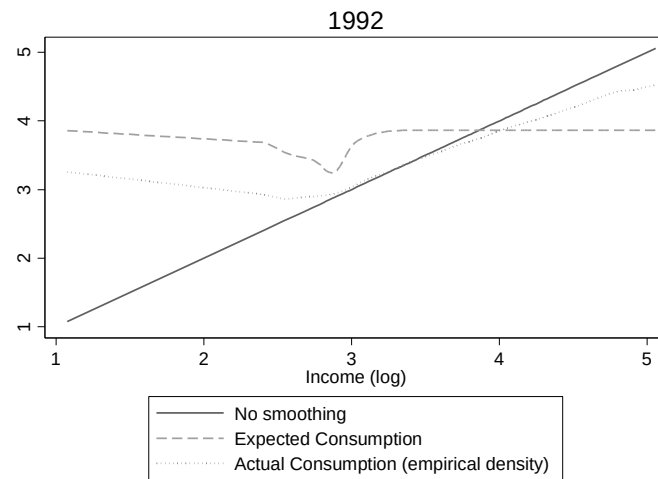
The empirical evidence shows that, even though our model tends to overestimate the degree of smoothing for the households in the lower part of the distribution of income and to underestimate it for the high income families, it is able to predict actual consumption relatively well; this is particularly true for the beginning and the end of the time period analysed. It is possible to notice that the line of the model's prediction flattens beyond a certain income level; this is probably due to the relatively large share of the households belonging to the higher income group. The declining trend of the model's line at low levels of income is likely to reflect the decreasing density of household with higher income levels in that part of the distribution – at the very end of it is more reasonable to expect to find households belonging to the high income group whose members are temporarily without earnings.

In terms of relationship between income and consumption, the graphs clearly show that low income families are able to smooth consumption, and the extent of this smoothing (or consumption subsidization) is increasing after the end of the 1990s.

Graph 8
Predicted Consumption vs. Income







6 Conclusions

This paper has shown how the use of a mixture model approach can contribute to the empirical analysis of consumption insurance by modelling the distributions of income and consumption and comparing them on a cross sectional basis and over time. The main results of the paper can be summarized as follows.

Firstly, mixtures of normal distributions proves to be a reliable way to fit income data and address the heterogeneity issue; in our case, the output of the ML estimation reveals that income is composed by three normal sub-distributions, which provides a satisfactory description of the multimodality features of the empirical income density, as illustrated by the first set of our graphical analysis.

Secondly, we have been able to simulate the density of consumption under the hypothesis that (full) insurance/perfect smoothing occur; by comparing this density with the densities of actual consumption and of the consumption with perfect smoothing – indeed the density of income estimated by our mixture model – one can observe to what extent household have been able to insure against consumption risk. One important finding here, as described by the relevant graphs, is that the distribution of actual consumption is closer to the counterfactual distribution of smoothed consumption rather than to the distribution reflecting consumption = income. Nevertheless, consumption insurance is certainly in place, at least to some extent, for household in the upper part of the distribution, especially after the mid 1980s. This underlies a likely positive effects deriving from the boost to the process of financial liberalization that has affected the UK in that period.

Overall, smoothing is taking place, to a larger extent, during the 1970s and the 1980s. It is also possible to notice that the degree of consumption insurance tends to increase after the beginning of the 1980s, when the distribution of income becomes more polarized. The degree of smoothing is higher for the household in the mid-bottom part of the distribution, especially near the convergence of the bottom and middle groups. Finally, we have computed the expected level of consumption for each household group, given the level of income. This allows us to analyse the trends in the two variables for each household group over time. It appears that the insurance mechanism is effectively working, especially for the households with the lower level of income, which are able to increase the degree of consumption during the 1990s. Trends of the so called middle group of households also show some degree of smoothing in the final part of the time period.

Bibliography

Asdrubali P. and S. Kim (2008). "Incomplete Intertemporal Consumption Smoothing and Incomplete Risk Sharing". *Journal of Money, Credit and Banking*, Vol. 40 Issue 7, pp 1521–1531.

Attanasio O. P. and Browning, M. (1995). "Consumption over the Life Cycle and over the Business Cycle". *American Economic Review*, Vol. 85, No. 5, pages 1118-37.

Attanasio O., P. K. Goldberg and E. Kyriazidou (2000), "Credit constraints in the market for consumer durables: evidence from micro data on car loans". NEBR Working Paper No 7694.

Autor D. H., L. F. Katz and M. S. Kearney (2008). "Trends in U.S. Wage Inequality: Revising the Revisionists". *Review of Economics and Statistics* Vol. 90:No. 2, pp 300-323.

Autor, David H., Lawrence F. Katz and Melissa S. Kearney. "The Polarization Of The U.S. Labor Market," *American Economic Review*, 2006, v96(2,May), pp. 189-194.

Bar-Ilan A. and Blinder A. S. (1988). "The Life Cycle Permanent-Income Model and Consumer Durables". *Annals of Economics and Statistics / Annales d'Économie et de Statistique* , No. 9, La Théorie du Cycle de vie pp. 71-91.

Bertola G., Guiso L. and Pistaferri L. (2005), "Uncertainty and Consumer Durables Adjustment", *Review of Economic Studies*, Vol. 72, No. 4, pp. 973-1007.

Blundell R. and I. Preston (1998). "Consumption Inequality and Income Uncertainty," *Quarterly Journal of Economics*, Vol. 113 No. 2, pp. 603-640.

Blundell, R, L. Pistaferri, and I. Preston (2008). "Consumption Inequality and Partial Insurance." *American Economic Review*, Vol. 98, No. 5, pp. 1887–1921.

Canova F. and M. O. Ravn (1996). "International consumption risk sharing". *International economic review*. Vol. 37, No. 3 pp 573-601.

Celeux, G. (2007). "Mixture Models for Classification". In: *Advances in Data*, D. Reinhold Editor. Springer Publisher.

Chesher A. (1984). "Testing for Neglected Heterogeneity". *Econometrica*, Vol. 52, No. 4, pp. 865-872.

Chiappori P.A. (1988). "Rational Household Labor Supply". *Econometrica*, Vol. 56, No. 1, pp. 63-89.

Clarida R. (1991). "Aggregate Stochastic Implications of the Life Cycle Hypothesis". *Quarterly Journal of Economics*, Vol 106, No. 3, pp 851-867.

Cowell F. A. (2011). "Measuring inequality". Oxford University Press.

Crucini M. J. (1999), "On international and national dimensions of risk sharing". *Review of Economics and Statistics*. Vol. 81, Issue 1, pp 73–84.

Cutler, D. M. and L. F. Katz (1992). "Rising Inequality? Changes in the Distribution of Income and Consumption in the 1980's". *American Economic Review*, Vol. 82 No. 2, pp. 546-551.

Deere D.R. and Vesovich J. (2006). "Educational wage premiums and the US income distribution: a survey". In: *Handbook of the economic of education*. Vol. 1.

Dempster A. P., N. M. Laird and D. B. Rubin (1977). "Maximum likelihood from incomplete data via the EM algorithm". *Journal of the Royal Statistical Society: Series B*, Vol. 39 No. 1, pp 1–38.

Demyanyka Y. and V. Volosovych (2008). "Gains from financial integration in the European Union: Evidence for new and old members". *Journal of International Money and Finance*. Vol. 27, Issue 2, pp. 277-294.

Flachaire E. and O. Nuñez (2007). "Estimation of the income distribution and detection of subpopulations: An explanatory model". *Computational Statistics & Data Analysis*. Vol. 51 Issue 7 pp.....

Fratzscher M. and J. Imbs (2009). "Risksharing, finance, and institutions in international portfolios". *Journal of Financial Economics*, Vol. 94, Issue 3, pp. 428–447.

Gallegati, M., A. Palestrini, D. Gatti and E. Scalas (2006). "Aggregation of Heterogeneous Interacting Agents: The Variant Representative Agent Framework". *Journal of Economic Interaction and Coordination*. Vol. 1, Issue 1, pp 5-19.

Goodman A., P. Johnson and S. Webb (1997). "Inequality in the UK". Oxford University Press.

Goos M and Alan Manning (2007) "Lousy and Lovely Jobs: The Rising Polarization of Work in Britain". *The review of economics and statistics* pp. 89-1.

Hagenaars, A., K. de Vos and M.A. Zaidi (1994). "Poverty Statistics in the Late 1980s: Research Based on Micro-data". Office for Official Publications of the European Communities. Luxembourg.

Heathcote J., K. Storesletten and G. Violante (2009). "Quantitative Macroeconomics with Heterogeneous Households". *Annual Review of Economics* Vol. 1, Issue 1, pp 319-354.

Jappelli T. and L. Pistaferri (2010a), "The Consumption Response to Income Changes". *The annual Review of Economics*. Vol. 2.

Jappelli T. and L. Pistaferri (2011), "Financial integration and consumption smoothing". *The Economic Journal*, Vol. 121, Issue 553, pp 678-706.

Jenkins, S. P. (1991), "Income Inequality and Living Standards: Changes in the 1970s and 1980s". *Fiscal Studies*, 12: 1–28.

Krueger D. and F. Perri, (2006). "Does Income Inequality Lead to Consumption Inequality? Evidence and Theory". *Review of Economic Studies*, Vol. 73 No. 1, pp. 163-193.

Kuznets, S. (1955), "Economic growth and income inequality". *American Economic Review*, 45, pp. 1-28.

Lindsay B. G. (1995). "Mixture models: theory, geometry, and applications". *Institute of mathematical statistics*.

Marron J. S. and M. P. Wand (1992). "Exact Mean Integrated Squared Error". *The Annals of Statistics*, Vol. 20, No. 2, pp. 712-736.

McLachlan G. J. and D. Peel (2000). "Finite mixture models". *John Wiley and Sons*.

Meghir C. and L. Pistaferri (2004). "Income Variance Dynamics and Heterogeneity". *Econometrica*, Vol. 72, No. 1, pp. 1–32.

Morduch J. (1995). "Income Smoothing and Consumption Smoothing". *Journal of Economic Perspectives*. Vol. 9 No. 3, pp 103-14.

Pittau M.G. (2005). "Fitting Regional Income Distributions in the European Union". *Oxford Bulletin of Economics and Statistics*. Vol. 67 No. 2, pp 135-161.

Quah D. T. (1996). "Empirics for economic growth and convergence". *European Economic Review*. Volume 40, Issue 6, pp. 1353-1375.

Quah D. T. (1996). "Twin peaks: growth and convergence in models of distribution dynamics". *Economic Journal* No.106, pp. 1045–1055.

Ventura L. (2008). "Risk sharing opportunities and macroeconomic factors in Latin American and Caribbean countries: a consumption insurance assessment". Policy Research Working Paper 4490, World Bank.