

Why Is Europe More Equal Than the United States?

Thomas Blanchet
Lucas Chancel
Amory Gethin

September 2020



World Inequality Lab

WHY IS EUROPE MORE EQUAL THAN THE UNITED STATES? *

THOMAS BLANCHET

LUCAS CHANCEL

AMORY GETHIN

Abstract

We combine all available household surveys, income tax and national accounts data in a systematic manner to produce comparable pretax and posttax income inequality series in 38 European countries between 1980 and 2017. Our estimates are consistent with macroeconomic growth rates and comparable with US Distributional National Accounts. We find that inequalities rose in most European countries since 1980 both before and after taxes, but much less than in the US. Between 1980 and 2017, the European top 1% pretax income

*Thomas Blanchet, Paris School of Economics – EHESS, World Inequality Lab: thomas.blanchet@wid.world; Lucas Chancel, World Inequality Lab, Paris School of Economics and IDDRI, SciencesPo lucas.chancel@psemail.eu; Amory Gethin, World Inequality Lab – Paris School of Economics: amory.gethin@psemail.eu. This is a revised and extended version of the work by the same authors that previously circulated under the title “How Unequal is Europe? Evidence from Distributional National Accounts, 1980-2017”. We thank Daron Acemoglu, Facundo Alvaredo, Andrea Brandolini, Bertrand Garbinti, Jonathan Goupille-Lebret, Antoine Bozio, Stefano Filauro, Ignacio Flores, Branko Milanovic, Thomas Piketty, Dani Rodrik, Emmanuel Saez, Wiemer Salverda, Daniel Waldenström, and Gabriel Zucman for helpful comments on this paper. We are also grateful to Alari Paulus for his help in collecting Estonian tax data. Matéo Moisan provided useful research assistance. We acknowledge financial support from the Ford Foundation, the Sloan Foundation, the United Nations Development Programme and the European Research Council (ERC Grant 340831).

share rose from 8% to 11% while it rose from 11% to 21% in the US. Europe's lower inequality levels are mainly explained by a more equal distribution of pretax incomes rather than by more equalizing taxes and transfers systems. "Predistribution" is found to play a much larger role in explaining Europe's relative resistance to inequality than "redistribution": it accounts for between two-thirds and ninety percent of the current inequality gap between the two regions.

JEL codes: H23, H24, H51, H52, H53, E01.

I. INTRODUCTION

While inequality has attracted a considerable amount of attention in recent academic and policy debates, basic questions about the distribution of growth and redistribution remain unanswered. To what extent was the remarkable rise in income inequality in the US over the past decades inevitable? What can be learned from the study of inequality and growth in other high income countries, and in particular European nations, to answer this question? The comparative study of the distribution of growth, taxes and transfers can provide critical insights into such debates. However, to date, inequality and redistribution estimates across countries have been hard to interpret because of a lack of conceptual and empirical consistency.

The standard source to measure economic growth across countries is the national accounts, while the standard source to measure and analyse inequality and redistribution across countries is household surveys. Surveys are known to underrepresent top incomes and do not add up to macroeconomic income totals, leading to potential inconsistencies in the study of growth, inequality and redistribution. In order to address some of these limitations, [Piketty and Saez \(2003\)](#), [Piketty \(2003\)](#) and several authors have relied on tax data to study long run trends in inequality. Yet tax data also has a lot of issues in terms of comparability and consistency. In particular, tax data measures concepts that vary between countries, and miss the large share of income that goes untaxed. To overcome this problem, [Alvaredo et al. \(2016\)](#) and [Piketty, Saez, and Zucman \(2018\)](#) developed "Distributional National Accounts" that combine various sources so as to distribute the entirety of a country's national

income, and established guidelines to carry out this work in other parts of the world.

This work has attracted interest from academics (triggering important methodological debates¹) and policymakers but unfortunately, with the exception of France (Bozio et al., 2018; Garbinti, Goupille-Lebret, and Piketty, 2018) similar work in comparable countries remains rare. This makes it difficult to compare the situation of the United States with other countries. In Europe, the lack of estimates of the income distribution that integrate surveys, tax and national accounts are not the result of a lack of data *per se*. In fact, there is a fair amount of data available, at least since the 1980s. The problem is that these data are scattered across a variety of sources, taking several forms, using diverse concepts and different methodologies. So we find ourselves with a disparate set of indicators that are not always comparable, are hard to aggregate, provide uneven coverage, and can tell conflicting stories.

As a result, the literature has so far struggled to answer basic questions, such as whether European inequality as a whole is similar to that of the US, whether this has changed over time, and whether the differences between Europe and the US are due to pretax inequality or to redistribution. This paper addresses these substantive and methodological issues by constructing the first distributional national accounts for 38 European countries since 1980, which are comparable with recently produced distributional national accounts for the US. While we still face considerable challenges in the construction of good estimates of the income distribution in some countries, we believe that the series described in this paper present major improvements over existing ones.

Our estimates combine virtually all the existing data on the income distribution of European countries in a consistent way. That includes, first and foremost, household surveys, national accounts, and tax data. It also includes additional databases on social contribution schedules, social benefits by function, and government spending on health that have been compiled by several institutions over the years (OECD, Eurostat, WHO). Our methodology exploits the strengths of each data source to correct for the weaknesses of the others. It avoids the kind of systematic errors that would arise from the comparison of different income concepts, different statistical units and different methodologies. As such, our estimates are meant to reflect the

¹ See for instance Auten and Splinter, 2019a, which we further discuss below

best of the literature's knowledge on what has been the evolution of inequality in Europe.

In line with the logic of Distributional National Accounts (DINA), we distribute the entirety of the national income. This includes money that never explicitly shows up on anyone's bank account, such as imputed rents, production taxes or the retained earnings of corporations, yet can account for a significant share of the income recorded in national accounts and official publications of macroeconomic growth. Therefore, our results are consistent with macroeconomic totals, and provide a comprehensive picture of how income accrues to individuals, both before and after government redistribution. Using a broad definition of income makes our results less sensitive to various legislative changes, and therefore more comparable both over time and between countries.

Rather than focusing on a handful of indicators, we cover the entire distribution from the poorest income groups (well covered by household surveys) to the top 0.001% (which we can capture thanks to tax data) for income both before and after redistribution. We can aggregate our distributions at different regional levels, and analyze the structure of inequality with a good level of detail. We can use our estimates to compute any set of synthetic indicators (such as top and bottom income shares, poverty rates or Gini coefficients) in a consistent way at the country or regional level. We can test how all these results are affected by alternative assumptions on the distribution of unreported income or on the mismeasurement of top incomes in surveys.

First, we find that inequality has increased in Europe as a whole and within nearly all European countries, before and after taxes since 1980. The share of pretax income that accrued to the richest 1% of Europeans rose from 8% to 11% before taxes and from 6% to 8% after taxes between 1980 and 2017. This rise of inequality is much more contained than in the US, where the top 1% pretax income share rose from 11% to 21% over the period, and from 9% to 16% after taxes. We also find that European inequality increased between 1980 and 1990, but less so afterwards, while the rise was sustained in the US.

Second, our results show that European income inequality levels are mostly explained by within-country inequality, rather than by average income differences

between Western, Northern and Eastern European countries. Between-country average income differentials are also found to explain none of the inequality trends observed in Europe as a whole over the past four decades. Still, regional dynamics vary: Eastern Europe has experienced the highest inequality increase, while the trend is more muted in Western Europe. Northern European countries did experience a significant increase of inequality but remain the most equal region, before and after redistribution.

Third, the main reason for the Old Continent's relative resistance to the rise of inequality has little to do with the direct impact of taxes and transfers. While Western and Northern European countries redistribute more than the US (about 48% of national income is taxed and redistributed in Europe vs. 30% in the US), the distribution of taxes and transfer does not explain the large gap between Europe and US posttax inequality levels. In fact, most of the gap can be explained by pretax inequality levels: using Gini coefficients, we find that pretax inequality levels account for around 90% of the difference between US and Europe posttax inequality levels. We do not find clear evidence either that taxes and transfers are more progressive in Europe than in the US. It follows that policies and institutional setups ensuring more equal pretax income distributions matter tremendously to keep inequality in check.

The rest of this paper is organized as follows: section II. reviews the existing literature on the measurement of inequality and redistribution in Europe and the US, section III. presents our conceptual framework, section IV. discusses the data sources and details of our methodology, section V. presents our findings on the distribution of pretax incomes in Europe while section VI. discusses the impact of taxes and transfers on inequality in Europe and the US and section VII. concludes.

II. RELATED LITERATURE

This paper contributes to the growing literature combining distributional data with national accounts to measure income and wealth inequalities. Following the contributions of [Piketty \(2003\)](#) and [Piketty and Saez \(2003\)](#), who used income tax tabulations to study the evolution of top incomes in France and the United States in the course of the twentieth century, a new body of research has combined income

tax returns and Pareto interpolation techniques to compute estimates of top income shares in a number of countries (see [Atkinson and Piketty, 2007; 2010](#) for a global perspective). This area of study has provided a number of insights into the long-run evolution of inequality. However, top income shares tend to rely on country-specific definitions of taxable income and tax units, and only cover a small fraction of the population (generally the top 10% or top 1%). Fiscal income also diverges from the national income, due to the existence of tax exempt income components, and is therefore inconsistent with macroeconomic growth figures. In contrast, household surveys are generally more consistent when it comes to measuring comparable income concepts. However, they tend to underestimate the incomes of top earners, because of small sample sizes ([Taleb and Douady, 2015](#)), and because the rich are less likely to answer surveys ([Korinek, Mistiaen, and Ravallion, 2006](#)) and more likely to underreport their income (e.g. [Cristia and Schwabish, 2009; Paulus, 2015; Angel, Heuberger, and Lamei, 2017](#)).

Recent studies have attempted to overcome these issues by combining surveys with tax and national accounts data. Statistical institutes and international organizations have increasingly recognized the need to bridge the micro-macro gap. Since 2011, an expert group on the Distribution of National Accounts mandated by the OECD has been working on methods to allocate gross disposable household income to income quintiles ([Fesseau and Mattonetti, 2013; Zwijnenburg, Bournot, and Giovannelli, 2019](#)). In a similar fashion, experimental statistics on the distribution of personal income and wealth have been recently published by [Eurostat \(2018\)](#), [Statistics Netherlands \(2014\)](#), [Statistics Canada \(2019\)](#) and the [Australian Bureau of Statistics \(2019\)](#). These exercises have improved upon traditional survey-based estimates, but do not make systematic use of tax data and are restricted to the household sector. This can make estimates of inequality sensitive to the tax base in ways that are not economically meaningful, since firms can have differential incentives to distribute dividends or accumulate retained earnings depending on local tax legislation.

[Piketty, Saez, and Zucman \(2018\)](#) were among the first to allocate all components of the US national income to individuals based on tax microdata and explicit assumptions about the distribution of tax exempt income. This work (hereafter

referred to as "PSZ") spurred debate in particular with [Auten and Splinter \(2019b\)](#) who argue that different assumptions made on the distribution of non-taxable income limit the secular rise of US pretax income inequality documented in [Piketty, Saez, and Zucman \(2018\)](#) by a factor of 2.² The recent academic debate around US inequality statistics sheds light on the need to develop common methods and concepts to distribute the national income. The interest of the DINA method precisely lies in its scope and international comparability.³ Several studies have recently followed a similar methodology to extend the DINA approach to other countries, including close work with national statistical institutions (e.g. with the French Statistical Institute, see [Germain, 2020](#)) to develop an internationally comparable integrated framework of distributional measures of income and wealth growth. This is the framework that we adopt in this paper: that is, we combine in a systematic manner data from surveys, income tax returns and national accounts to estimate the distribution of the national income in thirty-eight European countries between 1980 and 2017.

This paper contributes to the existing literature on the evolution of income inequality in Europe. It has generally been acknowledged that income inequality has grown in Europe since the 1980s ([OECD, 2008](#)), but little is known of how this rise compares across countries, across income groups in the distribution, or across time periods. The effort made by the Luxembourg Income Study (LIS) to harmonize existing surveys has been extremely helpful in improving the comparability of pre-2000 inequality statistics in Western Europe. Yet, because of sampling issues and misreporting at the top of the income and wealth distributions, surveys can reveal inequality trajectories which are inconsistent with those suggested by tax data. Similar limitations apply to Eastern Europe: the historical survey tabulations studied by [Milanovic \(1998\)](#), the EU-SILC surveys now conducted in new EU member countries and the top income shares recently estimated from income tax data (e.g.

²See [Auten and Splinter \(2019a\)](#), Table F1. More precisely, two thirds of the difference in top 1% income share levels in 2014 between [Piketty, Saez, and Zucman \(2018\)](#) and [Auten and Splinter \(2019a\)](#) are accounted for by differences in (i) the treatment of underreported income, (ii) the treatment of retirement income and (iii) adjustments made to the definition of taxable income (see [Auten and Splinter \(2019a\)](#), Table TB-6).

³A comprehensive discussion of the DINA methodology is presented in [Alvaredo et al. \(2016\)](#). Recent studies following the DINA approach include [Morgan \(2017\)](#) for Brazil and [Chancel and Piketty \(2019\)](#) for India.

[Novokmet, 2018](#); [Bukowski and Novokmet, 2019](#)) exhibit different dynamics. This paper is about combining all these sources in meaningful ways, using new techniques and a harmonized methodology.

Another question which has received much attention in recent years is that of the comparison between Europe and the United States. While it is acknowledged that posttax income inequality is greater in the US than in most European countries today, it remains unclear whether this gap is due to differences in pretax inequality or to differences in government redistribution, and how this gap evolved over time. International organization such as the OECD (e.g. [OECD, 2008](#); [2011](#)), in line with other researchers (e.g. [Jesuit and Mahler, 2010](#); [Immervoll and Richardson, 2011](#); [Guillaud, Oleckers, and Zemmour, 2019](#)) find that low posttax inequality levels in Europe compare to the US are essentially due to redistribution. These studies typically rely on household survey data only and tend to limit their analysis to direct taxes and transfers because of the difficulties associated to distributing the totality of national income before and after transfers. This contrasts with [Bozio et al. \(2018\)](#), who use the DINA methodology, distribute all taxes and transfers and find that redistribution reduces inequality less in France than in the United States, contrary to a widespread view. Whether the US are more unequal than Europe as a whole (i.e. as a region) also remains an open question.⁴

This article differs from existing studies on the distribution of income in Europe in a number of ways. First, we go beyond the available survey microdata by collecting and harmonizing a rich dataset of historical survey tabulations. This allows us to go back in time and consistently study the long-run evolution of income inequality in the large majority of European countries from the 1980s until today. Second, we use all available studies on the evolution of top income shares, as well as previously unused

⁴Works on the distribution of income in the EU-15 ([Atkinson, 1996](#)) or the Eurozone ([Beblo and Knaus, 2001](#)) suggested that income inequality was higher in the US, but recent studies extending the analysis to new, poorer Eastern European member states have found mixed results (e.g. [Brandolini, 2006](#); [Dauderstädt and Keltek, 2011](#); [Salverda, 2017](#); [Filauro and Parolin, 2018](#)). One of the limits of existing studies is that they are based on surveys. This may bias comparisons of income inequality between European countries and between Europe and the US if surveys capture more accurately top incomes in some countries than in others. The top-coding of incomes in the public-use samples of the US Current Population Survey, for instance, contrasts with the use of administrative data to fill in survey income components in several European countries, leading to important differences in the reliability of survey-based inequality estimates.

tax data sources, to correct for the under-representation of high-income earners. Third, we allocate all components of the national income to individuals, including tax exempt income, production taxes and collective government expenditure. This allows us to analyze the distribution of macroeconomic growth in Europe and the effects of different forms of redistribution on inequality. Importantly, we are able to test and document how all our results are robust to alternative assumptions and definitions of income, which we believe is important in the context of current debates on inequality statistics.

Methodologically, our approach also departs from previous distributional national accounts studies in the way we combine available data sources. [Piketty, Saez, and Zucman \(2018\)](#) and [Garbinti, Goupille-Lebret, and Piketty \(2018\)](#) start with tax data, to which they progressively add information from surveys and national accounts. This “top-down” approach exploits all the richness of the tax microdata and yields very detailed and precise estimates. However, while this type of work should be extended to as many European countries as possible, there are many countries and time periods (in Europe and other continents) for which tax microdata are simply not available. This justifies our “bottom-up” approach, which starts from surveys and gradually incorporates information from top income shares and unreported national income components. As such, we view our methodology as well-suited to estimating the distribution of the national income in countries gathering a mix of survey microdata, tabulated tax returns and a variety of heterogeneous data sources. This case corresponds to a majority of countries beyond Europe and the US.⁵

III. CONCEPTUAL FRAMEWORK

We study the distribution of the national incomes of thirty-eight European countries, spanning from Portugal to Cyprus and from Iceland to Malta, between 1980

⁵In a similar fashion, [Piketty, Saez, and Zucman \(2019\)](#) have recently proposed a simplified method for recovering estimates of top pretax national income shares based on the fiscal income shares of [Piketty and Saez \(2003\)](#) and very basic assumptions on the distribution of untaxed labour and capital income components. Our methodology is different because it starts from different sources, but follows a similar spirit. As we show in [Section IV.](#), we are able to reproduce very closely the results of [Garbinti, Goupille-Lebret, and Piketty \(2018\)](#) for France by combining their top fiscal income shares with available surveys and national accounts data.

and 2017. Our geographical area of interest includes the twenty-eight members of the European Union, five candidate countries (Bosnia and Herzegovina, Serbia, Montenegro, Macedonia, Albania), and five countries which are not part of the EU but have maintained tight relationships with it (Iceland, Norway, Switzerland, Kosovo and Moldova).

We follow as closely as possible the principles of the DINA guidelines ([Alvaredo et al., 2016](#)), which we briefly outline below. This allows us to compare our results with existing studies, including [Piketty, Saez, and Zucman \(2018\)](#) for the United States.

III.A. Macroeconomic Concepts

Net National Income Our preferred measure to compare income levels between countries and over time is the net national income. It is equal to gross domestic product (GDP) net of capital depreciation, plus net foreign income received from abroad. While GDP figures are most often discussed by academics and the general public, we believe national income to be more meaningful for our purpose, since capital depreciation is not earned by anyone, while foreign incomes are, on the contrary, received or paid by residents of a given country.⁶

From Survey and Taxable Income to National Income National income is the sum of the primary incomes of households, corporations, non-profit institutions serving households and the general government. Household income includes the compensation of employees, mixed income, and property income, which are generally — though imperfectly — covered by household surveys and tax data. It also includes the imputed rents of owner-occupied dwellings, which are much less often available from traditional sources but can represent a significant share of the capital income of households. The primary incomes of other institutional sectors can amount to a fifth of national income, but do not appear in either surveys or tax data. These mainly consist in the taxes on products and production received by the general government (net of subsidies) and the retained earnings of financial

⁶National income, like all national accounts concepts, and also like surveys, is about the *resident* population. So the income of a citizen of country A living in country B is attributed to country B.

and non-financial corporations. Taxes on production are a separate component of GDP and their distribution is a question of distributional tax analysis to which we come back below. Retained earnings correspond to profits that are kept within the company rather than distributed to shareholders as dividends. This income ultimately increases the wealth of shareholders and therefore represents a source of income to them.⁷ While the national income concept, contrary to GDP, does account for labor and capital foreign income flows, our estimates do not account for income received through wealth held offshore in tax havens, or for other forms of tax evasion. While offshore wealth is likely to matter more for the measurement of wealth than the measurement of income, we acknowledge that this limitation may lead to a moderate underestimation of top incomes. We leave this for future research.

III.B. Income Distribution Concepts

The DINA framework acknowledges three levels of distribution, called *factor income*, *pretax income* and *posttax income*. *Factor income* is the income that accrues to individuals as a result of their labor or their capital, before any type of government redistribution, be it through social insurance or social assistance schemes. *Pretax income* corresponds to income after the operation of the social insurance system (pension and unemployment), but before other types of redistribution. It is closest — though better harmonized and conceptually broader — to the “taxable income” in most countries. *Posttax income* accounts for other forms of redistribution of income operated by the government. We consider two types of posttax income. Posttax disposable income removes all taxes from pretax income, but only adds back cash transfers from the government, and therefore does not sum up to national in-

⁷Several papers have documented the impact of including retained earnings in the United States (Piketty, Saez, and Zucman, 2018), Canada (Wolfson, Veall, and Brooks, 2016), and Chile (Atria et al., 2018; Fairfield and Jorratt De Luis, 2016). In Norway, Alstadsæter et al. (2017) showed that the choice to keep profits within a company or to distribute them is highly dependent on tax incentives, and therefore that failing to include them in estimates of inequality makes top income shares and their composition artificially volatile. Previous work would sometimes include capital gains in their income definition, which indirectly accounts for this type of income. Yet this constitutes a poor proxy, because capital gains are recorded upon realization, rather than when they accrue to individuals. And whether capital gains are realized or not depends on their value and on tax incentives. Therefore, attributing retained earnings to individuals directly is more reliable, more consistent with macroeconomic measures of income, and more comparable across countries.

come. Posttax national income also adds back in-kind transfers (including collective consumption) and therefore adds up to national income.

In this paper, we will mostly focus on pretax and posttax national income (which we refer to below as posttax income). We also will discuss posttax disposable income, which is closer to posttax income concepts found elsewhere in the literature, and also more straightforward to compute. But posttax national income remains our concept of interest since it accounts for institutional differences between countries — in particular with respect to public health spending.

Factor Income On the labor side, factor income includes the entire compensation that firms pay to their employees, including social contributions paid by employees or employers, and mixed income. On the capital side, it includes the property income distributed to households, the imputed rents for owner-occupiers, and the primary income of the corporate sector (i.e. undistributed profits). We consider that the undistributed profits of privately owned corporations belong to the owners of the corresponding corporations. Factor income also includes the primary income of the government, which mostly corresponds to taxes on products and production, minus the interests that the government pays on its debt. We also add the retained earnings of publicly-owned corporations to that primary income. Our benchmark series distribute it proportionally to labor and capital income, in line with DINA recommendations ([Alvaredo et al., 2016](#)). But we also provide variants in which taxes on products are paid proportionally to consumption (see extended appendix).⁸

From Factor to Pretax Income Pretax income correspond to factor income, to which we add social insurance benefits in the form of unemployment and pension, and from which we remove the social contributions that pay for them. Note that for pretax income to sum up to national income, it is important to remove the same amount of social contributions as the amount of social benefits that we distribute. This way, we both avoid double counting and ensure that we look at the redistribution from social insurance in a way that is budget-balanced. In doing so, we observe

⁸Interest payments on government debt have no aggregate effect on national income because it represents a transfer from the government to households, but it does have a second-order distributional effect because ownership of government bonds is usually more concentrated than income.

significant heterogeneity between countries. In most countries, social contributions exceed pension and unemployment benefits, because social contributions also pay for health or family-related benefits that we classify as assistance-based (i.e. non insurance-based) redistribution. Therefore, we only deduct a fraction of social contributions from pretax income (the “contributory” part). But in some countries, like Denmark, social contributions are virtually non-existent. In these cases, we have to assume that social insurance is financed by the income tax, and therefore deduct a fraction of the income tax from factor income to get to pretax income.

From Pretax to Posttax Income To move from pretax to posttax income, we first remove all taxes and social contributions that remain to be paid by individuals. This includes the taxes on products and production that we previously added, and also the corporate tax that was added through undistributed profits. Then we add all types of government transfers, and government consumption. We distribute all of government consumption proportionally to income, with the exception of public health expenditures.⁹ We use the proportionality assumption as a benchmark for simplicity, transparency and comparability with earlier work on distributional national accounts, in particular in the United States (Piketty, Saez, and Zucman, 2018). But we consider that it is important to make an exception for health spending. Indeed, while many European countries have public health insurance systems, the United States have a mostly private one, with some public programs such as Medicaid and Medicare, which are explicitly distributed to their recipients in inequality statistics for the United States. Therefore, distributing health spending proportionally to income would, in comparison, understate the amount of redistribution that European countries engage in. We distribute health spending in a lump sum way, considering as a first approximation that the insurance value provided by health systems is similar for everyone.

We distribute the net saving of the government (the discrepancy between what the government collects in taxes and what it pays as transfers, consumption or interest)

⁹As a robustness check, we show in Appendix Figure A.7 the profile of redistribution in Europe and the US when making the polar assumptions that all government consumption is distributed as a lump-sum in Europe and proportionally in the US and find that even under these extreme assumptions, our main conclusions are unchanged, as we further discuss in section VI..

proportionally to the income of individuals so that posttax national income matches national income.

Unit of Analysis In our benchmark series, the statistical unit is the adult individual (defined as being 20 or older) and income is split equally among spouses.¹⁰

IV. SOURCES AND METHODOLOGY

This section describes the main steps followed to estimate the distribution of the national incomes of European countries. We refer to the appendix A for a more detailed and technical description of the methodology, and to the extended appendix for a detailed account of data sources country by country, methodological steps and robustness checks. In broad strokes, our methodology starts from a variety of household surveys. We harmonize them and correct them using tax data. Finally, we account for the various parts of national income that are absent from the usual sources.

IV.A. *National Accounts*

Main Aggregates For total national income, we use series compiled by the World Inequality Database based on data from national statistical institutes, macroeconomic tables from the United Nations System of National Accounts and other historical sources (see [Blanchet and Chancel, 2016](#)). For the various components of national income, we collect national accounts data from Eurostat, the

¹⁰Our notion of a “spouse” follows from that of the EU-SILC and includes any married people and partners in a consensual union (with or without a legal basis). We also compute additional series in which income is split between all adult household members, not just members of a couple (i.e. a “broad” rather than a “narrow” equal-split). The difference is not entirely negligible in certain Southern and Eastern European countries. Until now DINA studies have had a tendency to use the narrow equal-split in developed countries (e.g. [Garbinti, Goupille-Lebret, and Piketty, 2018](#); [Piketty, Saez, and Zucman, 2018](#)) and the broad equal-split in less developed ones (e.g. [Chancel and Piketty, 2019](#); [Novokmet, 2018](#); [Piketty, Yang, and Zucman, 2019](#)). We focus on the narrow equal-split in our benchmark for comparability with the United States, but also shed some light on this issue by providing both concepts – see figure A.2.7 in the extended appendix. We avoid using other, more complex equivalence scales because these creates complicated non-linearity issues that prevent “equalized income” from summing up to national income, and are therefore not well-suited to our problem.

OECD, and the UN. We use Eurostat and the OECD in priority, as they tend to have the most reliable data, but their coverage is less extensive than the UN.¹¹ We provide a detailed view of the coverage that these data provide in our extended appendix.¹²

Additional Sources In a few cases, we need to rely on additional sources to perform decompositions of the national income that are needed for our series and more precise than what is available through standard data portals. First, when we distribute the retained earnings of the corporate sector, we have to separate the share that is owned by private citizens from the share that belongs to the governments. To do so, we use the fractions of equities owned by the households sector in the financial balance sheets available from the OECD. We also need to separate the social benefits that correspond to pension and unemployment from the other types of social benefits in order to calculate pretax income.¹³ To do so, we rely on the OECD social expenditure database, which breaks down social benefits by function in great details since 1980. Finally, we need to separate health expenditures from the rest in the individual consumption of the government. For that, we extract health government spending from the System of Health Accounts, a database that emerged from a joint work between the OECD, Eurostat and the WHO.

IV.B. Survey Microdata

Sources We collect and harmonize household survey data from several international and country-specific datasets. Our most important source of survey data is the

¹¹We link together the various series, rescaling older and lower-quality series to match the newer and higher-quality ones in their latest year of overlap to avoid any structural break.

¹²Using these sources, we have a sufficiently detailed decomposition of national income that covers nearly 100% of the continent national income up until 1995. Before that, coverage becomes increasingly sparser: we have the full decomposition for about 50% of national income in 1990, decreasing to 20% at the very beginning of our series. We impute missing series by retropolating them using exponential smoothing with a coefficient of 0.9. As a last resort, we rely on regional averages.

¹³The DINA guidelines (Alvaredo et al., 2016) distinguish between social insurance benefits (D621 + D622 in the SNA) which individuals are entitled to if they have contributed and social assistance benefits in cash (D623), which individuals can receive without having contributed. Unfortunately that level of details is not commonly available in the national accounts of most countries, which only report the aggregate item D62. This is why we rely on alternative sources.

European Union Statistics on Income and Living Conditions (EU-SILC), which have been conducted on a yearly basis since 2004 in thirty-two countries. We complement EU-SILC by its predecessor, the European Community Household Panel (ECHP), which covers the 1994-2001 period for thirteen countries in Western Europe. Our second most important source of survey data is the Luxembourg Income Study (LIS) that provides access to harmonized survey microdata covering twenty-six countries since the 1970s. Most Western European countries are covered from 1985 until today, and several countries from Eastern Europe have been surveyed since the 1990s.

Imputations When we have access to survey microdata, we can usually estimate income concepts that are close to our concepts of interest (pretax and posttax income) with only a few components of income that remain to be added separately (see section IV.E.). A significant exception concerns social contributions in EU-SILC: while both employer and employee social contributions are recorded, employee contributions are combined with income and wealth taxes. We use the social contribution schedules published in the OECD Tax Database to impute employee social contributions separately. Before 2007, employer contributions may also not be recorded despite having information on income before taxes and employee contributions. In such cases, we also impute employer contributions based on schedules from the OECD Tax Database. Beside that, measures of income before and after taxes and transfers have been recorded consistently as part of EU-SILC. The Luxembourg Income Study also produces some historical data on pretax income, in many cases by imputing direct taxes and social contributions as part of their harmonization effort. As a result, we have survey microdata on both pretax income and posttax income in almost all countries since 2007, and for over a longer time period for a number of Western European countries (Germany, the United Kingdom, Switzerland, and Nordic countries).

IV.C. Survey Tabulations

Sources We complement the survey microdata with a number of tabulations available from the World Bank's PovcalNet portal, the World Income Inequality Database (WIID) and other sources. PovcalNet provides pre-calculated survey distributions

by percentile of posttax income or consumption per capita. The WIID gathers inequality estimates obtained from various studies, and gives information on the share of income received by each decile or quintile of the population. Finally, we collect historical survey data on posttax income inequality in former communist Eastern European countries provided by [Milanovic \(1998\)](#). In all cases, we use generalized Pareto interpolation ([Blanchet, Fournier, and Piketty, 2017](#)) to recover complete distributions from the tabulations. A detailed breakdown of available survey data sources by country is available in the extended appendix.

Harmonization Contrary to microdata, tabulations only provide distributions covering specific income concepts and equivalence scales. The majority of tabulations recorded in PovcalNet and WIID correspond to posttax income, while cases in which we only observe consumption are limited to a handful of Eastern European countries (Moldova, Kosovo, Montenegro).¹⁴ The equivalence scales available are more diverse, including households, adults, individuals, the OECD modified equivalence scale or the square root scale.¹⁵ For these data sources, as well as for survey microdata where information on taxes and transfers is incomplete, we have to develop a strategy to transform the distribution of the observed “source concept” (e.g. consumption per capita or posttax income among households) into an imputed distribution measured in a “target concept” (pretax or posttax income per adult).

The key idea behind our harmonization procedure is that, while the different income or consumption concepts that we observe are different, they are also related. Using all the cases where the income distribution is simultaneously observed for two different concepts, we can map the way they tend to relate to one another, and use that to convert any source concept to our concept of interest. In practice, we formalize this idea by writing the average income of each percentile for the

¹⁴The only exceptions correspond to a handful of Eastern European countries at the beginning of the period (Bosnia and Herzegovina, Moldova, Montenegro) for which we have no other source available. In these cases we use the survey distribution of pretax income as a proxy for the “true” pretax income.

¹⁵When computing inequality estimates with the OECD modified equivalence scale, the first adult in the household is given a weight of 1, other adults are given a weight of 0.5, and children are given a weight of 0.3 each. The square root scale divides total income by the square root of the size of the household.

distribution of interest as a function of all the percentiles of the distribution from which we wish to impute, and also as a function of various auxiliary variables that may potentially account for that relationship (average income, population and household structure, marginal tax rates, social expenditures). Finding that function amounts to a regression problem, albeit a high-dimensional, non-parametric one. To avoid making *ad hoc* restrictions, we rely on recent advances in non-parametric high dimensional statistics, also known as machine learning. We use XGBoost (Chen and Guestrin, 2016), a state-of-the-art implementation of a standard, robust and high performing algorithm called *boosted regression trees*. We provide a detailed view of the method and the results in section I.A. of appendix A. In particular, we show that this approach performs better than more naive ones, such as assuming a single correction coefficient by percentile, especially in countries where we have to resort to poor predictors of income, such as consumption.

Still, we stress that this approach is not perfect: the relationships between the different concepts are not deterministic, so that these imputations involve their share of uncertainty. However, the existing literature has often chosen to ignore these issue altogether, and directly combined, say, income and consumption data (e.g. Lakner and Milanovic, 2016). We feel that our approach is preferable, because it corrects at least for what can be corrected. Note that in the end, the output of the harmonization procedure is straightforward and intuitive: it mostly adjust the levels of the different series, but does not introduce any trend that was not already in the data. We show the impact of these corrections country by country in our extended appendix.

IV.D. Survey Corrections

Survey data are known to often miss the very rich. For our purpose it is important to distinguish two reasons for that: non-sampling and sampling error. Sampling error refers to problems that arise purely out of the limited sample size of survey data. Low sample sizes affect the variance of estimates, but they may also create biases, especially when measuring inequality at the top of the distribution. Non-sampling error refers to the systematic biases that affect survey estimates in a way that is not directly affected by the sample size. These mostly include people refusing to answer surveys and misreporting their income in ways that are not observed, and therefore

not corrected, by the survey producers. Estimates based on raw survey data do not account for any of these biases and therefore tend to underestimate incomes at the top end.

Non-sampling Error We correct survey data for non-sampling error using known top income shares estimated from administrative data. Following contributions by [Piketty \(2001\)](#) for France and [Piketty and Saez \(2003\)](#) for the United States, several authors have been using tax data to study top income inequality in the long run. Most of these studies have been published in two collective volumes ([Atkinson and Piketty, 2007; 2010](#)), and their results have been compiled in the World Inequality Database.¹⁶ In general, tax data is only reliable for the top of the distribution, and this is why these series do not cover anything below the top 10%. In that literature, researchers estimate the share of top income groups by dividing their income in the tax data by a corresponding measure of total income in the national accounts. At the time of writing, data series were available for nineteen European countries, providing information on the share of income received by various groups within the top 10%.

We complete this database by gathering and harmonizing a new collection of tabulated tax returns covering Austria (1980–2015), East Germany (1970–1988), Estonia (2002–2017), Iceland (1990–2016), Italy (2009–2016), Luxembourg (2010, 2012), Portugal (2005–2016), Romania (2013) and Serbia (2017). We use these tabulations to add new top income shares to our database. We provide a detailed account of the computations for each country in the extended appendix. In most cases, we directly correct the surveys with the tax data using the method of [Blanchet, Flores, and Morgan \(2018\)](#) rather than using a total income estimate from the national accounts. Direct correction of survey data is a more flexible and practical approach, at least for the recent period, and is now being preferred in the latest work on inequality (e.g. [Bukowski and Novokmet, 2017](#); [Morgan, 2017](#); [Piketty, Yang, and Zucman, 2019](#)). When extending existing series using that method, as in Italy or Portugal, our results are consistent with the work that was done previously, thus confirming the consistency and reliability of both approaches. Our results also reveal

¹⁶See <http://wid.world>.

that the underestimation of top incomes varies a lot across surveys and is typically higher in Eastern European countries. This points to the importance of correcting surveys with tax data to make comparisons between countries more reliable.

We correct the survey data using standard survey calibration methods. The principle of survey calibration is to reweight observations in the survey in the least distortive way so as to match some external information. Statistical institutes already routinely apply these methods to ensure survey representativity in terms of age or gender. We directly extend them to also ensure representativity in terms of income. The applicability of these methods to correct for the underrepresentation of the rich in surveys has been discussed at length by [Blanchet, Flores, and Morgan \(2018\)](#).

One difficulty is that our external source of information consists in top income shares. Because top income shares are a non-linear statistic, they cannot directly be used in standard calibration procedures. We tackle that issue using suggestions from [Lesage \(2009\)](#). They involve linearizing top income shares statistics by calculating their influence function, and introducing a nuisance parameter. We discuss that methodology in detail in appendix [A](#). In concrete terms, we increase weight at the top of the distribution so that survey top incomes match their value observed in the tax data.

One advantage of calibration procedures is that they allow to perform survey correction with a taxable income concept that may differ from the income concept of interest — either in terms of income definition or statistical unit. We always match concepts to the best of our ability between the tax data and the survey data to perform the correction. Then we use income concepts that are better defined and more economically meaningful to produce our inequality series. Confronting tax data and survey data as such is a very powerful way to make income tax statistics comparable between countries.¹⁷ It lets us account for top incomes while retaining the wealth of information included in the surveys, notably on taxes and transfers, so

¹⁷For older time periods from which we cannot perform that exercise directly due to lack of proper survey microdata, we retropolate the correction on the income tax series that is done over the more recent period. See extended appendix for details.

that we can study them and calculate both pretax and posttax incomes.¹⁸

When we do not directly observe tax data in a country, we still perform a correction based on the profile of nonresponse that we observe in other countries. This is only the case for a few small countries — Albania, Bosnia and Herzegovina, Bulgaria, Cyprus, Kosovo, Latvia, Lithuania, Macedonia, Malta, Moldova, Montenegro and Slovakia. To capture statistical regularities, we estimate the nonresponse profile as a function of the distribution of income in the uncorrected survey using the same machine learning algorithm as in section IV.C.. We stress that this remains a rough approximation and that in our view the proper estimation of top income inequality requires access to tax data. Fortunately, our tax data covers the majority of the European population, and an even larger share of European income, so that the impact of these corrections on our results is small. Indeed, as shown in figure A.2.11 in appendix, excluding these countries from the analysis has little impact on the results.

Sampling Error The sample size of surveys varies a lot and can sometimes be quite low: this, in itself, can seriously affect estimates of inequality at the top and, on average, will underestimate it (Taleb and Douady, 2015). Correcting sampling error requires some sort of statistical modeling. We borrow methods coming from extreme value theory, which is routinely used in actuarial sciences to estimate the probability of occurrence of very rare events, but can similarly be used to estimate the distribution of income at the very top.

The main tenet of extreme value theory is that a parametric family of distributions — the generalized Pareto distribution — more or less provides a universal approximation of distribution tails. It is a flexible model and includes the Pareto, the exponential, or the uniform distribution as a special case. We use it to model the top 10% of income distributions. We estimate the model using a simple and robust method known as probability-weighted moments (Hosking and Wallis, 1987). We

¹⁸Our ability to explore the composition of top incomes is still limited by the precision of the survey data. At the European level and over several years, the number of observations is sufficient to study the top 1%, but not much higher. When it comes to studying the marginal distribution of pretax and posttax incomes, however, statistical methods can address these sampling error issues (see below).

provide technical details for the method in appendix A. Note that by construction, this adjustment has absolutely no impact on the top 10% income share (which we know from the tax data), it only refines the income distribution within the top 10%.

IV.E. Incomes Traditionally Excluded from Tax and Survey Sources

Once we have harmonized and corrected our survey data using tax data, we find ourselves with more correct and comparable inequality series. But those series do not yet account for all of the national income because they lack some components from the household sector (imputed rents), the corporate sector (undistributed profits) and the government sector (taxes on products and government spending).¹⁹

Imputed Rents We extract the total value of imputed rents from the national accounts. To distribute them, we rely on (calibrated) EU-SILC data that does record imputed rents (although they are not included in the headline inequality figures). We perform a simple statistical matching procedure using income as a continuous variable to add imputed rents, which we describe in the appendix I.D.. The imputed rents total is rescaled to match national accounts. The method preserves the rank dependency (i.e. the copula) between income and imputed rents in EU-SILC, the distribution of imputed rents in EU-SILC, the distribution of income in the original data, and the imputed rents total in the national accounts.

Undistributed Profits We distribute the private share of undistributed profits to individuals proportionally to the ownership of corporate stock. This includes both private and public stocks that are held directly or indirectly through mutual funds and private pension plans. However, we exclude sole proprietorship, since in the national accounts they are not an entity separate from the household to which they belong.²⁰

¹⁹Some minor discrepancies remains (government surplus, FISIM, NPISH income, etc.). Given the size of these aggregates, they can only have minor impacts on the overall distribution. We distribute them proportionally. The DINA guidelines (Alvaredo et al., 2016) sometimes suggest slightly more refined imputation methods, but these are more difficult to apply given the nature of our data, and have only very limited impact on the results.

²⁰This income is part of “mixed income” and distributed as part of the rest of household sector income.

The amount of undistributed profits comes from the national accounts. We determine the share of these profits that is privately owned (rather than government-owned) using the financial balance sheets: we calculate the share of equity assets that is owned by the household sector out of the equity assets owned by both the household and the government sector. Government-owned undistributed profits are treated like a primary income of the government.

The distribution of stock ownership comes from the Household Finance and Consumption Survey (HFCS), the pan-European wealth survey of the European Central Bank. We calibrate that survey on the top income shares as we do for other surveys to make it representative in terms of income and to get consistent results. The HFCS only started around 2013, so before that year we keep the distribution of retained earnings constant and only change the amount of retained earnings to be distributed: this constitutes a reasonable approximation because stock ownership is always highly concentrated, so that the main impact of retained earnings on inequality comes from changes in their average amount rather than changes in the inequality of stock ownership. After 2013, we use the wave that is closest to the year under consideration.

The Corporate Income Tax Because the corporate income tax is paid out of corporations' profit, we distribute it similarly to undistributed profits in pretax income.

Taxes on Products and Production In our benchmark series we distribute taxes on products and production proportionally to pretax income, so we do not need any additional source. But we provide alternative series in which taxes on products are distributed proportionally to consumption (see figure A.2.4 in extended appendix). For that, we rely on the Household Budget Surveys (HBS) from Eurostat to get the distribution of consumption and its dependency to income. We use the same statistical matching procedure as before to attribute a consumption to people alongside the income distribution, and attribute the taxes on products proportionally to it.

In-kind Transfers We distribute public health spending in a lump-sum way to individuals. That is, we assume that healthcare that is free at the point of use provides an insurance value that is the same for everyone across the population. We distribute other types of in-kind transfers proportionally in our benchmark series, similarly to earlier DINA studies (Piketty, Saez, and Zucman, 2018).

IV.F. Validation of our Methodology

The Impact of the Different Methodological Steps Our estimates differ from existing, survey-based estimates for two groups of reasons: because we use tax data at the top of the distribution, and because we incorporate forms of income that are traditionally absent from inequality statistics. How do these elements impact our results? Figure 1a gives the answer.²¹ Based only on raw, survey-based estimates, we would conclude that inequality has been slightly going down in Europe after a one-time increase in the early 1990s: the top 10% income share has been stable after 1995, while the bottom 50% income share has been slightly but consistently on the rise. When using tax data to correct the top of the distribution, we get a fairly different picture: the increase of the top 10% income has been much more significant, while the share of the bottom 50% has been stable. Adding missing income components modifies the distribution of income further. Some income components (such as undistributed profits) have a strong unequalizing impact, while others (like imputed rents) have more equalizing tendencies. Overall, we distribute between one fifth and one quarter of national income in the form of additional income components. This leads to our “DINA” series, which show a slightly higher top 10% income share in recent years than survey and tax data alone. Most of the difference with raw survey estimates, however, comes from tax data.²²

[Figure 1 about here.]

Comparison with Earlier Works We wish to provide results that are conceptually similar to other works on distributional national accounts (Bozio et al., 2018;

²¹Our extended appendix shows the impact of the different steps country-by-country.

²²We provide a similar reconciliation of our results with official Eurostat data in figure A.2.2 of the extended appendix.

Garbinti, Goupille-Lebret, and Piketty, 2018; Piketty, Saez, and Zucman, 2018). But in practice our methodology is quite different. In France, Garbinti, Goupille-Lebret, and Piketty (2018) and Bozio et al. (2018) estimated the distribution of pretax national and posttax disposable income using detailed tax *microdata*, combined with various surveys, microsimulation models for taxes and benefits, rescaling income component by component to the national accounts. We, on the other hand, only use tax tabulations to correct survey data, and rescale our results to the national accounts at a coarser level. The advantage of our method is that it is applicable much more widely and rapidly, including in countries in which no tax data is available.

To what extent can our approach yield results that are comparable to more complex and detailed works? As figure Ib shows, we get results that are very similar to these earlier works in the case of France (see the extended appendix for detailed information on French data sources and methodological steps). Concretely, our methodological approach starts from the raw survey series shown on the bottom line, which suggest that the top 10% share has fluctuated between 22% and 26%. In a second step, we calibrate these distributions to the fiscal income shares measured from tax data, which reveal that this share has in fact remained approximately constant over time, reaching about 30%. In a third step, we impute additional sources of income, such as retained earnings and imputed rents. This gives us the DINA top 10% pretax income share, which follows closely the series estimated by Garbinti, Goupille-Lebret, and Piketty (2018). Finally, we impute all taxes, as well as cash transfers. This yields the distribution of posttax disposable income, which is also remarkably similar to that obtained by Bozio et al. (2018).²³ Note that we obtain these results in spite of the fact that our data sources for France are not of an especially high quality — and are also very different from the ones used in DINA

²³We show in figure A.2.10 in appendix that we also closely reproduce those studies' findings for the bottom 50% of the distribution.

projects.²⁴

V. THE DISTRIBUTION OF EUROPEAN NATIONAL INCOMES, 1980-2017

In this section, we show that pretax inequality has risen much less in Europe than in the US since 1980. This is true for most European countries taken separately but also for Europe taken as a whole — a block that is broadly similar in terms of population size and aggregate economic output as the US.²⁵

V.A. *The distribution of pretax income in 2017*

Table I shows the average incomes of various groups in Western and Northern Europe, Eastern Europe and the United States for comparison. Regional differences between Western Europe Northern Europe and Eastern Europe are mainly driven by differences among bottom income groups rather than by differences at the top of the distribution. The average income of Western and Northern Europeans is about 80% higher than the average income of Eastern Europeans. When comparing the incomes of the poorest 20% of adults in these regions, however, this gap exceeds 250%. Accordingly, table I reveals a remarkable degree of convergence in the top incomes of Western Europe, Northern Europe and Eastern Europe. The top 0.001% earns close to 9 million euros in Eastern Europe, compared to just above 10.5 million euros in Western and Northern Europe.

[Table 1 about here.]

²⁴The SILC statistics for France are a transcription of a survey (called SRCV) which is used for its extensive set of questions on material poverty, but is not considered the best survey for income inequality. For that purpose, the French statistical institute relies on another survey, called ERFS. But that survey is not part of any international scheme, such as EU-SILC, nor is it available through portals such as the Luxembourg Income Study. Therefore, we do not include it in our estimations. Before SILC is available, we rely on France's Household Budget Survey, which has been made available through LIS. While France's HBS is a key source for consumption data, it is not viewed as the best source for income data either. It is also separate from EU-SILC data, which explains the inconsistent trend. Therefore, there is no reason to think that our methodology would work better for France than other countries just because of the quality of the data in input.

²⁵In 2018, aggregate national income was €14.2tn in Europe and €13.1tn in the US. Total population was 432 million in Europe and 241 million in the US. In terms of population, Western and Northern Europe combined (334 million) are more comparable to the US than Europe as a whole. See Appendix Table A.6 for the composition of Europe and its subregions).

On the other hand, income differences between Europe and the United States are mostly driven by top incomes. The average national income per adult in the United States is 45% higher than in Western and Northern Europe, and more than 2.5 times higher than in Eastern Europe. Yet, the bottom 50% of the US population earns 15% less than the bottom 50% of Western and Northern Europe, and more than 40% less when looking at the bottom 20%. By contrast, the average pretax income of the top 1% exceeds one million euros in the United States, compared to €372,000 in Western and Northern Europe and €238,000 in Eastern Europe.

V.B. The rise of top incomes

[Figure 2 about here.]

Figure II shows the evolution of the pretax income distribution between 1980 and 2017 in Western Europe, Northern Europe, Eastern Europe and the US, with a further decomposition of groups within the top 1%. In all three European regions, growth has been markedly higher for the top 10% than for the bottom 90% of pretax income earners, and has been even higher at the very top of the distribution. The distribution of growth has been much more skewed in Eastern Europe than in the rest of the continent.²⁶ In the US, the bottom 50% did worse than its counterparts in all European regions.

V.C. The distribution of macroeconomic growth

[Table 2 about here.]

Table II shows the average annual real pretax income growth of selected income groups in Western Europe, Northern Europe, Eastern Europe and the United States over the 1980-2017 and 2007-2017 periods.

²⁶Let us stress here that we focus solely on monetary income inequality, which was unusually low in Russia and Eastern Europe under communism. Other forms of inequality prevalent at the time, in terms of access to public services or consumption of other forms of in-kind benefits, may have enabled local elites to enjoy higher standards of living than what their income levels suggest. That being said, the survey tabulations at our disposal do partially account for forms of in kind income, so this limitation should not be exaggerated (see [Milanovic, 1998](#)). Furthermore, the top 10% income share did continue to rise in many Eastern European countries after 1995 and is now higher than in other European regions.

While the long run picture reveals a clear increase in inequality, the period of stagnation that followed the 2007-2008 crisis has been less detrimental to the European middle class than to other income groups. In Western Europe and Northern Europe the average incomes increased for groups located at the middle of the distribution more than for any other groups. Eastern European countries were less affected by the crisis but experienced comparable evolutions: the bottom 20% grew at an annual rate of 1.6% between 2007 and 2017, lower than the regional average of 1.9%. Therefore, while the financial crisis has to some extent halted the rise in top income inequality in Europe, income gaps between the middle and the bottom of the distribution have continued to widen, and low incomes have consistently lagged behind the expansion of the overall economy.

The distribution of growth in the United States since 2007 has been unambiguously more skewed: the average national income grew faster than in Western and Northern Europe at an annual rate of 0.4%, but the pretax incomes of the bottom 50% dropped by 1.2% and that of the bottom 20% by as much as 2.9%.

V.D. The role of regional integration and spatial inequality

[Figure 3 about here.]

The top panel of figure III compares the levels and evolution of the top 1% and bottom 50% income shares in the US, Europe as a whole, and Western and Northern Europe between 1980 and 2017.²⁷ Income inequality was unambiguously larger in the US than in Europe in 2017, even after accounting for differences in average incomes between European countries. The share of regional income received by the top percentile was twice as high in the United States (21%) as in Western and

²⁷We choose to pool together Western and Northern Europe because of their similarities in terms of political institutions (most countries joined the EU by the mid-1990s), living standards and economic tax and transfer systems (as discussed in the next section). For Europe as a whole, as well as for Northern and Western Europe, we aggregate country-level distributions after converting average national incomes at market exchange rates euros, rather than at purchasing power parity. This approach is justified by the fact that PPP conversion factors exist for European countries but not for US states: it would be unclear why one would correct for spatial differences in the cost of living in the former case but not in the latter. Estimating the distribution of European-wide income at purchasing power parity slightly reduces European inequality levels, as well as the share of inequality explained by between-country income disparities, so it does not affect our main conclusions.

Northern Europe (10.5%) or even Europe at large (11%). Meanwhile, the income share of the bottom 50% reached 16% in Europe and 20% in Western and Northern Europe as compared to less than 12% in the US. This was not always the case: in 1980, the bottom 50% share was actually higher in the US than in Europe as a whole, amounting to 20% of the national income, and it was only two percentage points lower than in Western and Northern Europe.

A more detailed picture of the distribution of growth in Europe and the US is shown in the bottom panel of figure III, which plots the average annual income growth by percentile in the two regions between 1980 and 2017, with a further decomposition of the top percentile. Growth of the average national income per adult was slightly higher in the US than in Europe: it exceeded 1.35% per year in the former, compared to 1.14% in the latter. However, this average gap hides important differences across the distribution. The average pretax income of the top 0.001% grew at a rate of 3.25% in Europe as a whole and as much as 5.39% per year in the US.²⁸ Meanwhile, macro growth has benefited significantly more to low-income groups in Europe than in the US.

[Figure 4 about here.]

To what extent are these figures driven by differences in average incomes between US states and between European countries, rather than within states and within countries? A Theil decomposition of within-group and between-group inequality in Europe and the US is shown in figure IV. The Theil index has risen much more in the US than in the Europe, and this change has been entirely due to increases in inequality within US states. In 1980, the Theil index in the US was almost perfectly equal to that of Europe at large, amounting to about 0.45; in 2017, by contrast, it was higher than 1 in the US, whereas it did not exceed 0.6 in Europe. The overall Theil index and the Theil index of within-state inequality are almost indistinguishable in the US: within-state inequality explained 97% of overall US inequality in 1980, and as much as 99% in 2017. The share of inequality explained by the between-group component is larger in Europe, but it has decreased from about 30% in 1980 to 20% in 2017, due to the rise in income differences within countries. In other

²⁸Results for Western and Northern Europe are almost identical (see figure II).

words, average income convergence in Europe has become increasingly insufficient to reduce inequalities between European residents.

VI. THE IMPACT OF TAXES AND TRANSFERS ON INEQUALITY

Our findings suggest that taxes and transfers are not more redistributive in Europe than in the US. However, given the higher level of pretax inequality in the US, European countries remain more equal than the US after all taxes and transfers are taken into account. Defining *predistribution* as the entire set of policies and institutions impacting the distribution of pretax incomes, and *redistribution* the entire set of policies impacting the distribution of posttax incomes, we find that Europe is more equal than the US thanks to predistribution rather than redistribution, contrary to a widespread view.²⁹

VI.A. *The Structure and Distribution of Taxes*

We start from the aggregate value of taxes and transfers available in National accounts. Appendix table A.4 and section B present these values for the three European regions and for the US, with a break-down by type of taxes and transfers. We now seek to distribute all these taxes (including corporate, sales and value-added taxes) paid by the different income groups.

In what follows, we distinguish two ways to report total taxes paid by income groups: (i) non-contributory taxes paid as a share of pretax income (i.e. excluding social contributions that pay for social insurance schemes) and (ii) total taxes paid as a share of factor income. When presenting taxes paid as a share of pretax income, we remove contributory social contributions from the analysis, since our definition of pretax income is net of the operation of pension and unemployment insurance systems. This way to look at tax incidence is useful for international comparisons focusing on the entire support of the adult distribution (i.e. when looking at the working *and* non-working populations). Its downside is that it misses a share of payments that can legitimately be considered as taxes by individuals. In countries

²⁹See for instance OECD, 2008; 2011 who reach the opposite conclusion as we further discuss below. See also Hacker (2011) for a discussion of predistribution and redistribution.

that opted for public pension systems, contributory social contributions indeed represent fixed compulsory contributions made by individuals to the State or pension institutions. We address this issue by reporting *total* taxes paid by income groups as a share of factor income. By doing so and by narrowing down the analysis to the employed and working-age (20–64) population, the analysis remains consistent and cross-country comparisons meaningful.³⁰

[Figure 5 about here.]

Focusing on taxes paid as a share of pretax income, we find contrasting patterns. Income and wealth taxes are progressive, but less so in Eastern Europe than in Western and Northern Europe. Indirect taxes are regressive as they are paid proportionally to consumption. Overall, non-contributory taxes are progressive in Western and Northern Europe, while they are flat or even regressive in Eastern Europe. The pattern for the United States is similar to that of Western and Northern Europe, except that indirect taxes are less important, and the income tax more progressive, so that the overall taxation profile is more progressive than in Europe.

Looking at *total* taxes paid as a share of factor income, we observe a different picture. The average tax rate increases to about 40% in European regions, highlighting the importance of pension and unemployment insurance contributions. The tax profile becomes regressive at the top of the distribution in Eastern Europe, where the top 1% faces a tax rate below 30%. In Western and Northern Europe, total tax rates are marginally lower for bottom groups than for the P80P95 group, but the tax rate drops for the top 5% and particularly so among the top 1%. This is due to consumption taxes, and to the fact that contributory contributions are mostly based on labor incomes, which only represent a small fraction of factor incomes at the top of the distribution. Social contributions are less important in the United States, so that the overall profile is broadly flat, with an increase for the top 1%.³¹

³⁰See section IV. for a longer discussion. Alternative methods to report tax incidence are discussed in Bozio et al. (2018).

³¹Figure V does not present effective tax rates beyond the top 1%, because in Europe available survey data does not allow us to decompose top incomes into subcomponents with a sufficient level of precision.³² Country-level studies (see Bozio et al. (2018) for France, and Saez and Zucman (2019) for the United States) find that the total tax rate tends to further drop for top income groups.

VI.B. The structure and distribution of transfers

[Figure 6 about here.]

Figure VI presents the distribution of transfers across income groups in the three regions, including in-kind transfers and collective expenditures. Unsurprisingly, the distribution of transfers is progressive in all regions. We also note that health payments are the most progressive type of transfers. This is also true in the United States, whose public health spending is large and more targeted towards the very poor (via Medicaid) than in comparable European countries.

Assumptions made on the allocation of collective consumption expenditure can potentially have a large impact on posttax inequality levels across countries. As discussed in the methodology section, our benchmark series allocate government expenditures in a proportional way, with the exception of health expenditures which are allocated in a lump sum way. As a robustness check, we make polar assumptions that all government consumption is distributed lump-sum in Europe and proportionally in the US. Even under these extreme and highly implausible assumptions, we find that redistribution is not dramatically and unambiguously more progressive in Europe than in the US (see Appendix Figure A.7).

VI.C. The net impact of taxes and transfers on inequality

[Figure 7 about here.]

The first panel of figure VII presents the net impact of redistribution measured by the difference between posttax incomes and pretax incomes.³³ The bottom 80% turns out to be a net beneficiary of redistribution in all regions. The impact of redistribution on pretax incomes among the bottom 50% is slightly lower in Eastern Europe than in Western and Northern Europe (the bottom 50% gain on average 25% vs. 30%, respectively) and the figure in both regions is well under the US (nearly 55%). The net effect of taxes and transfers on the top 10% is relatively similar in the three regions (-15% to -19%).

³³We stress that in this representation, pretax incomes are ranked along the distribution of pretax incomes and posttax incomes are ranked along the distribution of posttax incomes.

Is the US tax and transfer system more progressive than European systems? We compare our redistribution incidence curves of Panel A with the actual amounts of national income shifted between different income groups (Panel B). The net share of national income transferred to the bottom 50% is lowest in Eastern Europe (4.5%). In Western and Northern Europe and in the US, this value reaches around 6% of national income. Western and Northern Europe redistribute the same amount to the poorest half of their population, meaning that the US' apparently generous redistribution profile of Panel A can be largely explained by the low pretax incomes of the bottom half of the population. When breaking down net redistribution according to smaller groups, we find that Europeans are indeed more generous at the very bottom of the distribution: the bottom 20% receives 2.6% of national income in net redistribution compared to only 1.8% in the US³⁴. At the top of the distribution, a symmetrical effect is at play: the net share of national income flowing out of the top 10% group is significantly larger in the US than in Europe given that the group's average income is significantly higher than in Europe.

[Figure 8 about here.]

What does redistribution implies for pretax and posttax incomes? Figure VIII plots these incomes in Europe and in the US over time. We restrict the analysis to Western and Northern Europe and the US because of the relative degree of homogeneity of these regions in terms of economic development. Average incomes grew at a moderate rate in both regions over the 1980-2017 period: 49% in Western and Northern Europe and 66% in the US, corresponding to a 1% and 1.4% average annual increase.

Both pretax and posttax disposable incomes among the bottom 50% evolved very differently in Europe and in the United States: they grew by about 30% over the period in Europe, while they remain mostly stagnant in the US. It is only when looking at posttax national income that we observe an increase in the incomes of the

³⁴See Extended Appendix Figure A2.8.

bottom 50% in both regions.³⁵

When comparing the importance of pretax income inequality and posttax income inequality in European countries and the US, landmark publications on inequality (see for example [OECD, 2008; 2011](#), referred to below as the "standard" view) have found that Scandinavian countries redistribute in general more than other Western European countries, which in turn redistribute more than the US.³⁶ The policy conclusions of these findings are relatively clear: by increasing redistribution, high income inequality countries are likely to distribute incomes more equally. While our research does not contradict this general conclusion (more progressive taxes and more progressive transfers reduce disposable income inequalities for a given pretax distribution), our new results shed light on the importance of predistribution mechanisms in explaining Europe's relatively low inequality levels.

Why do our conclusions contradict the standard view on redistribution in Europe and the US? We find that the reasons for such differences are threefold. First, the standard approach used to compare redistribution across countries is based on household survey data rather than on distributional national accounts. While distributing national accounts poses evident conceptual and methodological challenges discussed in this paper, we argue that our new dataset allows much more meaningful international comparisons of redistribution. First, a comprehensive analysis of redistribution must take into account all forms of transfers and hence all forms of taxes, which is not possible using income concepts available in household surveys alone. Second, we use more precise administrative data to capture the top of the distribution. Third, the choice of income concepts and indicators used to compare redistribution between countries is crucial to make meaningful comparisons.

³⁵We also compare our European DINA series with Auten and Splinter ([Auten and Splinter, 2019a](#)) (AS) inequalities series for the US (see Appendix Figure [A.1](#)). Our findings on redistribution hold when we compare our series with AS US series. Indeed, AS US pretax inequality levels are lower than DINA US levels (in 2014, the top 1% is about 6 percentage points lower in DINA than in AS), but AS US posttax inequality levels are also lower than US DINA (in 2014, the top 1% is about 8 percentage points lower in AS). It follows that the US also appears to have a more progressive system of taxes and transfers than Europe when using alternative AS US series.

³⁶In [OECD, 2011](#), redistribution as measured by the difference between market Gini and disposable income Gini is found to be 18% in the US vs. 40% in Sweden and 33% in Norway (p. 270). Similar findings are obtained in more recent OECD publications such as [Causa and Norlem Hermansen, 2017](#).

Figure IX presents redistribution in the US and Northern European countries along different types of income concepts and datasets. In Panel A, we reproduce, based on a subset of our data, what we referred above as the “standard” view on redistribution in Europe and the US: that European countries achieve a much more equal distribution of disposable incomes thanks to taxes and transfers. Income concepts mobilized in this panel are from household surveys and do not incorporate components from tax or national accounts data. Factor income inequality appears to be relatively similar in both regions. One of the reasons for this is the use of survey data, which underestimates inequality more in the US than in Europe.³⁷ We also find that redistribution from factor income to posttax income looks particularly high in Europe. As described in section IV., this is because a large share of the elderly population has no factor income in public pension systems, which artificially inflates the gap between factor and posttax inequality. Panel A thus informs more on the relative size of public pension systems than on the level of tax-and-transfer redistribution. While some pension systems have a redistributive component, it remains very limited in most European countries (OECD, 2008): pensions can essentially be seen as replacement income.

[Figure 9 about here.]

Panel B plots factor and posttax inequality using DINA data. Factor income inequality is in effect much larger in the US than in Europe, once top incomes are more precisely measured thanks to tax data. Panel C limits the analysis to the working age population, in order to avoid including retired individuals with near-zero factor income. Once pensions are controlled for, Europe’s redistribution level is substantially reduced, but European countries remain much more equal than the

³⁷In Northern Europe our series actually show a slightly lower level of inequality for factor income than survey-based estimates. This is due to the equalizing effect of imputed rents, which typically constitute income to retirees with few other sources of factor income.

US after taxes and transfers.³⁸ Overall, we find that “predistribution” accounts for about two-thirds of the current Europe-US inequality gap when pension systems are treated as redistribution (figure IX, panel B) and around 90% of the inequality gap when pensions are not treated as redistribution (figure IX, panel C).³⁹

European countries therefore have low inequality levels today not because of the net effect of taxes and transfers, but primarily because they manage to distribute pretax incomes more equally than the US. To what extent can these differences in pretax inequality levels be due to trade and technology? The two regions have been exposed in the same way to imports of goods from low-income and emerging countries since the 1980s: from about 1.5% of GDP in the late 1980s to around 7% today (World Bank, 2019). It is unlikely that the rate of penetration of new technologies played a huge role in the US-Europe inequality divergence as well, since in most industrial sectors, robot penetration appears to be lower in the US than in Western European countries (Acemoglu and Restrepo, 2017).

What policies and institutions could have impacted the distribution of pretax incomes in such a different way in the Europe and the US since 1980? The range of policies which contribute to the formation and distribution of labor and capital incomes before taxes and transfers is wide, from education and health to minimum wages and corporate governance regulations, and it goes beyond the scope of this paper to identify how each of these contributed to the remarkable pretax inequality divergence we have documented.

³⁸We also stress that unemployment insurance (UI) payments play little role in explaining Europe’s relatively low inequality levels. Our definition of pretax income treats UI benefits and contributions as a replacement income, like pensions. Therefore, low pretax inequality levels could potentially mask high levels factor income inequality, before the operation of UI systems. In fact, figure IX (panel C) shows that UI systems play a minor role in explaining Europe’s low inequality levels. We also present bottom 50% income shares with and without UI benefits in Europe and the US, from 2007 to 2015, period during which unemployment rose sharply in most Western European countries. We find that UI systems play little role in the Europe-US pretax income inequality divergence observed (see extended appendix figure A2.9).

³⁹The percentage of the posttax inequality gap that can be accounted for by pretax inequality ($(\Delta Gini_{pretax}^{US-EU} / \Delta Gini_{posttax}^{US-EU})$) is equal to 69.7% for Northern Europe (cf. Figure IX Panel B.) when pensions are treated as a form of redistribution. The value is 63% for Western Europe. Focusing on the working age population only, the percentage of the posttax inequality gap between Northern Europe and the US that can be accounted for by pretax inequality is 91% (cf. Panel C.). The value is 87% when comparing US with Western Europe.

We nevertheless stress that several policies can both enter in the formation of posttax and pretax incomes and in that respect, the opposition between “predistribution” and “redistribution should not be overemphasized. Indeed, in the Distributional National Accounts framework, health payments enter in the formation of posttax disposable and posttax national incomes. The way health systems are organized across countries can therefore have notable effects on posttax income inequality, but differences in the organization of health systems across countries also impact the distribution of pretax incomes significantly.⁴⁰ Similarly, educational transfers appear in posttax incomes in DINA, but educational systems may be one of the most important determinants of pretax income inequalities. Available evidence suggests that universal access to higher education systems (as is the case of most Western and Northern European countries, in contrast to the US) tends to be associated with lower inequalities in access to education (Martins et al., 2010, Chetty et al., 2017), which can in return shape pretax income inequality.⁴¹ Finally, tax progressivity, which directly impacts disposable and posttax income inequality, can indirectly impact the formation of pretax incomes by slowing down the process of capital accumulation or possibly discouraging the bargaining for top compensations (Piketty, Saez, and Stantcheva, 2014). Countries which faced a large reduction in tax progressivity tend to be the ones which faced the highest increase in top pretax income shares (Alvaredo et al., 2018). The Distributional National Accounts developed in this study can help improve our understanding of these important questions in the future.

⁴⁰One of the most salient differences between the US and Western and Northern European health systems is that, for a broadly similar level of government spending on healthcare, European countries are characterized by public universal access, which tends to limit inequalities in access to healthcare, with better health results overall. Case and Deaton (2015) find that, after a historical decline, morbidity rates among white men have increased in the US since the late 1990s, contrary to European countries. Poor health is associated with reduced capabilities for the worse-off, as well as lower pretax incomes and social mobility (Case, Lubotsky, and Paxson, 2002).

⁴¹In the US, the share of private expenditure on tertiary educational institutions is over 65%, and this value is around 60% in other Anglo-Saxon countries, compared to 30% in France, Spain or Italy and as low as 8% in Germany and Scandinavian countries (Piketty, 2020).

VII. CONCLUSION

We have developed a novel methodology which combines virtually all available surveys, tax data and national accounts in 38 European countries in a systematic manner and which distributes the totality of national income before and after taxes. We use our methodology to compare inequality, growth and redistribution in Europe and in the US over the 1980-2017 period.

Our results show that pretax and posttax income inequality increased in almost all European countries since 1980. While this rise is widespread, we do observe heterogeneity in its intensity and timing across countries. But despite a generalized rise of inequalities in Europe over the past decades, European countries have been much more successful at containing the rise of income disparities than the US. Taxes and transfers contributed to limit the rise of inequality in European countries since 1980, but we do not find evidence that taxes and transfers are more progressive in Europe than in the US.

In future research, our data could also be used to test the distributional impact of changes in taxes and transfers on inequality; it could for instance be used to simulate the effect of the US adopting the tax codes of specific European countries on the distribution of pretax and posttax incomes.

WORLD INEQUALITY LAB, PARIS SCHOOL OF ECONOMICS

WORLD INEQUALITY LAB, PARIS SCHOOL OF ECONOMICS AND IDDRI-SCIENCESPO

WORLD INEQUALITY LAB, PARIS SCHOOL OF ECONOMICS

REFERENCES

- Acemoglu, Daron and Pascual Restrepo (2017). “Robots and jobs: Evidence from US labor markets”. In: *NBER working paper w23285*.
- Alstadsæter, Annette, Martin Jacob, Wojciech Kopczuk, and Kjetil Telle (2017). “Accounting for Business Income in Measuring Top Income Shares: Integrated Accrual Approach Using Individual and Firm Data from Norway”.

- Alvaredo, Facundo, Anthony B. Atkinson, Lucas Chancel, Thomas Piketty, Emmanuel Saez, and Gabriel Zucman (2016). “Distributional National Accounts (DINA) Guidelines: Concepts and Methods used in WID.world”. In: *WID.working paper series 2016/1*.
- Alvaredo, Facundo, Lucas Chancel, Thomas Piketty, Emmanuel Saez, and Gabriel Zucman (2018). *World inequality report 2018*. Belknap Press.
- Angel, Stefan, Richard Heuberger, and Nadja Lamei (2017). “Differences Between Household Income from Surveys and Registers and How These Affect the Poverty Headcount: Evidence from the Austrian SILC”. In: *Social Indicators Research* 138.2, pp. 1–29.
- Atkinson, Anthony B. (1996). “Income distribution in Europe and the United States”. In: *Oxford Review of Economic Policy* 12.1.
- Atkinson, Anthony B. and Thomas Piketty (2007). *Top Incomes Over the Twentieth Century*. Oxford University Press.
- (2010). *Top Incomes - A Global Perspective*. Oxford University Press.
- Atria, Jorge, Ignacio Flores, Ricardo Mayer, and Claudia Sanhueza (2018). “Top incomes in Chile: A Historical Perspective of Income Inequality (1964-2015)”. Australian Bureau of Statistics (2019). “Australian National Accounts: Distribution of Household Income, Consumption and Wealth, 2003-04 to 2017-18”. In: Auten, Gerald and David Splinter (2019a). “Income inequality in the United States: Using tax data to measure long-term trends”. In: *Working Paper*.
- (2019b). “Top 1 Percent Income Shares: Comparing Estimates Using Tax Data”. In: *AEA Papers and Proceedings*. Vol. 109, pp. 307–11.
- Beblo, Miriam and Thomas Knaus (2001). “Measuring income inequality in Euro-land”. In: *Review of Income and Wealth* 47.3.
- Blanchet, Thomas and Lucas Chancel (2016). “National Accounts Series Methodology”.
- Blanchet, Thomas, Ignacio Flores, and Marc Morgan (2018). “The Weight of the Rich: Improving Surveys Using Tax Data”.
- Blanchet, Thomas, Juliette Fournier, and Thomas Piketty (2017). “Generalized Pareto Curves: Theory and Applications”.

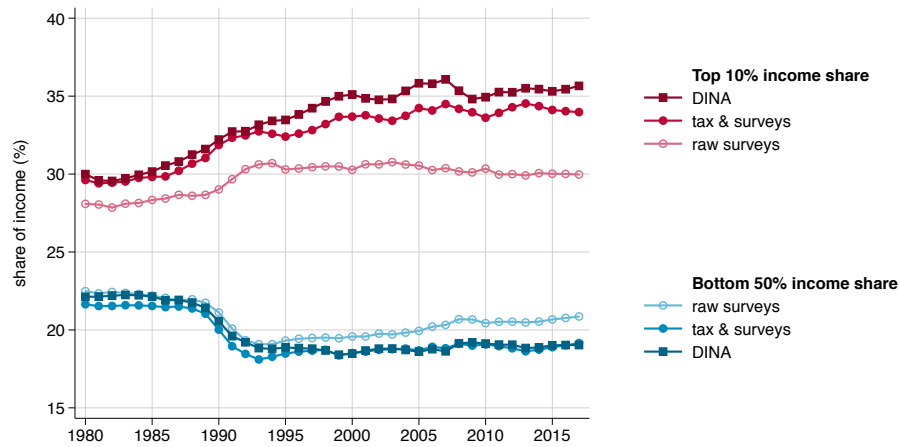
- Bozio, Antoine, Bertrand Garbinti, Jonathan Goupille-Lebret, Malka Guillot, and Thomas Piketty (2018). “Inequality and Redistribution in France, 1990-2018: Evidence from Post-Tax Distributional National Accounts (DINA)”. In: *WID.world Working Paper 2018/10*.
- Brandolini, Andrea (2006). “Measurement of income distribution in supranational entities: The case of the European Union”. In: *LIS Working Paper Series, No. 452*.
- Bukowski, Pawel and Filip Novokmet (2017). “Top incomes during wars, communism and capitalism: Poland 1892-2015”. In: *WID.world Working Paper Series 2017/22*.
- (2019). “Between Communism and Capitalism: Long-Term Inequality in Poland, 1892-2015”.
- Case, Anne and Angus Deaton (2015). “Rising morbidity and mortality in midlife among white non-Hispanic Americans in the 21st century”. In: *Proceedings of the National Academy of Sciences* 112.49, pp. 15078–15083.
- Case, Anne, Darren Lubotsky, and Christina Paxson (2002). “Economic status and health in childhood: The origins of the gradient”. In: *American Economic Review* 92.5, pp. 1308–1334.
- Causa, Orsetta and Mikkel Norlem Hermansen (2017). *Income redistribution through taxes and transfers across OECD countries*. Tech. rep. OECD Economics Department Working Papers, No. 1453.
- Chancel, Lucas and Thomas Piketty (2019). “Indian Income Inequality, 1922-2015: From British Raj to Billionaire Raj?” In: *Review of Income and Wealth* 65.S1, S33–S62.
- Chen, Tianqi and Carlos Guestrin (2016). “XGBoost: A Scalable Tree Boosting System”.
- Chetty, Raj, John N Friedman, Emmanuel Saez, Nicholas Turner, and Danny Yagan (2017). *Mobility report cards: The role of colleges in intergenerational mobility*. Tech. rep. national bureau of economic research.
- Cristia, Julian and Jonathan A. Schwabish (2009). “Measurement error in the SIPP: Evidence from administrative matched records”. In: *Journal of Economic and Social Measurement* 34.1, pp. 1–17.

- Dauderstädt, Michael and Cem Kelttek (2011). “Immeasurable Inequality in the European Union”. In: *Intereconomics* 46.1, pp. 44–51.
- Eurostat (2018). “Income comparison: social surveys and national accounts”. In: Fairfield, Tasha and Michel Jorratt De Luis (2016). “Top Income Shares, Business Profits, and Effective Tax Rates in Contemporary Chile”. In: *Review of Income and Wealth* 62.August, S120–S144.
- Fesseau, Maryse and Maria Liviana Mattonetti (2013). “Distributional Measures Across Household Groups in a National Accounts Framework: Results from an Experimental Cross-country Exercise on Household Income, Consumption and Saving”. In: *OECD Statistics Working Papers* 2013/4.
- Filauro, Stefano and Zachary Parolin (2018). “Unequal unions? A comparative decomposition of income inequality in the European Union and United States”. In: *Journal of European Social Policy*.
- Garbinti, B., J. Goupille-Lebret, and T. Piketty (2018). “Income inequality in France, 1900-2014: Evidence from Distributional National Accounts (DINA)”. In: *Journal of Public Economics* 162.1, pp. 63–77.
- Germain, et al. (2020). “Rapport sur les comptes nationaux distributionnels”. In: *INSEE Rapport*.
- Guillaud, Elvire, Matthew Olckers, and Michaël Zemmour (2019). “Four Levers of Redistribution: The Impact of Tax and Transfer Systems on Inequality Reduction”. In: *Review of Income and Wealth*.
- Hacker, Jacob (2011). “The institutional foundations of middle-class democracy”. In: *Essay*.
- Hosking, J R M and J R Wallis (1987). “Parameter and Quantile Estimation for the Generalized Pareto Distribution”. In: *Technometrics* 29.3, pp. 339–349.
- Immervoll, Herwig and Linda Richardson (2011). “Redistribution Policy and Inequality Reduction in OECD Countries: What has changed in two decades?” In: *OECD Social, Employment and Migration Working Papers No. 122*.
- Jesuit, David K. and Vincent A. Mahler (2010). “Comparing Government Redistribution Across Countries: The Problem of Second-Order Effects”. In: *Social Science Quarterly* 91.5, pp. 1390–1404.

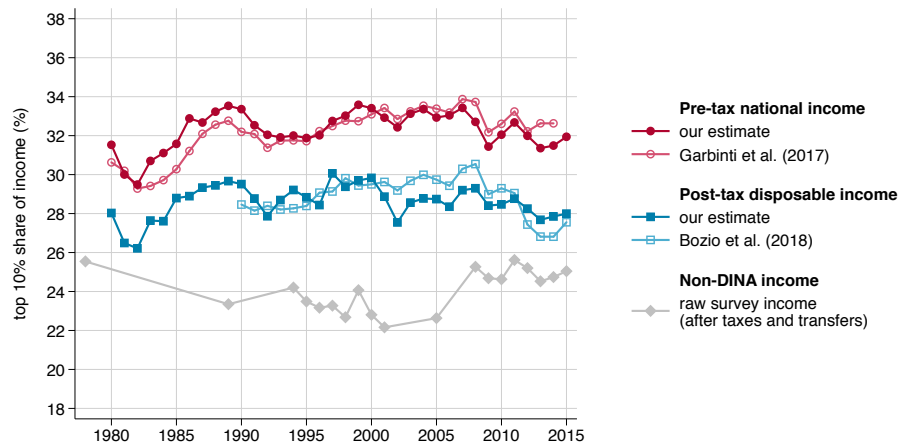
- Korinek, Anton, Johan A. Mistiaen, and Martin Ravallion (2006). “Survey nonresponse and the distribution of income”. In: *Journal of Economic Inequality* 4.1, pp. 33–55.
- Lakner, Christoph and Branko Milanovic (2016). “Global Income Distribution: From the Fall of the Berlin Wall to the Great Recession”. In: *The World Bank Economic Review* 30.2, pp. 203–232.
- Lesage, Éric (2009). “Calage non linéaire”.
- Martins, Joaquim Oliveira, Romina Boarini, Hubert Strauss, and Christine De La Maisonneuve (2010). “The policy determinants of investment in tertiary education”. In: *OECD Journal: Economic Studies* 2009.1, pp. 1–37.
- Milanovic, Branko (1998). *Income, Inequality, and Poverty during the Transition from Planned to Market Economy*. World Bank Regional and Sectoral Studies.
- Morgan, Marc (2017). “Extreme and Persistent Inequality: New Evidence for Brazil Combining National Accounts , Surveys and Fiscal Data, 2001-2015”.
- Novokmet, Filip (2018). “The long-run evolution of inequality in the Czech lands, 1898-2015”. In: *WID.world Working Paper Series* 2018/6.
- OECD (2008). *Growing unequal: Income Distribution and Poverty in OECD countries*. OECD Publishing.
- (2011). *Divided we stand: Why inequality keeps rising?* OECD Publishing.
- Paulus, Alari (2015). “Income underreporting based on income expenditure gaps: Survey vs tax records”.
- Piketty, Thomas (2001). *Les hauts revenus en France au XXe siècle*. Seuil.
- (2003). “Income inequality in France, 1901-1998”. In: *Journal of Political Economy* 111.5, pp. 1004–1042.
- (2020). *Capital and Ideology*. Harvard University Press Cambridge.
- Piketty, Thomas and Emmanuel Saez (2003). “Income inequality in the United States, 1913-1998”. In: *Quarterly Journal of Economics* 118.1, pp. 1–39.
- Piketty, Thomas, Emmanuel Saez, and Stefanie Stantcheva (2014). “Optimal taxation of top labor incomes: A tale of three elasticities”. In: *American economic journal: economic policy* 6.1, pp. 230–71.

- Piketty, Thomas, Emmanuel Saez, and Gabriel Zucman (2018). “Distributional National Accounts: Methods and Estimates for the United States”. In: *Quarterly Journal of Economics* 133.2, pp. 553–609.
- (2019). “Simplified Distributional National Accounts”. In: *AEA Papers and Proceedings* 109, pp. 289–295.
- Piketty, Thomas, Li Yang, and Gabriel Zucman (2019). “Capital Accumulation, Private Property, and Rising Inequality in China, 1978-2015”. In: *American Economic Review* 109.7, pp. 2469–2496.
- Saez, Emmanuel and Gabriel Zucman (2019). *The Triumph of Injustice: How the Rich Dodge Taxes and How to Make Them Pay*. W. W. Norton & Company.
- Salverda, Wiemer (2017). “The European Union’s income inequalities”.
- Statistics Canada (2019). “Distributions of Household Economic Accounts, estimates of asset, liability and net worth distributions, 2010 to 2018, technical methodology and quality report”. In:
- Statistics Netherlands (2014). “Measuring Inequalities in the Dutch Household Sector”.
- Taleb, Nassim Nicholas and Raphael Douady (2015). “On the super-additivity and estimation biases of quantile contributions”. In: *Physica A: Statistical Mechanics and its Applications* 429, pp. 252–260.
- Wolfson, Michael, Mike Veall, and Neil Brooks (2016). “Piercing the Veil – Private Corporations and the Income of the Affluent”. In: *Canadian Tax Journal* 64.1, pp. 1–25.
- World Bank (2019). “World Bank country by country trade database”.
- Zwijnenburg, Jorrit, Sophie Bournot, and Federico Giovannelli (2019). “OECD Expert Group on Disparities in a National Accounts Framework: Results from the 2015 Exercise.” In: *OECD Statistics Working Papers 2019/76*.

(a) Pretax Income Inequality in All 38 European Countries, 1980–2017



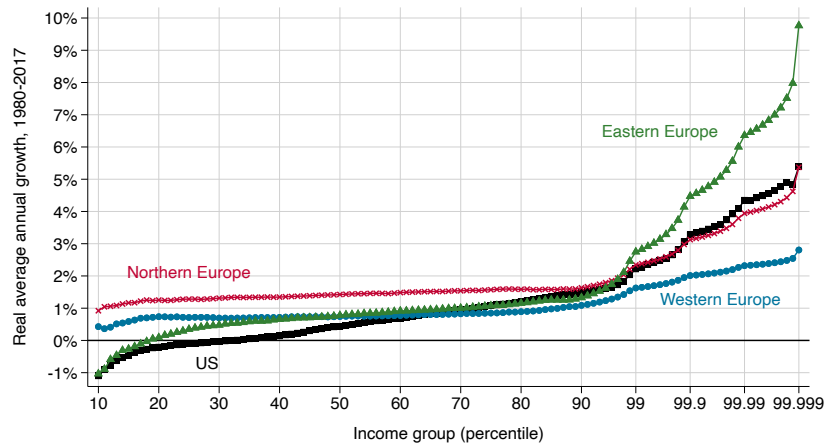
(b) Top 10% income shares in France, 1978-2015



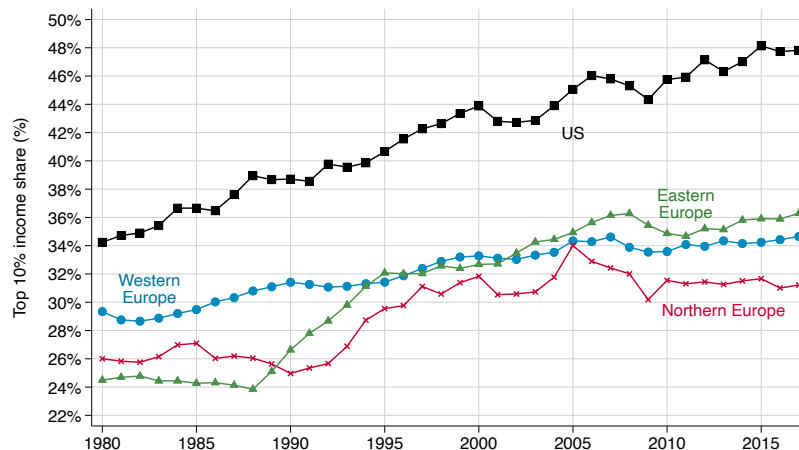
Source: Authors' computations combining surveys, tax data and national accounts. *Note:* Incomes measured at purchasing power parity. The unit of observation is the adult individual aged 20 or above. Income is split equally among spouses, except for the "raw survey income" series in the bottom panel for which income is split equally among all adult household members. See appendix table A.6 for country-by-country data sources.

Figure I
Validation of our Methodology

(a) Average annual income growth by percentile, 1980-2017



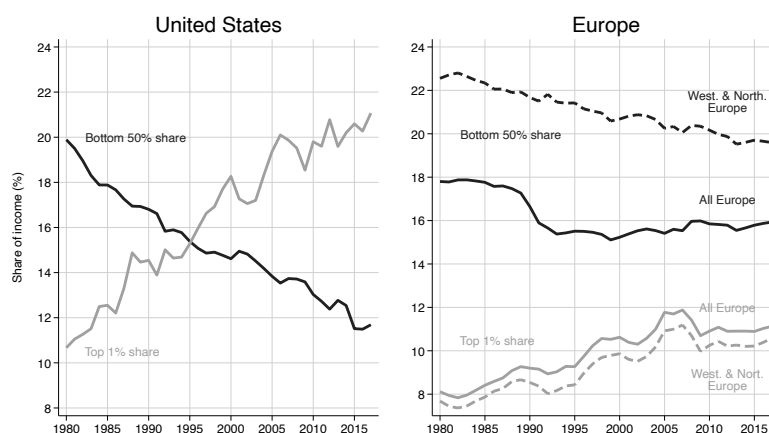
(b) Top 10% pre-tax income shares, 1980-2017



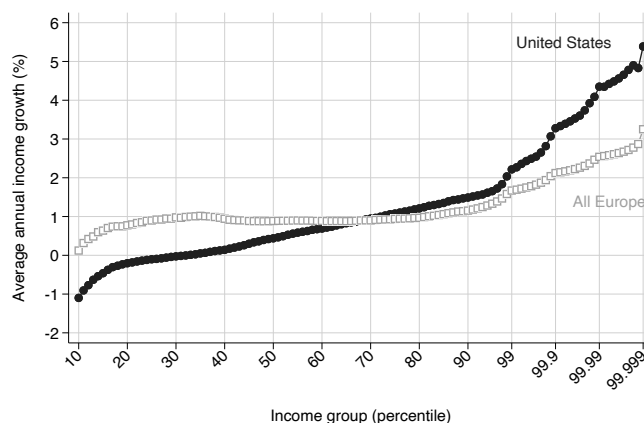
Source: Authors' computations combining surveys, tax data and national accounts. *Notes:* The top panel shows the average annual growth rate of pretax national income by percentile in Western Europe, Northern Europe and Eastern Europe, with a further decomposition of the top percentile. The bottom panel plots the share of regional income received by the top 10% between 1980 and 2017. The unit of observation is the adult individual aged 20 or above. Income is split equally among spouses. See Appendix Table A.6 for the composition of European regions.

Figure II
The rise of top incomes in Europe and the US, 1980-2017

(a) Top 1% and Bottom 50% pretax income shares in Europe and the US



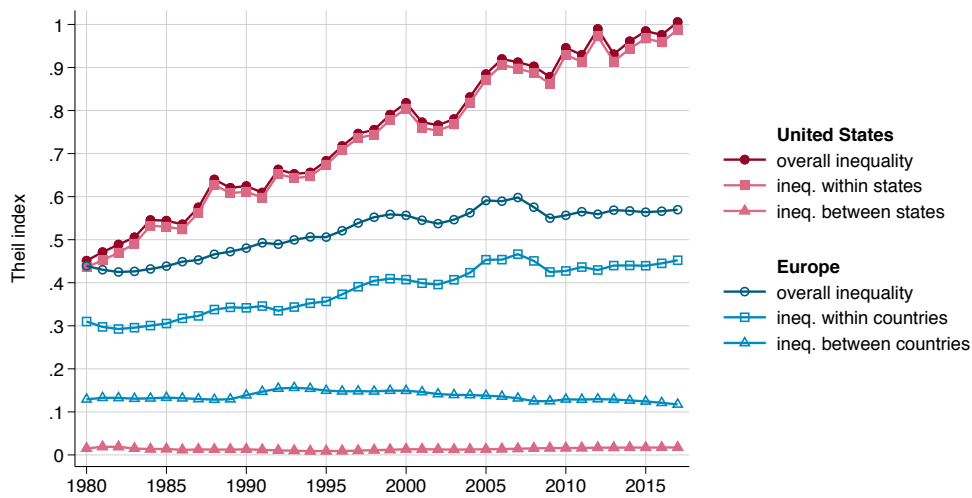
(b) Average annual pretax income growth by percentile, 1980-2017



Source: Authors' computations combining surveys, tax data and national accounts. *Notes:* The figure compares the evolution of pretax income inequality in Europe and the United States between 1980 and 2017. The top panel compares the share of pretax income received by the bottom 50% to that received by the top 1% of the regional population. The bottom panel plots the average annual pretax income growth rate by percentile, with a further decomposition of the top percentile. Figures for the US come from [Piketty, Saez, and Zucman \(2018\)](#). Figures for Europe correspond to Europe at large, that is after accounting for differences in average national incomes between European countries, measured at market exchange rates. The same holds for Western and Northern Europe. The unit of observation is the adult individual aged 20 or above. Income is split equally among spouses. See Appendix Table [A.6](#) for the composition of European regions. Appendix Figure [A.5](#) shows the distribution of posttax incomes in the two regions.

Figure III

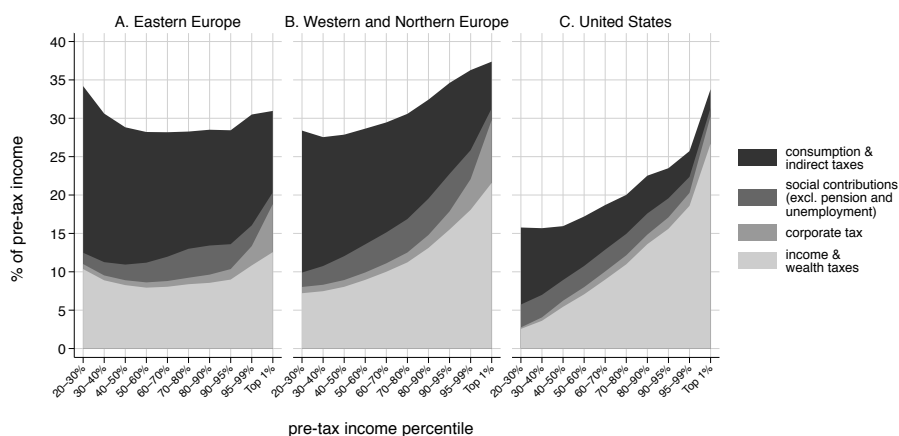
The distribution of pretax income growth in Europe and the United States, 1980-2017



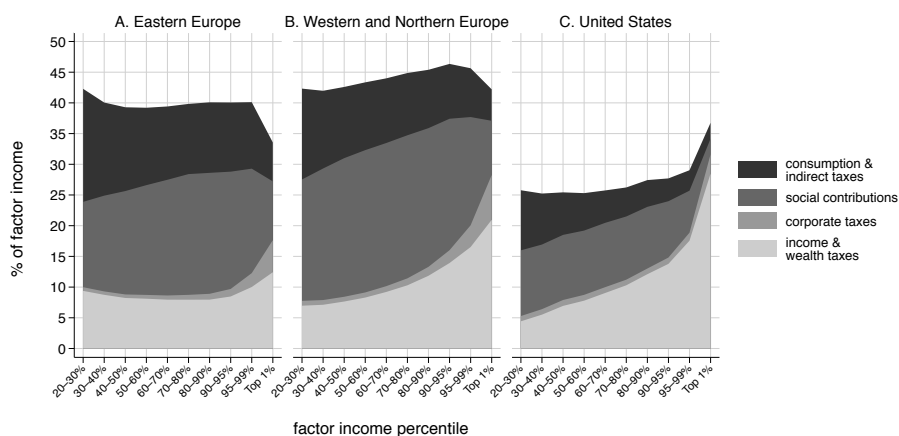
Source: Authors' computations combining surveys, tax data and national accounts for European countries. Figures for the US come from [Piketty, Saez, and Zucman \(2018\)](#) for the overall Theil index, and from state GDP estimates of the Bureau of Economic Analysis for the US between-group component. See appendix A for details. *Notes:* Figures for Europe correspond to Europe at large, that is after accounting for differences in average national incomes between European countries, measured at market exchange rates. The income concept is pretax income. The unit of observation is the adult individual aged 20 or above. Income is split equally among spouses. See Appendix Table A.6 for the composition of Europe.

Figure IV
Pretax income inequality in Europe and the United States, 1980-2017: Theil decomposition

(a) Non-contributory taxes paid as a share of pretax income

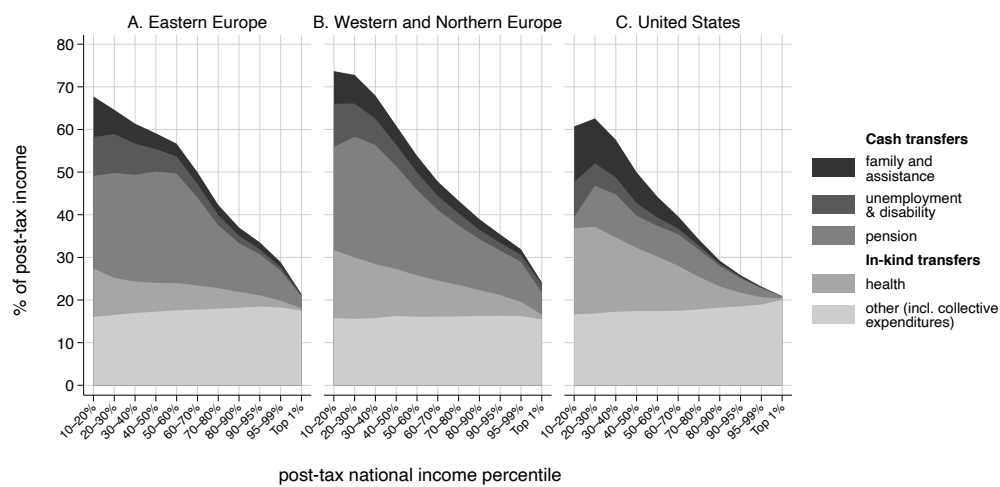


(b) Total taxes paid as a share of factor income



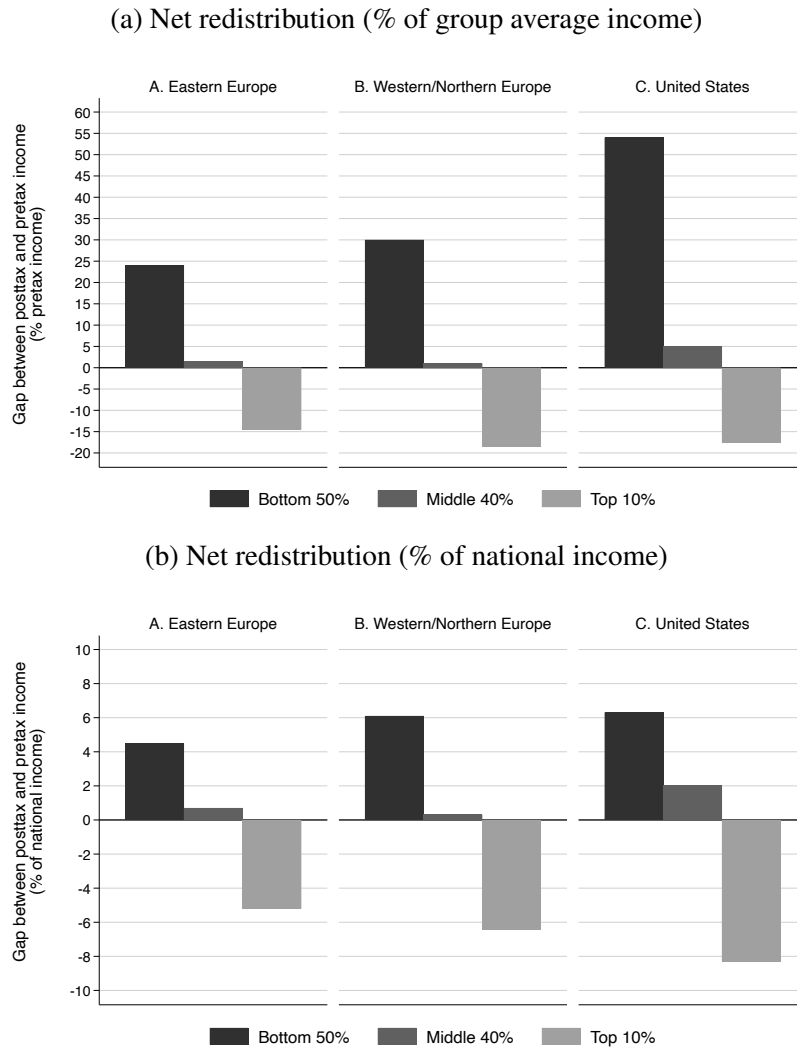
Source: Authors' computations combining surveys, tax data and national accounts for European countries and [Saez and Zucman, 2019](#) for the US. *Notes:* The unit of observation is the adult individual aged 20 or above. Income is split equally among spouses. The data correspond to population-weighted averages over the period 2007–2015 for Europe, and to 2017–2018 for the US. Taxes on products are distributed proportionally to consumption. See Appendix Table A.6 for the composition of European regions.

Figure V
Tax structure by income group in Europe and the United States



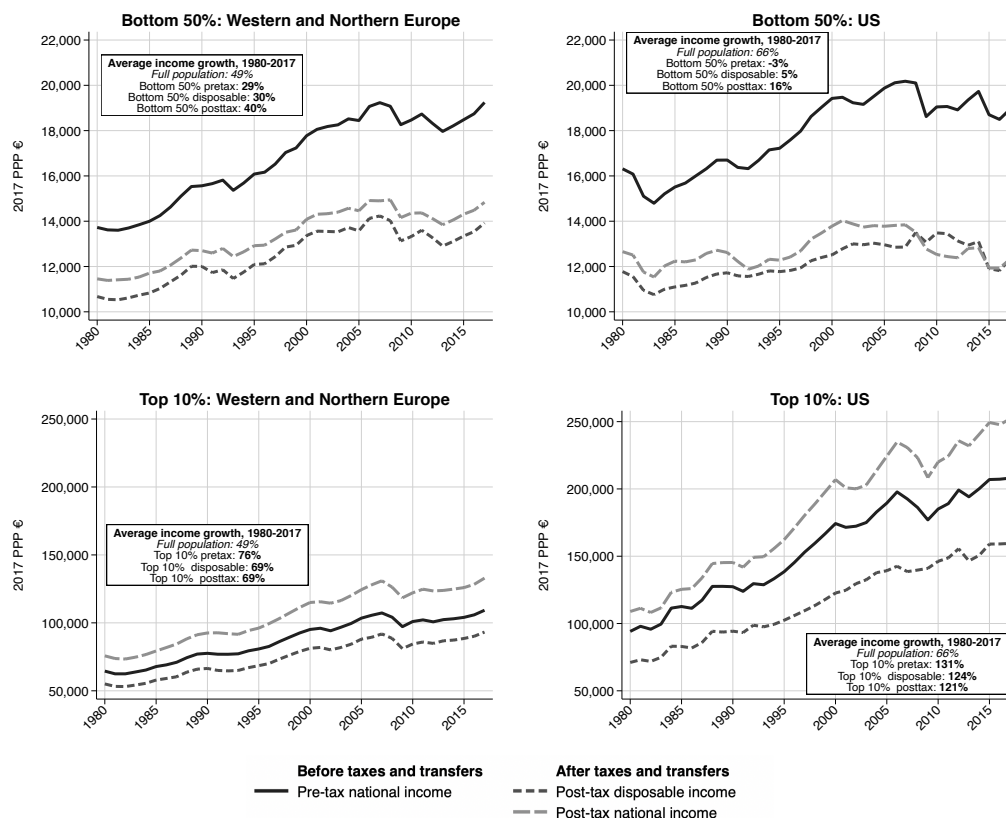
Source: Authors' computations combining surveys, tax data and national accounts for European countries and [Piketty, Saez, and Zucman, 2018](#) for the US. *Notes:* The unit of observation is the adult individual aged 20. Income is split equally among spouses. The data correspond to population-weighted averages over the period 2007–2015 for Europe, and to 2017–2018 for the US. See Appendix Table A.6 for the composition of European regions.

Figure VI
Structure of transfers by income group in Europe and the United States



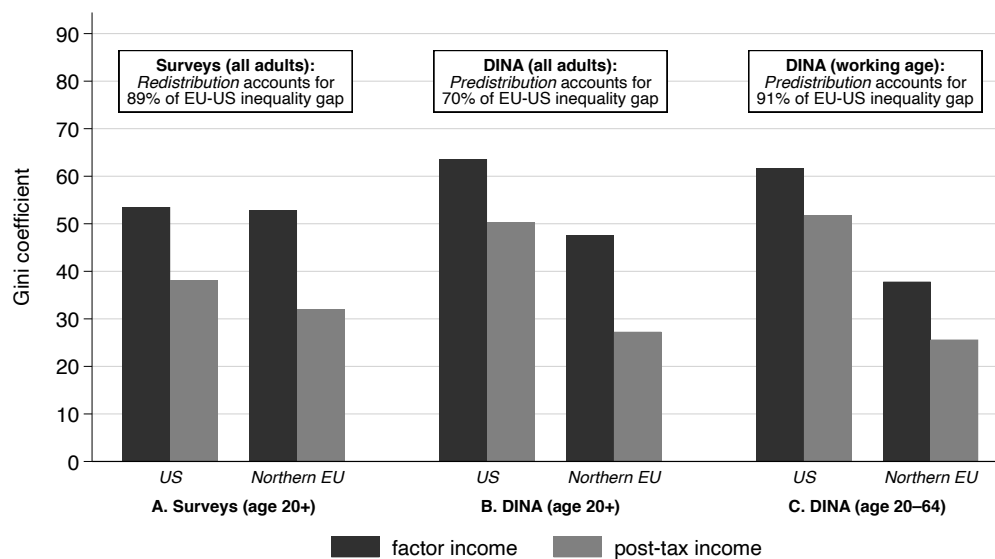
Source: Authors' computations combining surveys, tax data and national accounts for European countries and [Piketty, Saez, and Zucman, 2018](#) for the US. *Notes:* The unit of observation is the adult individual aged 20. Income is split equally among spouses. See Appendix Table A.6 for the composition of European regions.

Figure VII
Net redistribution in Europe and the US, 2017



Source: Authors' computations combining surveys, tax data and national accounts for European countries and [Piketty, Saez, and Zucman, 2018](#) for the US. *Notes:* Incomes are measured at Purchasing Power Parity in real 2017 Euros. PPP €1 = PPP\$ 1.3. The unit of observation is the adult individual aged 20. See Appendix Table A.6 for the composition of European regions.

Figure VIII
 Bottom 50% and Top 10% real incomes in Europe and the US, 1980-2017



Source: survey data for the United States comes from the Luxembourg Income Study, which is a transcription of the Current Population Survey (2007, 2010, 2013, 2017). Survey data for Europe comes from EU-SILC. DINA data for the United States comes from [Piketty, Saez, and Zucman \(2018\)](#). DINA Europe are based on authors' computations combining surveys, tax data and national accounts. *Notes:* Data are population-weighted average of the Gini coefficient in each region. The unit of observation is the adult individual aged 20 or above. Income is split equally among spouses. See Appendix Table [A.6](#) for the composition of European regions.

Figure IX
Measuring Redistribution: Impact of Data Sources and Pension Systems

Table I
The distribution of pretax income in Europe and the United States, 2017

	Eastern Europe		Northern Europe		Western Europe		United States	
	Average income	Income share	Average income	Income share	Average income	Income share		
Full population	€19,700	100%	€46,300	100%	€35,200	100%	€52,700	100%
Bottom 50%	€7,400	18.8%	€21,500	23.2%	€14,300	20.3%	€12,300	11.7%
Bottom 20%	€2,500	2.6%	€11,400	4.9%	€6,300	3.5%	€3,800	1.4%
Next 30%	€10,600	16.2%	€28,300	18.3%	€19,700	16.8%	€18,000	10.2%
Middle 40%	€22,200	45.0%	€52,700	45.5%	€39,700	45.0%	€53,300	40.5%
Top 10%	€71,600	36.3%	€145,000	31.2%	€122,000	34.7%	€252,000	47.8%
Top 1%	€238,000	12.1%	€416,000	9.0%	€367,000	10.4%	€1,110,000	21.1%
Top 0.1%	€796,000	4.0%	€1,200,000	2.6%	€1,150,000	3.3%	€5,190,000	9.8%
Top 0.01%	€2,660,000	1.3%	€3,450,000	0.7%	€3,640,000	1.0%	€23,830,000	4.5%
Top 0.001%	€8,920,000	0.5%	€9,950,000	0.2%	€11,570,000	0.3%	€92,020,000	1.7%

Source: Authors' computations combining surveys, tax data and national accounts. *Notes:* The table shows the average annual real pretax income of various groups of the population in Western and Northern Europe, Eastern Europe and the United States in 2017. Incomes measured at purchasing power parity, €1 = \$1.3. The unit of observation is the adult individual aged 20 or above. Income is split equally among spouses. See Appendix Table A.6 for the composition of European regions.

Table II
Average annual pretax income growth in Europe and the United States, 1980-2017

	Eastern Europe		Northern Europe		Western Europe		United States	
	1980-2017	2007-2017	1980-2017	2007-2017	1980-2017	2007-2017	1980-2017	2007-2017
Full population	1.3%	1.9%	1.6%	0.3%	1.0%	0.1%	1.4%	0.4%
Bottom 50%	0.3%	2.2%	1.3%	0.0%	0.7%	-0.1%	-0.1%	-1.2%
Bottom 20%	-0.9%	1.6%	0.9%	-0.7%	0.4%	-1.2%	-1.1%	-2.9%
Next 30%	0.6%	2.3%	1.3%	0.2%	0.7%	0.1%	0.1%	-0.9%
Middle 40%	1.1%	2.0%	1.5%	0.7%	0.9%	0.2%	1.0%	0.5%
Top 10%	2.1%	1.7%	2.1%	0.0%	1.5%	0.1%	2.3%	0.9%
Top 1%	3.7%	1.7%	2.9%	-1.3%	1.9%	-0.4%	3.3%	1.0%
Top 0.1%	5.7%	1.7%	3.6%	-3.1%	2.2%	-1.2%	4.2%	1.3%
Top 0.01%	7.9%	1.6%	4.4%	-5.0%	2.4%	-2.1%	4.9%	1.4%
Top 0.001%	10.2%	1.6%	5.3%	-6.9%	2.6%	-3.0%	5.4%	0.5%

Source: Authors' computations combining surveys, tax data and national accounts. *Notes:* The table shows the average annual real pretax income growth of various groups of the population in Western and Northern Europe, Eastern Europe and the United States over the 1980-2017 and 2007-2018 periods. Incomes measured at purchasing power parity. The unit of observation is the adult individual aged 20 or above. Income is split equally among spouses. See Appendix Table A.6 for the composition of European regions.

MAIN APPENDIX

FOR ONLINE PUBLICATION

A DETAILED METHODOLOGY

The issues that affect the validity and the comparability of existing income inequality estimates may be divided into three categories: conceptual discrepancies, nonsampling error, and sampling error.

Conceptual discrepancies are not errors in themselves but refer to differences as to what, precisely, is being measured. Existing estimates of income inequality may be concerned with different types of income and different populations units. While there may be a case for measuring inequality using any of these concepts and units, the existence of such a wide range of definitions makes it hard to compare inequality estimates both over time and between countries. As we have seen, both survey tabulations and fiscal data suffer from important conceptual discrepancies, as they are measured on different groups of individuals and using different income concepts. One of the contributions of this paper is to provide a new method to harmonize these different distributions.⁴²

Sampling and non-sampling errors apply to surveys. Sampling error refers to problems that arise purely out of the limited sample size of survey data. Low sample sizes affect the variance of estimates, which means they may vary a lot around their expected value. But low sample sizes may also create biases, especially when measuring inequality at the top of the distribution (Taleb and Douady, 2015). Estimates based on raw survey data do not account for any of these biases and therefore tend to underestimate incomes at the top end. Non-sampling error refers to the systematic biases that affect survey estimates in a way that is not directly affected by the sample size. These mostly include people refusing to answer surveys and misreporting their income in ways that are not observed, and therefore not corrected, by the survey producers.

The general methodology we introduce in this paper aims at correcting all three biases. We correct conceptual discrepancies by training a machine learning

⁴²Previous studies on European or global income distributions typically relied on a combination of non-harmonized income and consumption sources, see for instance Lakner and Milanovic (2016).

algorithm (Chen and Guestrin, 2016) that systematically analyzes how they affect estimates of the income distribution. We correct for non-sampling error in survey data by combining them with harmonized top income shares using a nonlinear survey calibration method (Deville and Särndal, 1992; Lesage, 2009). And we correct for sampling error by modeling the top tail of the income distribution based on extreme value theory (Ferreira and Haan, 2006). We view this methodology as a consistent and straightforward framework to exploit all published survey and tax information, while correcting for the weaknesses of these different sources. We feed to our methodology virtually all the data available and obtain estimates of inequality in Europe that reflect latest data and methodological developments.

I.A. Machine Learning Algorithm to Harmonize the Survey Data

The first step of our methodology consists in harmonizing surveys for which we are unable to recover directly the distribution of pre-tax and post-tax incomes among equal-split adults. This is the case of all survey tabulations, as well as some surveys for which we have microdata but for which pre-tax income or post-tax income was not measured. For these data sources, we have to develop a strategy to transform the distribution of the observed “source concept” (such as consumption per capita or pre-tax income among households, for instance) into an imputed distribution measured in a “target concept” (pre-tax or post-tax income per adult).

The distributions for the different income concepts across country-years are correlated: therefore, we can use the distribution for one income concept to impute the distribution for another whenever the former is observed but not the latter. To do so, we use all the cases where the income distribution is simultaneously observed for two different concepts to learn how one tends to relate to another. In practice, we use survey microdata (EU-SILC, LIS and ECHP) to compute distributions for all equivalence scales and all income concepts available in a given country-year. We then use these estimates — as well as survey tabulations observed in similar country-years but measured using different concepts — to model how different income concepts and population units relate to one another at different points of the distribution.

To clarify this idea, we can first consider a straightforward, but naive approach.

We can observe the p -th quantile of both the source and the target distributions for a variety of countries i and a variety of years t : write them $Q_{it}^{\text{target}}(p)$ and $Q_{it}^{\text{source}}(p)$. Therefore, we can estimate the average ratio between the two distributions for each percentile as $\alpha(p) = \mathbb{E}[Q_{it}^{\text{target}}(p)/Q_{it}^{\text{source}}(p)]$. Say that for a country j in year s , we only observe the source concept $Q_{js}^{\text{source}}(p)$. Then we can approximate the target concept as $Q_{js}^{\text{target}}(p) = \alpha(p)Q_{js}^{\text{source}}(p)$. While this remains an approximation, it at least corrects for some systematic discrepancies that we can observe in the data.

That approach has the merit of simplicity. When we tried it with our data, it gave passable results (see table A.3 and discussion below). But there are several problems with it, both in theory and in practice. The main issue is that it makes a very restrictive assumption about the way different income concepts may relate to one another: it considers that the sole predictor of, say, the 25th percentile of income for equal-split adults is the 25th percentile of income for households. Furthermore, it assumes that the relationship is entirely linear. There is no good theoretical reason for any of that to be true: a better, more general model would allow that 25th percentile of the target distribution to depend on any percentile of the source distribution, including but not limited to the 25th. It would also allow these relationships to be nonlinear and potentially with interactions. That relationship could also depend on auxiliary variables Z_{it} capturing demographic, political and institutional factors. The simple approach also cannot ensure that the estimated distribution for the target concept will be increasing, which creates problems that have to be dealt with in an *ad hoc* way (e.g. by re-ranking percentiles) and imply inefficient use of information. This is particularly true for the bottom of the distribution for which incomes can be close to zero and the ratios may therefore be very unstable.

Therefore, to construct the best mappings between the different concepts, we consider a much more general model. In that model, each percentile of the target distribution is an arbitrary function of every percentile of the source distribution, and of additional covariates. We write:

$$\mathbb{E}[Q_{it}^{\text{target}}(p)] = \varphi(Q_{it}^{\text{source}}(p_1), \dots, Q_{it}^{\text{source}}(p_m), p, t, Z_{it})$$

for a grid $0 \leq p_1 < \dots < p_m < 1$ of fractiles. Estimating such a model raises

some challenges. Linear regression will not be flexible enough due to its parametric assumptions and will tend to overfit the data if m is large due to the number of covariates.

To estimate this model, we therefore rely on more recent advances in high-dimensional, nonparametric regression, also known as *machine learning* methods. The algorithm we use is known as *boosted regression trees*, a powerful and commonly used method introduced by Friedman (2001). We rely on an implementation known as XGBoost (Chen and Guestrin, 2016), which has enjoyed great success due to its speed and performance, to the point that it has earned a reputation for “winning every machine learning competition” (Nielsen, 2016). On top of their performance, boosted regression makes it easy to deal with missing values, or to impose certain constraints, such as the fact that the quantile function $Q(p)$ must be increasing with p .

The algorithm starts from regression trees, a fast and simple nonlinear prediction method that successively cuts the space of predictors into two subspaces in which the predicted variable has lower variance. This leads to a “tree” of simple decision rules based on the value of the predictors. Following these rules the algorithm places any observation into a subspace where the predictor should have a relatively low variance, and the predicted value for that observation is the average of the predictor within that subspace.

Regression trees provide predictions that are simple, but rough. “Boosting” is a method that combines many of these simple but low accuracy prediction methods into a high-accuracy one. It starts by estimating a regression tree. It then runs a second regression tree to predict the residual from the previous regression: this is called a “boosting round.” The process is repeated several times: each round of boosting forces the algorithm to concentrate on the part of the data where the previous predictions failed. In the end, all the regression trees are combined into a single prediction.

The appropriate number of boosting rounds is determined by cross-validation: the sample is divided into K subsamples. For each subsample, we train the algorithm on the data excluding the subsample, and we test the prediction on the excluded subsample: we use the number of boosting rounds for which the cross-validation

prediction error is lowest. By excluding the sample on which we perform the prediction, we make sure to avoid overfitting to the data on which we estimate the model.

[Table A.1 about here.]

[Table A.2 about here.]

[Table A.3 about here.]

Since our dataset is made up of countries that we follow over the years, it has a panel dimension, which we take into account as follows. We assume that the country-specific prediction error is independent conditional on all observed variables (i.e. that it is a *random* rather than a *fixed* effect.) Under that assumption, the imputation method remains valid because the error term remains exogenous. However, there is a risk of over-fitting if we do not make sure that the different subsamples used in the cross-validation are not independent, because then we would force the algorithm to try to predict the country random effect. To avoid that problem, we perform the cross-validation by making sure that all the observations for one country are in the same cross-validation subsample, which is known as leave-one-cluster-out cross validation (Fang, 2011). When possible, we also estimate and include the country random effect into our imputation. The random effect is estimated as a function of the percentile using the mean prediction error by country and percentile.

In the end, for any target concept of interest, we get as many predictions as there are sources available. Let $\mathbf{y} = (\hat{Q}_{it}^{\text{target},1}, \dots, \hat{Q}_{it}^{\text{target},n})'$ the n different predictions. Using the cross-validation estimation of the prediction error, we can estimate the variance-covariance matrix Σ between the different predictions. Following the logic of generalized least squares, the optimal way of combining the n predictions into one is to average them, weighted by the row or column sums of the symmetric matrix Σ . This yields our harmonized estimate of the distribution, taking into account observed regularities across concepts and percentile groups.

As table A.1 shows, the mean (cross-validation) prediction error for the value of the average of a percentile is between 2% and 11% depending on the concept that

was used for the prediction.⁴³ Adjusting for the statistical unit while keeping the income concept identical creates the least difficulties. Consumption, on the other hand, is a rather poor predictor of income. Moving from post-tax to pre-tax income is a somewhat intermediary situation. The auxiliary variables that we use to improve the performance of the prediction are: regional dummies, average national income per adult (PPP), share of household with size 1 to 6, gross saving rate (% of GDP), overall social expenditures (% of GDP), top marginal income tax rate, income tax revenue (% of GDP), overall tax revenue (% of GDP), share of population by 10-year age bands and sex, corporate tax rate, VAT tax rate. Table A.2 shows the performance of a model that does not include these variables. While their inclusion has only second-order effects on our harmonized series, they do improve the prediction error, especially when trying to impute based on consumption: we improve the mean relative error by up to 2 pp.

Table A.3 shows the performance of a much more simple imputation method, namely using a single correction coefficient by percentile to move from one concept to another. This coefficient is computed as the mean ratio between two concepts for a given percentile. While this method performs reasonably well for concepts that are close to one another, it exhibit much worse performance when using a poor predictor such as consumption. In such cases, the prediction can be 50% or even 100% worse than our benchmark algorithm.

I.B. Calibration on Top Income Shares to Correct for Non-sampling Error

We correct survey data for non-sampling error using known top income shares estimated from administrative tax data. We do so by adjusting the survey weights using survey calibration methods (Deville and Särndal, 1992). Statistical institutes already routinely use these methods to ensure that their surveys are representative,

⁴³Before training the model, we transform the data using the transform $y \mapsto \text{asinh}(y)$ for the value of the quantiles and $x \mapsto -\log(1 - x)$ for the corresponding rank. This stabilizes the mode without changing the nature of the data. The use of asinh rather than the logarithm avoid issues with having zero or near-zero incomes at the bottom of the distribution. All distributions are normalized by their average since we are only concerned with the distribution of income. When we report prediction errors, these are computed for distributions that have been properly transformed back to their original value.

typically in terms of age and gender. Our approach is a natural extension of theirs, in the sense that we enforce representativity in terms of taxable income in addition to age and gender.

We start by explaining how survey calibration works in the standard, linear framework. This framework, however, does not apply directly to this situation because the top income shares that we use for calibration are not linear statistics. Therefore, we explain how to use an extension of the usual calibration framework suggested by [Lesage \(2009\)](#) to apply calibration to nonlinear constraints.⁴⁴

Statistically, survey calibration can be interpreted as the estimation of a non-response function, in which non-response depends on the variables introduced in the constraints. In that interpretation, we are assuming that nonresponse has the same shape as the influence function for top shares. This shape is that of a continuous, piecewise linear function with a kink at the threshold corresponding to the top share. It is almost flat below that threshold, meaning that the bottom 90% of the distribution is virtually unchanged. Above the threshold, nonresponse increases linearly with income — though we can capture non-linearity of nonresponse at the top by including several top income groups in the calibration, for example top 10%, 5% and 1%. That shape is what we expect if the richest households refuse to answer surveys at a higher rate, and also corresponds to the share of the nonresponse that we observe with access to richer data ([Blanchet, Flores, and Morgan, 2018](#)). Because the nonresponse function is continuous, our correction method preserves the continuity of the density function of income.

The average estimated nonresponse profile over all the survey and tax data is mostly flat for most of the distribution, meaning that survey distribution is mostly preserved. But observations in the top 0.1% are underrepresented by a factor of 3 on average.

When we do not directly observe tax data in a country, we still perform a correction based on the profile of nonresponse that we observe in other countries. To capture potential statistical regularities, we estimate the nonresponse profile as

⁴⁴Here it is important to distinguish nonlinear calibration in the sense that the distance function chosen is not quadratic — and therefore requires nonlinear optimisation methods — and nonlinear in the sense that the constraints of the problem are not linear. The former case is relatively standard, but here we face the latter.

a function of the distribution of income in the uncorrected survey using the same machine learning algorithm as in section I.A.. We stress that this remains a rough approximation and that in our view the proper estimation of top income inequality requires access to tax data. Fortunately, our tax data covers a large majority of the European population and an even larger majority of European income, so that the impact of these corrections on our results remain limited.

For the most recent years, in a few cases where survey data is more up-to-date than tax data, we updates the top income shares as follows. We estimate the change in median income by decile over the years, and use it to extrapolate the top income shares in most recent years. Given surveys' difficulties in tracking top incomes, we do not attempt to account for changes in very top income in this extrapolation (hence the use of the median). However, we still account for changes in the bottom of the distribution that would impact top shares.

Survey Calibration in the Linear Case Let U be a finite population of size N , with units indexed by $k \in \{1, \dots, N\}$. Let y be a variable (say, income) that takes the value y_k for the unit k . Let s be a random survey subsample drawn from U of size n . Let $\pi_k = \mathbb{P}\{k \in s\}$ be the probability that k is included in the sample s . The value $d_k = 1/\pi_k$ is called the *design weight* of observation k . If a statistic T over the complete population U can be written:

$$T = \sum_{k \in U} \phi(y_k)$$

Then the Horvitz-Thompson estimator of this quantity over the subsample s is:

$$\hat{T} = \sum_{k \in s} d_k \phi(y_k)$$

Assume that we know, from an external source, the value of m statistics $(C_1, C_2, \dots, C_m)$ over the complete population U that can be written, for $p \in \{1, \dots, m\}$:

$$C_p = \sum_{k \in U} \phi_p(y_k)$$

The Horvitz-Thompson estimator $\hat{C}_p$ of C_p is:

$$\hat{C}_p = \sum_{k \in s} d_k \phi_p(y_k)$$

In general, $\hat{C}_p$ will not be exactly equal to C_p , either because the design weights d_k are invalid (because of, say, unit nonresponse correlated with y), or simply because of sampling variability.

This raises the question: is it possible to improve the design weights $\{d_k, k \in s\}$ using the information contained in $(C_1, C_2, \dots, C_m)$? The question was answered positively by [Deville and Särndal \(1992\)](#), using the calibration procedure.

The Standard Calibration Procedure Let $\delta : (x, y) \mapsto d(x, y)$ be a distance function. For any statistic T written as:

$$T = \sum_{k \in U} \phi(y_k)$$

and any set of *calibration weights* $\{w_k, k \in s\}$, define the calibration estimator:

$$\tilde{T} = \sum_{k \in s} w_k \phi(y_k)$$

The calibration procedure finds $\{w_k, k \in s\}$ to solve the constrained optimization program:

$$\min \sum_{k \in s} \delta(w_k, d_k) \quad \text{s.t.} \quad \forall p \in \{1, \dots, m\} \quad \tilde{C}_p = C_p$$

That is, it finds a set of weights as close as possible from the initial weights (thus minimizing distortions from the original distribution), such that the constraints $\tilde{C}_p = C_p$, known as the *calibration margins*, are satisfied.

Because of the margins are a linear function of the data, the solution can be written:

$$w_k = d_k F[\lambda_1 \phi_1(y_k) + \dots + \lambda_p \phi_p(y_p)]$$

where $d_k F(w)$ is the inverse of $\frac{\partial}{\partial w} \delta(w, d_k)$. The $\lambda_1, \dots, \lambda_p$ are Lagrange multipliers whose value is determined by solving the equations associated to the equality constraints.

Interpretation as a Nonresponse Model This result can further be interpreted in terms of a nonresponse model. Indeed, $1/d_k$ is the probability of unit k being selected for inclusion in the sample. Ideally, w_k should correspond to the probability of unit k being effectively included in the sample, taking into account the possibility of unit nonresponse. Therefore,

$$\frac{w_k}{d_k} = F[\lambda_1 \phi_1(y_k) + \dots + \lambda_p \phi_p(y_p)]$$

is the probability of nonresponse as a function of $\phi_1(y_k), \dots, \phi_p(y_p)$. The $\lambda_1, \dots, \lambda_p$ are the parameters of the model, and F is the link function. If $\delta(w, d) = \frac{1}{2}(w - d)^2/d$ is the χ^2 distance, then F is linear and we get a linear probability model. But we could also choose δ to get, say, a logit model (which avoids the risk of estimated probabilities below zero, and therefore negative weights).

Calibration has been shown to reduce both the variance and the bias in survey data. The variance reduction is asymptotically identical regardless of the distance used for the procedure. The bias reduction may depend of the specific distance chosen because it determines the nonresponse model. But even if the model is not exactly right we can expect significant improvement.

However, standard calibration methods only apply to linear statistics — meaning statistics that can be written as a linear function of the data. This is not the case of inequality statistics, including top shares. [Lesage \(2009\)](#) suggested two methods to solve the nonlinear calibration problem. The first one involves linearizing the top shares using their *influence function*. Informally, the influence measures the marginal contribution of the weight of each observation to the overall statistic. In the case of top income shares, we have the following derivation:

Partial Sums and Top Shares in Survey Data Let $\alpha \in [0, 1]$. Over the complete population U , we define the partial sum of the bottom $100\alpha\%$ as:

$$Y_\alpha = \sum_{k \in U} y_k H(\alpha N - k + 1)$$

where $H(x) = 0$ if $x < 0$, $H(x) = x$ if $0 \leq x < 1$ and $H(x) = 1$ if $x \geq 1$. Its survey sample counterpart is:

$$\hat{Y}_\alpha = \sum_{k \in s} y_k H\left(\frac{\alpha N - W_{k-1}}{w_k}\right)$$

where $W_k = \sum_{l \in s} w_l \mathbb{1}_{y_l \leq y_k}$ and $N = W_n$. The income share of the top $100(1-\alpha)\%$ is:

$$S_\alpha = \frac{Y - Y_\alpha}{Y} \quad \text{and} \quad \hat{S}_\alpha = \frac{\hat{Y} - \hat{Y}_\alpha}{\hat{Y}}$$

Linearization of Top Shares Neither the top share or the partial sum can be rewritten as $\sum_{k \in s} w_k \phi(y_k)$ for some ϕ . Therefore we cannot directly apply the calibration method to it. A solution suggested by [Lesage \(2009\)](#) is to linearize the statistic using Deville's (1999) concept of *influence*.

The influence measures the effect of a marginal change in the weight of an observation on the statistic of interest. Formally, let M be the measure that puts a weight equal to 1 on each individual in U , and $M + t\delta_k$ the measure that puts a weight equal to 1 on each individual except k , which has a weight $1 + t$. Let $T(M)$ and $T(M + t\delta_k)$ be the corresponding values of an arbitrary statistic T . The influence of observation k is:

$$I(T)_k = z_k = \lim_{t \rightarrow 0} \frac{T(M + t\delta_k) - T(M)}{t}$$

Here, $\sum_{k \in U} z_k$ can be viewed as the linearized version on the original statistic. As [Lesage \(2009\)](#) explains, we can perform an approximate calibration on the statistic T in the nonlinear case by using the variable z_k instead of y_k in the standard calibration method. [Langel and Tillé \(2011\)](#) show that the linearized partial sum can be written

as:

$$I(Y_\alpha)_k = y_k H(\alpha N - k + 1) + (\alpha - \mathbb{1}_{y_k < Q_\alpha}) Q_\alpha$$

where $Q_\alpha = y_i$ with $N_{i-1} < \alpha N \leq N_i$. Its survey sample counterpart is:

$$I(\hat{Y}_\alpha)_k = y_k H\left(\frac{\alpha N - W_{k-1}}{w_k}\right) + (\alpha - \mathbb{1}_{y_k < \hat{Q}_\alpha}) \hat{Q}_\alpha$$

where $\hat{Q}_\alpha = y_i$ with $W_{i-1} < \alpha W_n \leq W_i$.

Enforcing the constraint $\hat{S}_\alpha = S_\alpha$ is equivalent to enforcing $\hat{Y}_\alpha - (1 - S_\alpha)\hat{Y} = 0$. Therefore we can calibrate the survey directly on the top share by setting:

$$z_k = y_k H\left(\frac{\alpha N - W_{k-1}}{w_k}\right) + (\alpha - \mathbb{1}_{y_k < \hat{Q}_\alpha}) \hat{Q}_\alpha - (1 - S_\alpha) y_k$$

and calibrate so that the sum of the z_k is zero.

Introduction of a Nuisance Parameter The second solution of [Lesage \(2009\)](#) involves the introduction of a *nuisance parameter*. For the top $(1 - \alpha) \times 100\%$ share, the nuisance parameter is the true value of the α -quantile of income. Given that value, one can apply standard calibration methods to impose the proper number of people and their proper amount of income on both sides of the quantile. The advantage is that this leads to the constraint being exactly satisfied. But for that method to give acceptable results, we need a good guess for the value of the nuisance parameter. [Lesage \(2009\)](#) suggests using its value in the original survey.

We obtained the best results by combining both methods. In the first step, we use the influence function method. This performs the majority of the required adjustment, but still leaves a small discrepancy between the survey and the tax data. In the second step, we get rid of the remaining discrepancy by applying the second approach, with the nuisance parameter estimated in the survey corrected through the first step.

I.C. Extreme Value Theory to Correct for Sampling Error

The sample size of surveys varies a lot and can sometimes be quite low: this, in itself, can seriously affect estimates of inequality at the top and, in general, will

underestimate it ([Taleb and Douady, 2015](#)). Correcting sampling error requires some sort of statistical modeling. We chose to use methods coming from extreme value theory, which is routinely used in actuarial sciences to estimate the probability of occurrence of very rare events, but can similarly be used to estimate the distribution of income at the very top.

The main tenet of extreme value theory can be understood in analogy to the central limit theorem. According to the central limit theorem, under some regularity assumptions, but regardless of the exact distribution of iid. variables $X_1, \dots, X_n$, the distribution of the sum $\sum_{i=1}^n X_i$ as n goes to infinity will belong to a tightly parametrized family of distributions (a Gaussian one). Similarly, under mild regularity assumptions, the distribution of the largest value of the sample $\max(X_1, \dots, X_n)$ as n goes to infinity will belong to a certain parametric family. The same holds for the second-largest value, the third-largest value, and so on. As a result, the top k largest values will approximately follow a distribution known as the generalized Pareto distribution, which has the cumulative distribution function:

$$F(x) = 1 - \left\{ 1 + \xi \left(\frac{x - \mu}{\sigma} \right) \right\}^{-1/\xi}$$

That result is known as the Pickands–Balkema–de Haan theorem (e.g. [Ferreira and Haan, 2006](#)). The generalized Pareto distribution therefore more or less provides a universal approximation of the distribution of the tails of distributions. It includes the Pareto or the exponential distribution as a special case. We use it to model the top 10% of income distributions. Because the likelihood surface of the generalized Pareto distribution is very flat, maximum likelihood estimation often gives poor results unless the sample size is very large. The standard method of moments also fails if the distribution has infinite variance, which can often occur with income distributions. We use a simple and robust alternative known as probability-weighted moments ([Hosking and Wallis, 1987](#)). For X following a generalized Pareto distribution, define $a = \mathbb{E}[X]$ and $b = \mathbb{E}[X(1 - F(x))]$. Then we have $\xi = (a - 4b + \mu)/(a - 2b)$ and $\sigma = (a - \mu)(2b - \mu)/(a - 2b)$, while μ is determined a priori from the threshold from which we start to use the model. We obtain the complete distribution by combining the empirical distribution for the bottom 90% with the generalized Pareto

model for the top 10%.

I.D. Statistical Matching Procedure

To combine data from several sources, we use the following matching procedure. We define a single (continuous) matching variable in common in both datasets. We rank both datasets according to that variable, and then match observations one-by-one on their rank. In practice, between different datasets have different weights and different sample sizes, observations have to be partially matched with one another. For example, imagine that the first (sorted) dataset has the weights $\{3, 1, \dots\}$ and the second one the weights $\{2, 4, \dots\}$. The matched dataset starts with one observation with weight 2 that has the characteristics of the first observation of each dataset. However, the first observation of the first dataset cannot be fully matched because its weight (3) is larger than the weight of the first observation from the second dataset (2). So we keep the first observation in the first dataset with its remaining weight (1), and match it to the second observation of the second dataset. That observation's weight (4) is in turn larger than 1, so we follow the same procedure. We continue the process until all the probability mass from both datasets has been matched. One can show that, if the initial datasets have sizes N and M , the matched dataset will at most have size $N + M - 1$.

This procedure is akin to finding the optimal transport map between both datasets with a cost that is a convex function of the distance between the value of the matching variable in each datasets ([Rachev and Rüschendorf, 1998](#)). By construction, it preserve the distribution of the variables in each samples, and preserve the copulas between the variables.

I.E. Income Distribution in the US

To compare the geography of inequality in Europe with that of the United States, we use distributional national accounts data from [Piketty, Saez, and Zucman \(2018\)](#) and national accounts data by US state.

We attribute national income to each state based on their share of GDP (the only national account aggregate available at the state level). To that end, we use data on

total state domestic products from the Bureau of Economic Analysis, along with state adult populations series from the United States Census Bureau. (State domestic products provided by the Bureau of Economic Analysis go back as far as 1967. For the historical part, we extrapolate these series back to 1929 by using the growth rates in personal income per capita available from [Barro and Sala-i-Martin \(1992\)](#).)

This provides us with an estimate of national income by state, which lets us compute between-state inequality in the United States. Using the data from [Piketty, Saez, and Zucman \(2018\)](#), we can calculate the overall Theil index for the United States. Using the decomposability of the Theil index, we can then estimate within-state component of inequality for the United States as a residual.

B DETAILED STRUCTURE OF TAXES AND TRANSFERS IN EUROPE AND THE US

II.A. Structure of taxes

As per National Accounts, Western and Northern European countries are characterized by a higher macroeconomic tax rate (i.e. total taxes and social contributions paid as a share of national income) than Eastern European countries (49.8% vs. 41.0%, respectively). The total macroeconomic tax rate in the US is significantly lower, at 28.2%. We group taxes and contributions into four broad categories: (i) social contributions, (ii) indirect and consumption taxes, (iii) taxes on personal income and wealth, and (iv) taxes on corporate profits. Social contributions represent 20.7% of the national income in Western and Northern Europe, 16.4% in Eastern Europe and 7.6% in the US. This category can be divided into a contributory component (representing about 80-90% of social contributions) and a non-contributory component. Contributory contributions are primarily made of pension payments, representing about 80% of the category, with unemployment insurance making up the rest. Non-contributory contributions typically represent health insurance and disability payments. Indirect and consumption taxes represent 13.7% of the national income in Western and Northern Europe, 15.8% in Eastern Europe and 6.5% in the US. Income and wealth taxes represent 12.2%, 5.8%, and 11.2% of the national income in Western and Northern Europe, Eastern Europe and the US respectively, whereas corporate taxes represent 3.2%, 3.0% and 3.0% of national income in the

three regions respectively (see Appendix Table A.4). Appendix table A.5 presents separate results for Western and Northern European countries. The overall level of taxes and transfers is slightly higher in Northern Europe, but the main difference between the two regions is that social insurance is mainly financed by social contributions in Western European countries, while they are financed by income taxes in Northern European countries.

[Table A.4 about here.]

II.B. Structure of transfers

Transfers represent 46.9% of the national income in Western and Northern Europe, 40.5% in Eastern Europe, and 34.5% in the US. We distinguish five types of transfers, which can be grouped into two main categories: cash transfers and in-kind transfers. Cash transfers represent 22.2%, 17.7% and 8.8% of the national income in Western and Northern Europe, Eastern Europe, and the United States respectively. Cash transfers can take the form of social insurance (unemployment insurance and disability benefits) or of social assistance (family benefits and other cash transfers). Pensions represent the largest part of cash transfers in rich countries, and also account for most of the gap in total transfers between European countries and the US. They represent about 16% of the national income in Western and Northern Europe, compared to 5% in the US.

Unemployment insurance and disability payments represent 1.8%, 0.8 % and 1.4% of national income in the three regions, respectively. Family and other cash social assistance benefits represent 4.8% vs. 3.4% vs. 2.7% of national income in the three regions, respectively.

In-kind transfers represent a similar share of national income in the three regions (22% to 26%). They can be divided in two categories: health insurance and non-individualizable collective consumption expenditures.⁴⁵ Public health expenditures

⁴⁵In the national accounts, government consumption expenditure (P3S13) is divided into individual consumption expenditure (P31S13) and collective consumption expenditure (P32S13). P31S13 includes education and health expenditures and P32S13 includes all other expenditures (justice, police, transportation, research etc.). We effectively treat education as a collective consumption expenditure.

represent a quarter to a third of in-kind transfers. They are slightly higher in Western and Northern Europe than in the US, and significantly lower in Eastern Europe. Collective expenditure (e.g. general administration, defense, justice) and education spending represent the bulk of in-kind transfers (16% to 18% of national income).⁴⁶

C APPENDIX FIGURES AND TABLES

[Table A.5 about here.]

[Table A.6 about here.]

[Figure A.1 about here.]

[Figure A.2 about here.]

[Figure A.3 about here.]

[Figure A.4 about here.]

[Figure A.5 about here.]

[Figure A.6 about here.]

[Figure A.7 about here.]

[Figure A.8 about here.]

[Figure A.9 about here.]

⁴⁶While total transfers are very close in Northern and Western Europe, the two regions differ on the allocation of transfers. Northern European countries spend less on pensions and relatively more on other transfers such as unemployment and disability, health and education.

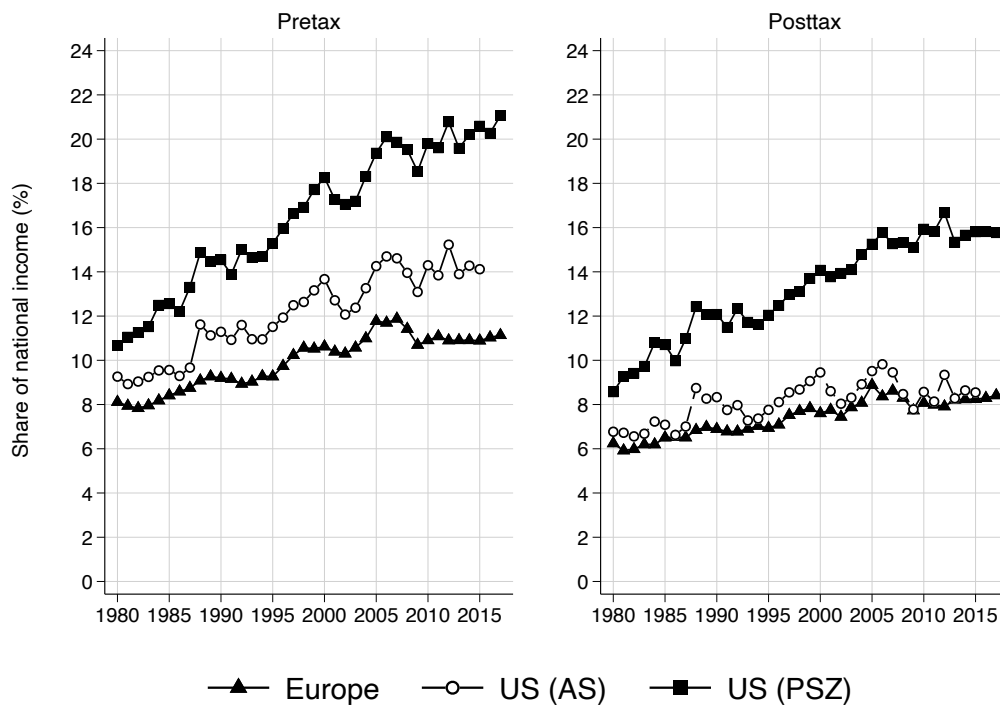
REFERENCES

- Aaberge, R. and A. B. Atkinson (2010). “Top incomes in Norway”. In: *Top incomes: a global perspective*. Ed. by A. B. Atkinson and Thomas Piketty. Oxford University Press. Chap. 9, pp. 448–481.
- Alvaredo, Facundo (2009). “Top incomes and earnings in Portugal 1936-2005”. In: *Explorations in Economic History* 46.1, pp. 404–417.
- Alvaredo, Facundo and Elena Pisano (2010). “Top incomes in Italy, 1974-2004”. In: *Top incomes: a global perspective*. Ed. by A. B. Atkinson and Thomas Piketty. Oxford University Press. Chap. 12, pp. 625–663.
- Alvaredo, Facundo and Emmanuel Saez (2010). “Income and wealth concentration in Spain in a historical and fiscal perspective”. In: *Top incomes: a global perspective*. Ed. by A. B. Atkinson and Thomas Piketty. Oxford University Press. Chap. 10, pp. 482–559.
- Atkinson, A. B. (2007). “The distribution of top incomes in the United Kingdom 1908-2000”. In: *Top incomes over the 20th century*. Ed. by A. B. Atkinson and Thomas Piketty. Oxford University Press. Chap. 4, pp. 82–140.
- Atkinson, A.B. and J.E. Sørensen (2013). “The long-run history of income inequality in Denmark: Top incomes from 1870 to 2010”. In: *EPRU Working Paper Series 2013-01*.
- Auten, Gerald and David Splinter (2019). “Income inequality in the United States: Using tax data to measure long-term trends”. In: *Working Paper*.
- Barro, Robert and Xavier Sala-i-Martin (1992). “Convergence”. In: *Journal of Political Economy* 100.2, pp. 223–251.
- Bartels, Charlotte (2017). “Top incomes in Germany, 1871-2013”. In: *WID.world Working Paper Series 2017/18*.
- Blanchet, Thomas, Ignacio Flores, and Marc Morgan (2018). “The Weight of the Rich: Improving Surveys Using Tax Data”.
- Bukowski, Pawel and Filip Novokmet (2017). “Top incomes during wars, communism and capitalism: Poland 1892-2015”. In: *WID.world Working Paper Series 2017/22*.

- Chen, Tianqi and Carlos Guestrin (2016). “XGBoost: A Scalable Tree Boosting System”.
- Chrissis, Kostas and Franciscos Koutentakis (2017). “From dictatorship to crisis: the evolution of top income shares in Greece (1967-2013)”. In: *Working Paper*.
- Decoster, André, Koen Dobbeleer, and Sebastiaan Maes (2017). “Using fiscal data to estimate the evolution of top income shares in Belgium from 1990 to 2013”. In: *Ku Leuven Discussion Paper Series DPS17.18*.
- Deville, Jean-Claude (1999). “Estimation de variance pour des statistiques et des estimateurs complexes: linéarisation et techniques des résidus”. In: *Techniques d’enquête* 25.2, pp. 219–230.
- Deville, Jean-Claude and Carl-Erik Särndal (1992). “Calibration Estimators in Survey Sampling”. In: *Journal of the American Statistical Association* 87.418, pp. 376–382.
- Fang, Yixin (2011). “Asymptotic Equivalence between Cross-Validations and Akaike Information Criteria in Mixed-Effects Models”. In: *Journal of Data Science* 9, pp. 15–21.
- Ferreira, Ana and Laurens de Haan (2006). *Extreme Value Theory: An Introduction*. Springer Series in Operations Research. Springer.
- Foellmi, Reto and Isabel Z. Martínez (2017). “Volatile top income shares in Switzerland? Reassessing the evolution between 1981 and 2010”. In: *Review of Economics and Statistics* 99.5, pp. 793–809.
- Friedman, Jerome H. (2001). “Greedy Function Approximation: A Gradient Boosting Machine”. In: *The Annals of Statistics* 29.5, pp. 1189–1232.
- Garbinti, B., J. Goupille-Lebret, and T. Piketty (2018). “Income inequality in France, 1900-2014: Evidence from Distributional National Accounts (DINA)”. In: *Journal of Public Economics* 162.1, pp. 63–77.
- Hosking, J R M and J R Wallis (1987). “Parameter and Quantile Estimation for the Generalized Pareto Distribution”. In: *Technometrics* 29.3, pp. 339–349.
- Jäntti, M., M. Riihelä, R. Sullström, and M. Tuomala (2007). “Long-term trends in top income shares in Ireland”. In: *Top incomes over the 20th century*. Ed. by A. B. Atkinson and Thomas Piketty. Oxford University Press. Chap. 12, pp. 501–530.

- Jäntti, M., M. Riihelä, R. Sullström, and M. Tuomala (2010). “Trends in top income shares in Finland”. In: *Top incomes: a global perspective*. Ed. by A.B. Atkinson and Thomas Piketty. Oxford University Press. Chap. 8, pp. 371–447.
- Kump, Nataša and Filip Novokmet (2018). “Top incomes in Croatia and Slovenia, from 1960s until today”. In: *WID.world Working Paper Series 2018/8*.
- Lakner, Christoph and Branko Milanovic (2016). “Global Income Distribution: From the Fall of the Berlin Wall to the Great Recession”. In: *The World Bank Economic Review* 30.2, pp. 203–232.
- Langel, Matti and Yves Tillé (2011). “Statistical inference for the quintile share ratio”. In: *Journal of Statistical Planning and Inference* 141.8, pp. 2976–2985.
- Lesage, Éric (2009). “Calage non linéaire”.
- Mavridis, Dimitris and Pálma Mosberger (2017). “Income inequality and incentives: the quasi-natural experiment of Hungary, 1914-2008”. In: *WID.world Working Paper Series 2017/17*.
- Nielsen, Didrik (2016). *Tree Boosting With XGBoost Why Does XGBoost Win “Every” Machine Learning Competition?* Tech. rep. December, p. 2016.
- Novokmet, Filip (2018). “The long-run evolution of inequality in the Czech lands, 1898-2015”. In: *WID.world Working Paper Series 2018/6*.
- Oancea, B., T. Andrei, and D. Pirjol (2017). “Income inequality in Romania: The exponential-Pareto distribution”. In: *Physica A: Statistical Mechanics and its Applications* 469.1, pp. 486–498.
- Piketty, Thomas, Emmanuel Saez, and Gabriel Zucman (2018). “Distributional National Accounts: Methods and Estimates for the United States”. In: *Quarterly Journal of Economics* 133.2, pp. 553–609.
- Rachev, Svetlozar T and Ludger Rüschendorf (1998). *Mass Transportation Problems*. Vol. 1. Probability and its Applications. New York: Springer.
- Roine, Jesper and Daniel Waldenström (2010). “Top incomes in Sweden over the twentieth century”. In: *Top incomes: a global perspective*. Ed. by A. B. Atkinson and Thomas Piketty. Oxford University Press. Chap. 7, pp. 299–370.
- Salverda, W. and A. B. Atkinson (2007). “Top incomes in the Netherlands over the twentieth century”. In: *Top incomes over the 20th century*. Ed. by A. B. Atkinson and Thomas Piketty. Oxford University Press. Chap. 10, pp. 426–471.

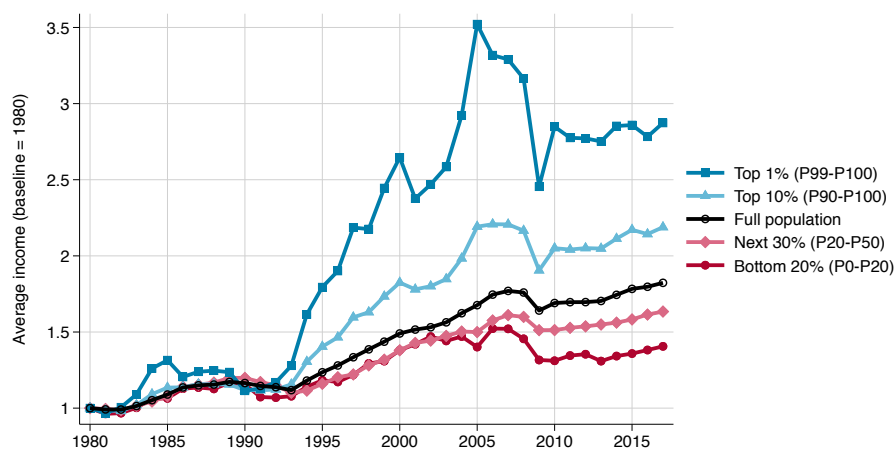
Taleb, Nassim Nicholas and Raphael Douady (2015). “On the super-additivity and estimation biases of quantile contributions”. In: *Physica A: Statistical Mechanics and its Applications* 429, pp. 252–260.



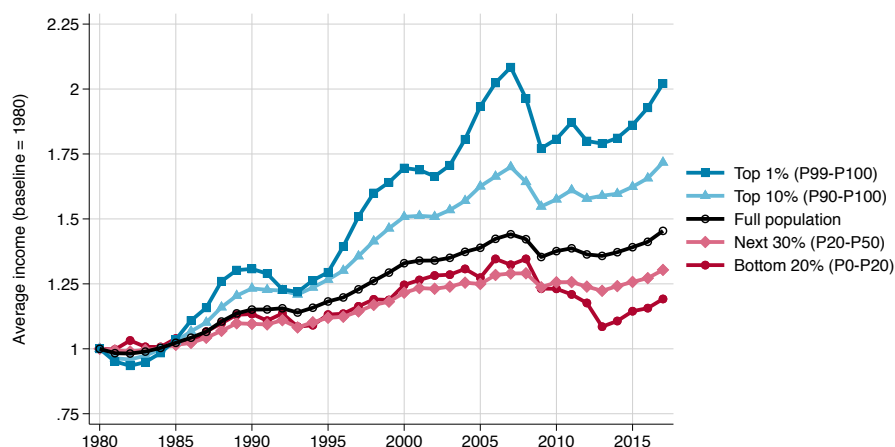
Source: Authors' computations combining surveys, tax data and national accounts for European countries and [Piketty, Saez, and Zucman, 2018](#) (US-PSZ) as well as [Auten and Splinter, 2019](#) (US-AS) for the US.

Figure A.1
Top 1% income share in Europe and the US

(a) Cumulated growth by income group: Northern Europe



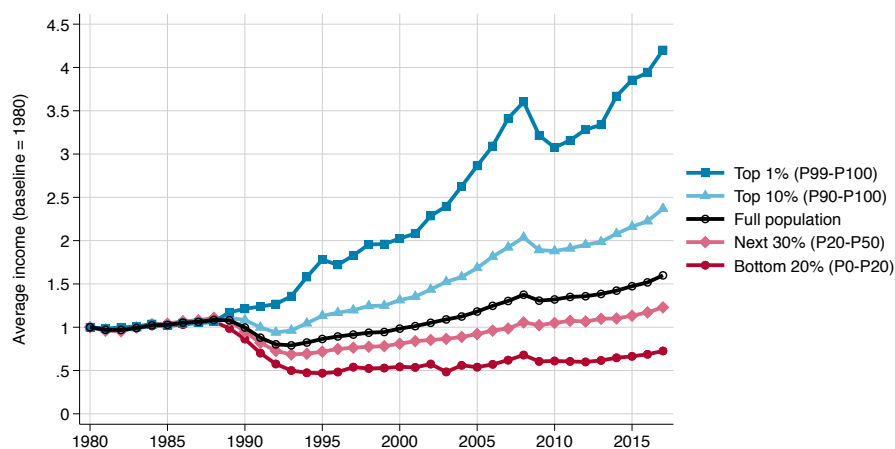
(b) Cumulated growth by income group: Western Europe



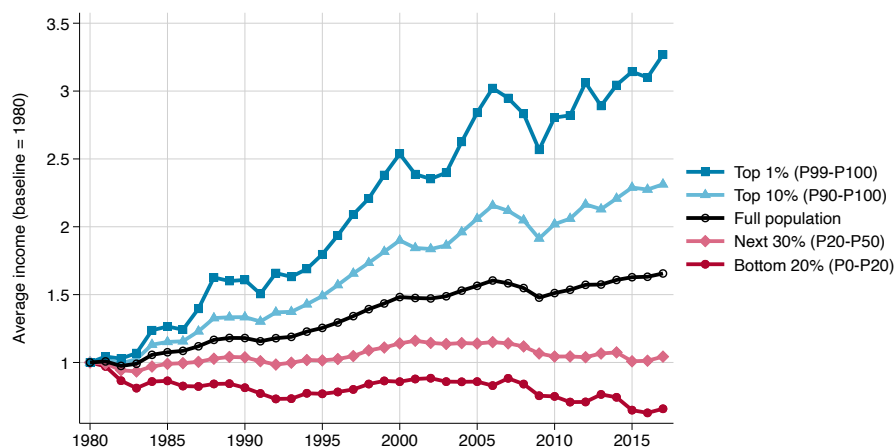
Source: Authors' computations combining surveys, tax data and national accounts for Europe. [Piketty, Saez, and Zucman, 2018](#) for the US. *Notes:* This figure shows the evolution of the average pretax income of the top 1% (p99p100), the top 10% (p90p100), the bottom 20% (p0p20), the next 30% (p20p50) and the average regional income relative to 1980, in Eastern Europe, Northern, Western Europe and the US. The unit of observation is the adult individual aged 20 or above. Income is split equally among spouses. Incomes measured at purchasing power parity. See Appendix Table A.6 for the composition of European regions and country-by-country data sources.

Figure A.2
Cumulated pretax income growth in Northern and Western Europe, 1980-2017

(a) Cumulated growth by income group: Eastern Europe

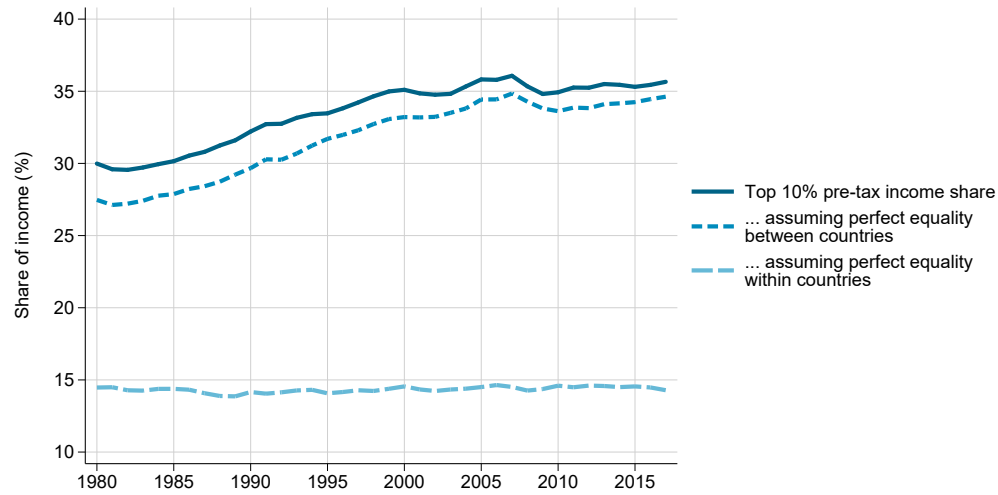


(b) Cumulated growth by income group: US



Source: Authors' computations combining surveys, tax data and national accounts for Europe. [Piketty, Saez, and Zucman, 2018](#) for the US. *Notes:* This figure shows the evolution of the average pretax income of the top 1% (p99p100), the top 10% (p90p100), the bottom 20% (p0p20), the next 30% (p20p50) and the average regional income relative to 1980, in Eastern Europe, Northern, Western Europe and the US. The unit of observation is the adult individual aged 20 or above. Income is split equally among spouses. Incomes measured at purchasing power parity. See Appendix Table A.6 for the composition of European regions and country-by-country data sources.

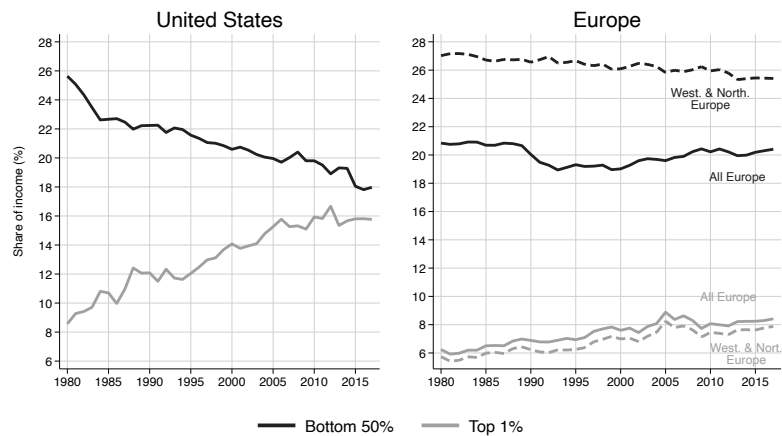
Figure A.3
Cumulated pretax income growth in Eastern Europe and the US, 1980-2017



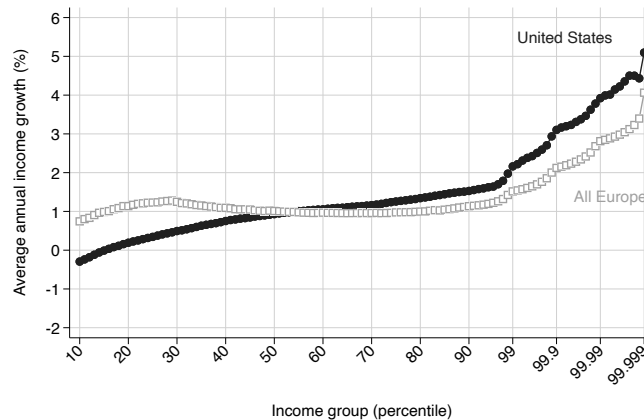
Source: Authors' computations combining surveys, tax data and national accounts. *Notes:* Incomes are measured at Purchasing Power Parity in real 2017 Euros. PPP Euro 1 = PPP\$ 1.3. The unit of observation is the adult individual aged 20 or above. See Appendix Table A.6 for the composition of European regions.

Figure A.4
Top 10% pre-tax income share in Europe: Geographical decomposition

(a) Top 1% and Bottom 50% posttax income shares in Europe and the US

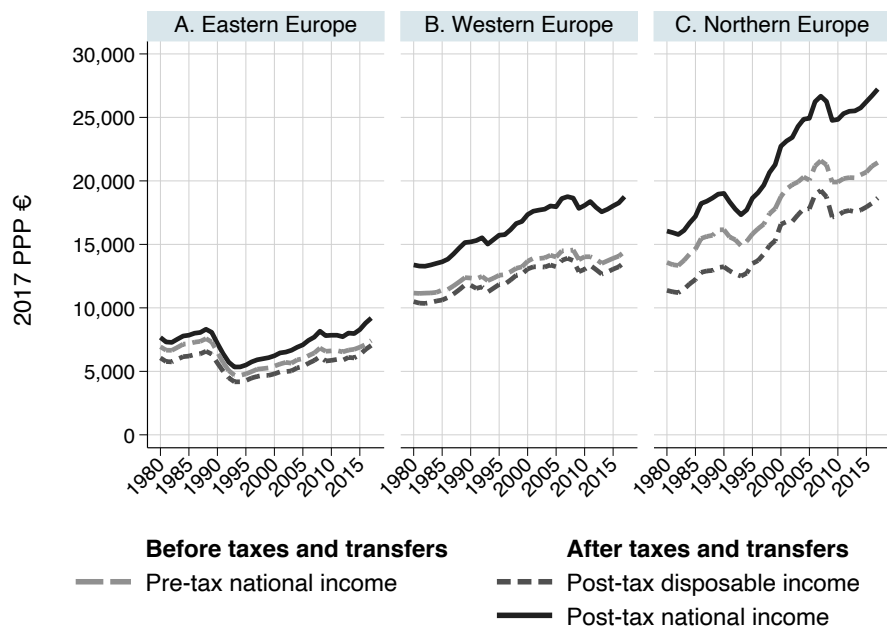


(b) Average annual posttax income growth by percentile, 1980-2017



Source: Authors' computations combining surveys, tax data and national accounts. *Notes:* The figure compares the evolution of posttax income inequality in Europe and the United States between 1980 and 2017. The top panel compares the share of posttax income received by the bottom 50% to that received by the top 1% of the regional population. The bottom panel plots the average annual posttax income growth rate by percentile, with a further decomposition of the top percentile. Figures for the US come from [Piketty, Saez, and Zucman \(2018\)](#). Figures for Europe are aggregated using market exchange rates. The unit of observation is the adult individual aged 20 or above. Income is split equally among spouses. See Appendix Table A.6 for the composition of European regions.

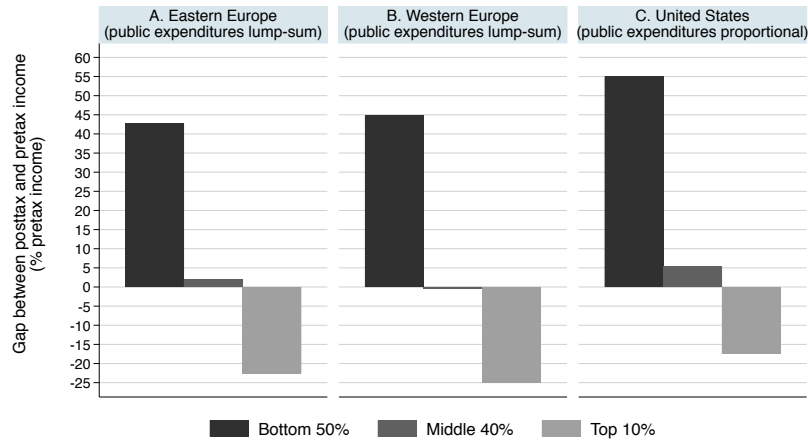
Figure A.5
The distribution of posttax income growth in Europe and the United States, 1980-2017



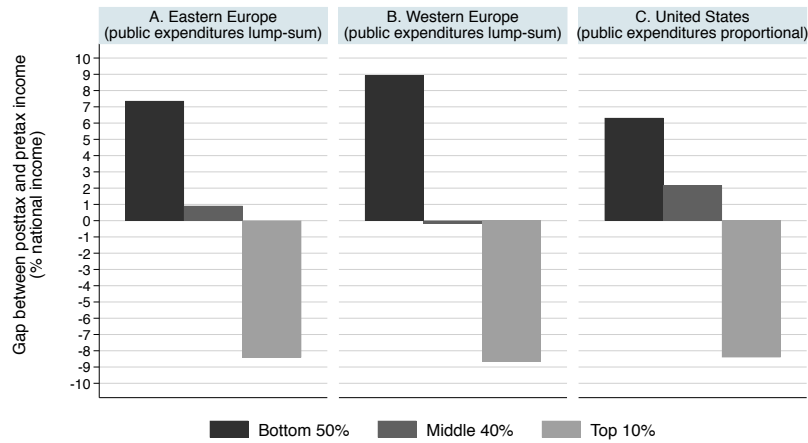
Source: Authors' computations combining surveys, tax data and national accounts for European countries and [Piketty, Saez, and Zucman, 2018](#) for the US. *Notes:* Incomes are measured at Purchasing Power Parity in real 2017 Euros. PPP Euro 1 = PPP\$ 1.3. The unit of observation is the adult individual aged 20. See Appendix Table [A.6](#) for the composition of European regions.

Figure A.6
Bottom 50% incomes in Europe, 1980-2017

(a) Net redistribution (% of group average income)



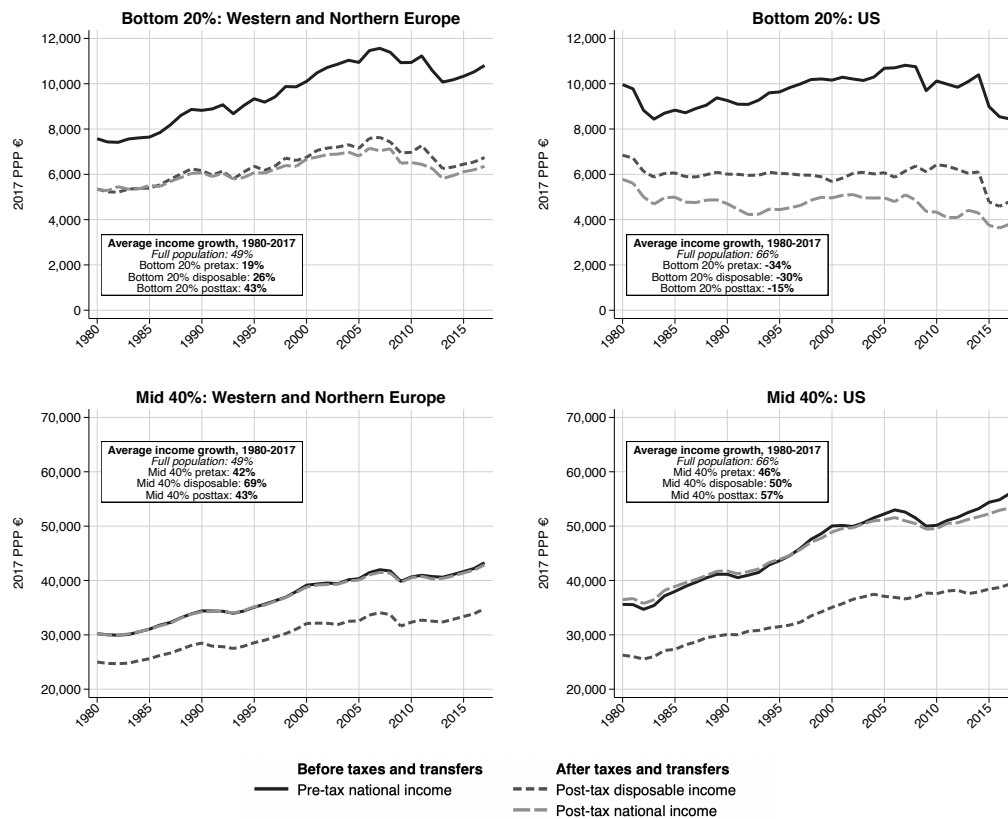
(b) Net redistribution (% of national income)



Source: Authors' computations combining surveys, tax data and national accounts for European countries and [Piketty, Saez, and Zucman, 2018](#) for the US. *Notes:* The unit of observation is the adult individual aged 20. Income is split equally among spouses. See Appendix Table A.6 for the composition of European regions.

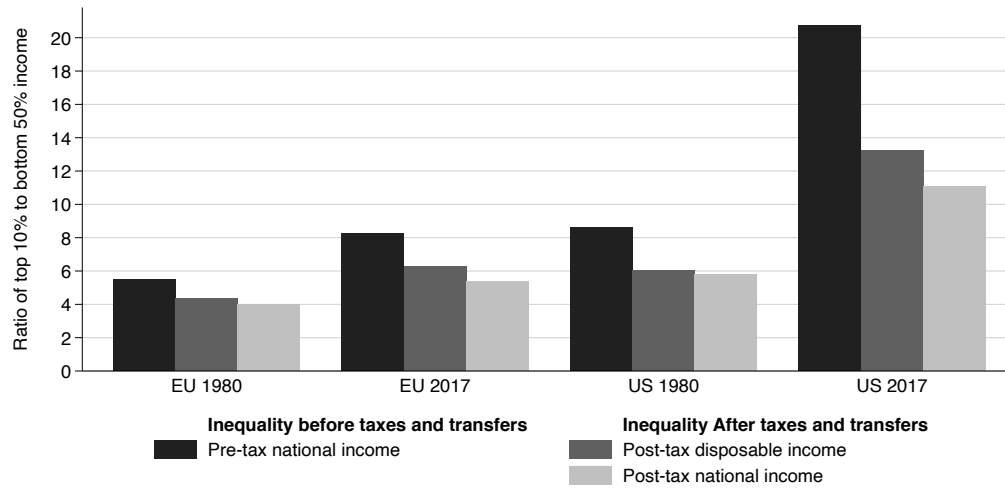
Figure A.7

Robustness Check: Distribution of Public Expenditures using Polar Assumptions for Europe and the US



Source: Authors' computations combining surveys, tax data and national accounts for European countries and [Piketty, Saez, and Zucman, 2018](#) for the US. *Notes:* Incomes are measured at Purchasing Power Parity in real 2017 Euros. PPP Euro 1 = PPP\$ 1.3. The unit of observation is the adult individual aged 20. See Appendix Table [A.6](#) for the composition of European regions.

Figure A.8
Middle 40% and Bottom 20% incomes in Europe and the US, 1980-2017



Source: Authors' computations combining surveys, tax data and national accounts for European countries and [Piketty, Saez, and Zucman, 2018](#) for the US. *Notes:* The unit of observation is the adult individual aged 20. Indicators are population weighted. European inequality estimates contain all Western, Northern and Eastern European countries. See Appendix Table [A.6](#) for more details.

Figure A.9
Redistribution in Europe and the US, 1980-2017

Table A.1
5-fold cross validation mean relative error on the average by percentile when imputing pre-tax and post-tax incomes from different concepts using our benchmark machine learning algorithm

	predictor	predicted concept			
		pretax income (broad equal-split)	pretax income (narrow equal-split)	posttax income (broad equal-split)	posttax income (narrow equal-split)
consumption	equal-split (broad)	9.9%	11.0%	8.4%	11.1%
	per capita	8.7%	11.1%	9.5%	12.0%
	households	9.2%	10.8%	7.9%	10.2%
	OECD scale	9.7%	10.4%	8.8%	11.7%
	square root scale	9.3%	10.7%	8.2%	11.7%
pretax income	equal-split (broad)	n/a	3.3%	5.8%	6.0%
	equal-split (narrow)	2.9%	n/a	5.6%	4.7%
	per capita	3.7%	5.1%	6.3%	6.4%
	households	3.9%	4.8%	7.2%	6.7%
	OECD scale	2.4%	3.8%	6.2%	6.2%
	square root scale	2.7%	4.1%	6.4%	6.5%
posttax income	equal-split (broad)	5.6%	6.4%	n/a	4.3%
	equal-split (narrow)	5.3%	4.8%	3.9%	n/a
	per capita	6.8%	7.6%	3.6%	5.5%
	households	6.4%	7.0%	3.9%	5.5%
	OECD scale	5.7%	6.5%	2.2%	4.5%
	square root scale	5.6%	6.5%	2.7%	4.7%

Source: authors' computations. *Note:* Error calculated only for the top 80% of the distributions to avoid problems of denominator near zero. The algorithm is XGBoost's implementation of boosted regression trees using $\eta = 0.1$ (Chen and Guestrin, 2016). Auxiliary variables included in the model are: regional dummies, average national income per adult (PPP), share of household with size 1 to 6, gross saving rate (% of GDP), overall social expenditures (% of GDP), top marginal income tax rate, income tax revenue (% of GDP), overall tax revenue (% of GDP), share of population by 10-year age bands and sex, corporate tax rate, VAT tax rate. *Interpretation:* When imputing pre-tax income per equal-split adult (broad) from consumption per household, the mean relative error for the average income of a given percentile is 9.2%.

Table A.2

5-fold cross validation mean relative error on the average by percentile when imputing pre-tax and post-tax incomes from different concepts using a machine learning algorithm without auxiliary variables

	predictor	predicted concept			
		pretax income (broad equal-split)	pretax income (narrow equal-split)	posttax income (broad equal-split)	posttax income (narrow equal-split)
consumption	equal-split (broad)	11.1%	12.2%	10.7%	11.8%
	per capita	11.0%	12.7%	9.2%	12.1%
	households	9.9%	11.8%	9.2%	11.7%
	OECD scale	10.8%	12.5%	9.9%	12.3%
	square root scale	10.6%	12.3%	9.3%	11.9%
pretax income	equal-split (broad)	n/a	3.7%	6.3%	6.5%
	equal-split (narrow)	3.1%	n/a	5.5%	4.5%
	per capita	3.9%	5.5%	6.8%	7.6%
	households	3.7%	5.4%	7.5%	7.5%
	OECD scale	2.4%	4.2%	6.4%	6.6%
	square root scale	2.6%	4.3%	6.6%	6.7%
posttax income	equal-split (broad)	5.8%	6.4%	n/a	4.4%
	equal-split (narrow)	5.4%	4.8%	4.0%	n/a
	per capita	7.3%	7.8%	3.8%	5.8%
	households	6.6%	6.7%	3.8%	5.7%
	OECD scale	6.2%	6.5%	2.3%	4.6%
	square root scale	6.2%	6.5%	2.7%	5.0%

Source: authors' computations. *Note:* Error calculated only for the top 80% of the distributions to avoid problems of denominator near zero. The algorithm is XGBoost's implementation of boosted regression trees using $\eta = 0.1$ (Chen and Guestrin, 2016). No auxiliary variables are included in this model. *Interpretation:* When trying to impute pre-tax income per equal-split adult from consumption per household, the mean relative error for the average income of a given percentile is 11%.

Table A.3
5-fold cross validation mean relative error on the average by percentile when imputing pre-tax and post-tax incomes from different concepts using a single correction coefficient by percentile

	predictor	predicted concept			
		pretax income (broad equal-split)	pretax income (narrow equal-split)	posttax income (broad equal-split)	posttax income (narrow equal-split)
consumption	equal-split (broad)	15.2%	17.2%	10.5%	15.2%
	per capita	20.3%	23.7%	11.0%	19.3%
	households	15.9%	18.2%	11.7%	16.2%
	OECD scale	16.7%	19.1%	11.0%	16.6%
	square root scale	14.9%	17.3%	11.1%	15.3%
pretax income	equal-split (broad)	n/a	3.7%	5.9%	6.1%
	equal-split (narrow)	3.7%	n/a	6.3%	4.5%
	per capita	3.9%	5.7%	6.7%	7.2%
	households	4.6%	5.9%	8.1%	8.0%
	OECD scale	2.4%	4.5%	6.3%	6.5%
	square root scale	2.8%	4.7%	6.6%	6.8%
posttax income	equal-split (broad)	5.8%	6.4%	n/a	4.9%
	equal-split (narrow)	6.1%	4.6%	4.8%	n/a
	per capita	6.7%	7.5%	3.9%	6.2%
	households	7.3%	7.6%	4.7%	6.6%
	OECD scale	6.1%	6.6%	2.2%	5.1%
	square root scale	6.2%	6.8%	2.7%	5.5%

Source: authors' computations. *Note:* Error calculated only for the top 80% of the distributions to avoid problems of denominator near zero. *Interpretation:* When trying to impute pre-tax income per equal-split adult from consumption per household, the mean relative error for the average income of a given percentile is 11%.

Table A.4
Total taxes and transfers in Europe and the US, 2007-2017
(% National income)

	Western & Northern Europe	Eastern Europe	United States
All taxes & social contributions	49.8%	41.0%	28.2%
Social contributions	20.7%	16.4%	7.6%
<i>Inc. contributory contributions</i>	16.6%	13.8%	5.3%
<i>Inc. non-contributory contributions</i>	4.1%	2.6%	2.2%
Taxes	29.1%	24.6%	20.7%
<i>Inc. Income & wealth taxes</i>	12.2%	5.8%	11.2%
<i>Inc. Corporate tax</i>	3.2%	3.0%	3.0%
<i>Inc. Indirect & consumption taxes</i>	13.7%	15.8%	6.5%
All non-contributory taxes & contributions	33.2%	27.2%	22.9%
All transfers	46.9%	40.5%	34.5%
Cash transfers	22.2%	17.7%	8.8%
<i>Inc. Pensions</i>	15.7%	13.5%	4.7%
<i>Inc. Unemployment & disability</i>	1.8%	0.8%	1.4%
<i>Inc. Other cash transfers</i>	4.8%	3.4%	2.7%
In-kind transfers	24.6%	22.8%	25.7%
<i>Inc. Health</i>	8.6%	5.4%	7.3%
<i>Inc. Other in-kind transfers</i>	16.0%	17.4%	18.3%

Source: Authors' computations based on national accounts and OECD social expenditures database. *Notes:* The table shows the structure of taxes and transfers in the US and Europe, as a percentage of national income. Values are population-weighted and averaged over the 2007-2017 period. See Appendix Table A.6 for the composition of European regions. See Appendix Table A.5 for separate Western Europe and Northern Europe groupings. The difference between taxes and transfers is different from the "government deficit" as measured traditionally, meaning net savings (B8n) or net lending/net borrowing (B9n), due to the exclusion of government interest payments, other current transfers and, in the case of B9n, fixed capital formation.

Table A.5
Total taxes and transfers in Europe and the US, 2007-2017
(% National income)

	Western Europe	Northern Europe	Eastern Europe	United States
All taxes & social contributions	49.5%	55.0%	41.0%	28.2%
Social contributions	21.2%	12.8%	16.4%	7.6%
<i>Inc. contributory contributions</i>	16.9%	11.7%	13.8%	5.3%
<i>Inc. non-contributory contributions</i>	4.3%	1.0%	2.6%	2.2%
Taxes	28.3%	42.3%	24.6%	20.7%
<i>Inc. Income & wealth taxes</i>	11.8%	19.5%	5.8%	11.2%
<i>Inc. Corporate tax</i>	3.1%	4.7%	3.0%	3.0%
<i>Inc. Indirect & consumption taxes</i>	13.4%	18.1%	15.8%	6.5%
All non-contributory taxes & contributions	32.6%	43.3%	27.2%	22.9%
All transfers	46.7%	49.2%	40.5%	34.5%
Cash transfers	22.4%	20.3%	17.7%	8.8%
<i>Inc. Pensions</i>	15.7%	14.5%	13.5%	4.7%
<i>Inc. Unemployment & disability</i>	1.8%	0.9%	0.8%	1.4%
<i>Inc. Other cash transfers</i>	4.8%	4.8%	3.4%	2.7%
In-kind transfers	24.4%	28.9%	22.8%	25.7%
<i>Inc. Health</i>	8.6%	9.2%	5.4%	7.3%
<i>Inc. Other in-kind transfers</i>	15.8%	19.7%	17.4%	18.3%

Source: Authors' based on national accounts. *Notes:* The table shows the structure of taxes and transfers in the US and Europe, as a percentage of national income. Values are population-weighted and averaged over the 2007-2017 period. See Appendix Table A.6 for the composition of European regions.

Table A.6
Coverage of data sources

Country	Surveys	Tax data	Undistrib. prof.	Imp. rents	Tax data source	Quality score
Western Europe						
Austria	1987-2015	1976-2015	1995-2017	1995-2017	Authors	Medium
Belgium	1985-2015	1990-2016	1985-2017	1985-2017	Decoster, Dobbelaer, and Maes (2017)	High
France	1989-2015	1980-2014	1980-2017	1980-2017	Garbinti, Goupille-Lebret, and Piketty (2018)	Very high
Germany	1981-2015	1980-2013	1991-2017	1991-2017	Bartels (2017)	High
Ireland	1980-2015	1980-2015	1995-2017	1995-2017	Jäntti et al. (2007)	High
Italy	1981-2015	1980-2009	1980-2017	1980-2017	Alvaredo and Pisano (2010)	High
Luxembourg	1985-2015	2010-2012	1995-2017	1995-2017	Authors	High
Netherlands	1983-2015	1981-2012	1980-2017		Salverda and Atkinson (2007)	High
Portugal	1980-2015	1980-2005	1995-2017	1995-2017	Alvaredo (2009)	High
Spain	1980-2015	1981-2012	1995-2017	1995-2017	Alvaredo and Saez (2010)	High
Switzerland	1982-2015	1981-2014	1990-2016		Foellmi and Martínez (2017)	High
United Kingdom	1986-2015	1981-2014	1989-2017	1990-2017	Atkinson (2007)	High
Northern Europe						
Denmark	1981-2015	1980-2010	1981-2017	1990-2017	Atkinson and Søgaard (2013)	High
Finland	1981-2015	1980-2009	1980-2017	1980-2017	Jäntti et al. (2010)	High
Iceland	2003-2014	1990-2016	2000-2014	2005-2014	Authors	High
Norway	1986-2015	1981-2011	1980-2017	1980-2017	Aaberge and Atkinson (2010)	High
Sweden	1981-2015	1980-2013	1980-2017	1980-2017	Roine and Waldenström (2010)	High
Eastern Europe						
Albania	1996-2012					Low
Bosn. & Herz.	1983-2015					Medium Low
Bulgaria	1980-2015		1995-2017			Medium
Croatia	1983-2015	1983-2013	1997-2014	2002-2012	Kump and Novokmet (2018)	High
Cyprus	1990-2015		1995-2017	1995-2017		Medium Low
Czech Republic	1980-2015	1980-2015	1993-2017	1993-2017	Novokmet (2018)	High
Estonia	1988-2015	2002-2017	1994-2017		Authors	High
Greece	1981-2015	2004-2011	1995-2017	1995-2016	Chrissis and Koutentakis (2017)	High
Hungary	1982-2015	1980-2008	1995-2017	1995-2017	Mavridis and Mosberger (2017)	High
Kosovo	2003-2017					Medium Low
Latvia	1988-2015		1994-2017	1995-2017		Medium
Lithuania	1988-2015		1995-2017	1995-2017		Medium
Malta	2006-2015			2000-2017		Medium Low
Moldova	1988-2017					Low
Montenegro	1983-2014					Medium Low
Macedonia	1983-2015					Medium Low
Poland	1983-2016	1983-2015	1995-2017	1995-2016	Bukowski and Novokmet (2017)	High
Romania	1989-2016	2013	1995-2017	2004-2013	Oancea, Andrei, and Pirjol (2017)	Medium
Serbia	1983-2015	2017	2010-2011	1997-2011	Authors	Medium
Slovakia	1980-2015		1995-2017	1995-2017		Medium
Slovenia	1987-2015	1991-2012	1995-2017	1995-2017	Kump and Novokmet (2018)	High

Table A.7
European countries and US performance in reaching SDG 10.1, 1980-2017

	1980-2017			2007-2017		
	Bottom 40%	Average	Difference	Bottom 40%	Average	Difference
Albania	1.4 %	1.4 %	0.0 p.p.	2.0 %	1.6 %	0.4 p.p.
Austria	1.2 %	1.2 %	0.0 p.p.	-0.1 %	-0.0 %	-0.1 p.p.
Bosn. & Herz.	3.3 %	3.9 %	-0.7 p.p.	1.6 %	1.6 %	0.1 p.p.
Belgium	0.8 %	1.1 %	-0.3 p.p.	-0.2 %	0.2 %	-0.4 p.p.
Bulgaria	0.6 %	1.9 %	-1.3 p.p.	3.0 %	3.2 %	-0.1 p.p.
Switzerland	0.6 %	0.6 %	-0.0 p.p.	0.5 %	0.1 %	0.5 p.p.
Cyprus	1.5 %	1.7 %	-0.2 p.p.	1.4 %	-1.7 %	3.1 p.p.
Czech Republic	0.1 %	0.9 %	-0.8 p.p.	0.8 %	1.0 %	-0.2 p.p.
Germany	0.1 %	0.9 %	-0.8 p.p.	-0.2 %	0.9 %	-1.1 p.p.
Denmark	1.0 %	1.4 %	-0.4 p.p.	-0.6 %	0.2 %	-0.9 p.p.
Estonia	1.0 %	1.7 %	-0.7 p.p.	2.0 %	0.7 %	1.2 p.p.
Spain	1.2 %	1.3 %	-0.1 p.p.	0.2 %	0.3 %	-0.1 p.p.
Finland	1.0 %	1.4 %	-0.4 p.p.	-1.4 %	-0.7 %	-0.7 p.p.
France	1.4 %	1.0 %	0.4 p.p.	0.4 %	0.1 %	0.4 p.p.
United Kingdom	1.2 %	1.6 %	-0.4 p.p.	1.2 %	0.1 %	1.0 p.p.
Greece	-0.4 %	-0.1 %	-0.2 p.p.	-4.9 %	-3.7 %	-1.2 p.p.
Croatia	-0.0 %	0.1 %	-0.1 p.p.	1.0 %	0.1 %	0.9 p.p.
Hungary	-0.6 %	1.0 %	-1.6 p.p.	0.4 %	1.1 %	-0.7 p.p.
Ireland	2.6 %	2.8 %	-0.3 p.p.	-0.1 %	0.3 %	-0.4 p.p.
Iceland	1.7 %	1.4 %	0.3 p.p.	2.7 %	0.7 %	2.1 p.p.
Italy	-0.5 %	0.4 %	-0.9 p.p.	-2.3 %	-1.1 %	-1.1 p.p.
Kosovo				3.4 %	2.9 %	0.5 p.p.
Lithuania	0.1 %	1.4 %	-1.3 p.p.	1.2 %	1.9 %	-0.7 p.p.
Luxembourg	1.4 %	1.8 %	-0.4 p.p.	-4.2 %	-3.9 %	-0.4 p.p.
Latvia	0.3 %	1.1 %	-0.8 p.p.	3.4 %	1.2 %	2.3 p.p.
Moldova	-1.4 %	-0.6 %	-0.8 p.p.	3.9 %	3.2 %	0.8 p.p.
Montenegro	-1.0 %	-0.6 %	-0.4 p.p.	1.1 %	1.5 %	-0.4 p.p.
North Macedonia	-0.8 %	-0.0 %	-0.8 p.p.	3.6 %	2.0 %	1.6 p.p.
Malta	1.9 %	2.3 %	-0.4 p.p.	1.9 %	2.6 %	-0.7 p.p.
Netherlands	0.3 %	0.8 %	-0.5 p.p.	-0.2 %	-0.1 %	-0.2 p.p.
Norway	1.5 %	1.7 %	-0.2 p.p.	-0.2 %	-0.2 %	0.0 p.p.
Poland	0.7 %	1.8 %	-1.2 p.p.	2.7 %	2.7 %	-0.0 p.p.
Portugal	0.6 %	1.3 %	-0.7 p.p.	0.3 %	-0.0 %	0.3 p.p.

Table A.7
European countries and US performance in reaching SDG 10.1, 1980-2017

	1980-2017			2007-2017		
	Bottom 40%	Average	Difference	Bottom 40%	Average	Difference
Romania	-0.4 %	1.4 %	-1.8 p.p.	3.4 %	2.7 %	0.7 p.p.
Serbia	-1.8 %	-0.2 %	-1.6 p.p.	-1.7 %	1.0 %	-2.7 p.p.
Sweden	1.3 %	1.8 %	-0.5 p.p.	0.9 %	1.0 %	-0.1 p.p.
Slovenia	-0.5 %	0.3 %	-0.8 p.p.	-0.4 %	-0.1 %	-0.3 p.p.
Slovakia	1.2 %	1.4 %	-0.2 p.p.	2.9 %	1.8 %	1.2 p.p.
United States	-0.3 %	1.4 %	-1.6 p.p.	-1.4 %	0.4 %	-1.9 p.p.

Source. Authors' computations combining surveys, tax data and national accounts. *Notes.* The table shows the average annual real growth of the pretax income of the bottom 40%, the average annual real growth of the average national income per adult, and the percentage points difference between the two growth rates over the 1980-2017 and 2007-2017 periods. Negative differences imply that the income of the bottom 40% grew slower than the average national income. The unit of observation is the adult individual aged 20 or above.