



Hierarchy and the power-law income distribution tail

Blair Fix¹

Received: 22 April 2018 / Accepted: 9 July 2018
© Springer Nature Singapore Pte Ltd. 2018

Abstract

What explains the power-law distribution of top incomes? This paper tests the hypothesis that it is firm hierarchy that creates the power-law income distribution tail. Using the available case-study evidence on firm hierarchy, I create the first large-scale simulation of the hierarchical structure of the US private sector. Although not tuned to do so, this model reproduces the power-law scaling of top US incomes. I show that this is purely an effect of firm hierarchy. This raises the possibility that the ubiquity of power-law income distribution tails is due to the ubiquity of hierarchical organization in human societies.

Keywords Power law · Income distribution · Firm hierarchy · Economic modeling

Introduction

In the late nineteenth century, Pareto [1] discovered that top incomes could be modeled with a power-law distribution. This scaling behavior meant that the income distribution tail could be approximated by the simple probability function:

$$p(x) = \frac{c}{x^\alpha}. \quad (1)$$

Here, $p(x)$ is the probability of finding an individual (in the tail) with income x , c is a constant,¹ and α is the scaling exponent, which captures the ‘fatness’ of the income distribution tail. The beauty of a power law is its simplicity. The important

¹ The constant c is equal to $(\alpha - 1)/(x_{\min})^{1-\alpha}$, where $x_{\min}$ is the lower bound of the power law (i.e., the beginning of the income distribution tail).

Electronic supplementary material The online version of this article (<https://doi.org/10.1007/s42001-018-0019-8>) contains supplementary material, which is available to authorized users.

✉ Blair Fix
blairfix@gmail.com

¹ York University, Toronto, ON, Canada

properties of the distribution tail are captured by a single parameter—the power-law exponent. Since Pareto’s initial discovery, the power-law scaling of top incomes has been re-confirmed many times (for a non-exhaustive list, see [2–12]).

What causes this nearly universal behavior? Is there a universal generation mechanism at work? Over the century since Pareto’s landmark discovery, many power-law generation mechanisms have been suggested to explain the scaling of top incomes. While the various mechanisms (reviewed below) differ in their mathematical properties, most are united by a shared focus on the stochastic growth of individual income.

This paper investigates a very different explanation for the power-law scaling of top incomes. I test the hypothesis that it is *firm hierarchy* that creates the power-law tail. This approach was first proposed by Lydall [13], who used a simple model to show that the hierarchical structure of firms could create a power-law distribution. At the time, Lydall’s work was largely speculative since little was known about the internal structure of firms. However, in the last two decades the empirical study of firm hierarchy has blossomed (for case studies, see [14–20]; for aggregate studies, see [21–29]). Enough evidence now exists that we can begin to explore the distributional consequences of hierarchy. To conduct this investigation, I use the existing case-study evidence to build a large-scale simulation of firm hierarchy. This model generalizes the trends found in case-study firms to create the first simulation of the hierarchical structure of the US private sector.

I verify the accuracy of the hierarchy model, in two ways. I first compare the model’s income distribution to that of the USA. I find that the hierarchy model does a reasonably good job of reproducing the key properties of US income distribution. Importantly, the model produces (without tuning it to do so) a power-law tail that is statistically identical to US empirical data. Next, I test a key feature of the hierarchy model—that top-earning individuals should be concentrated in large firms. I find that the model’s prediction is consistent with the available US evidence.

Having established the model’s accuracy, I then use the model to investigate the distributional effects of hierarchy. I find that it is firm hierarchy alone (and not any of the other income dispersion factors included in the model) that is responsible for generating the power-law tail. To summarize, the hierarchy model suggests that it is firm hierarchy (and its associated properties) that creates the power-law scaling of top incomes. This finding has important implications for both the empirical and theoretical study of income distribution. On the empirical side, these results indicate that the income effects of hierarchy are significant and need to be studied in more detail. On the theoretical side, these results suggest that hierarchy is a plausible mechanism for generating the power-law scaling of top incomes. This raises the possibility that the ubiquity of power-law income distribution tails is due to the ubiquity of hierarchical organization in human societies.

The paper is laid out as follows: “[Power-law generation mechanisms](#)” reviews the different mechanisms for generating power-law distributions. “[A firm hierarchy model](#)” outlines (in non-technical terms) the basic properties of the hierarchy model. (For a technical discussion, see the Online Appendices). “[Testing the hierarchy model: macro predictions](#)” and “[Testing the hierarchy model: micro predictions](#)” test the model against empirical data. “[Isolating the distributional role of firm hierarchy](#)” demonstrates that it is firm hierarchy (alone) that is responsible for creating the model’s power-law tail, and “[How hierarchy generates the power-law tail](#)” analyzes

how this is achieved. I conclude, in “[Discussion](#)” and “[Conclusions](#)”, with a discussion of the significance of these results and propose avenues for future research.

Power-law generation mechanisms

I review here in non-technical terms the various mechanisms for generating power-law distributions, with an emphasis on those that have been applied to modeling income. For a good technical review of these mechanisms, see [30–32]. One way to generate power laws is through a stochastic, multiplicative growth process. This mechanism was identified by Gibrat [33], but was first applied to income distribution by Champenowne [34], followed by many others [35–38]. The basic idea is that individual income is subjected to stochastic, multiplicative ‘shocks’. Under the condition that these multiplicative shocks are scale free (they do not depend on income size) and there is a minimum (reflective) lower bound on income, this process will produce a power-law distribution of income.

Closely related to this process is the mechanism of ‘preferential attachment’, sometimes called the ‘rich get richer’. Developed independently by Yule [39], Simon [40], Price [41] and Barabasi and Albert [42], this process involves stochastic addition with conditional probability. It is most easily applicable to the distribution of wealth (not income). We imagine a society in which units of wealth are added at random. If the probability of an individual receiving an additional unit of wealth is proportional to his/her existing wealth, the result (after many iterations) will be a power-law distribution of wealth.

In both multiplicative growth and preferential attachment models, the source for stochastic changes in income/wealth is left unexplained. More recently, econophysicists have developed a more sophisticated class of model that attempts to explain these ‘shocks’ in terms of the random exchange of money between pairs of interacting agents [11, 43–51]. These ‘kinetic-exchange models’ draw explicitly on the statistical mechanics of gases. Agents exchange money much like gas particles exchange kinetic energy. Given certain assumptions about these interactions, kinetic-exchange models can produce a power-law distribution.

Lastly, a very simple way to produce a power law is to *exponentially* transform an *exponential* distribution. This is the mechanism by which Lydall [13] showed that firm hierarchy could create a power law. If a firm hierarchy has a constant ‘span of control’ (the number of subordinates controlled by each superior), then relative employment will decrease exponentially by rank. If, at the same time, income increases exponentially by rank, the result will be a power-law distribution of income. (For a technical discussion, see “[How hierarchy generates the power-law tail](#)”).

The advantage of this hierarchy mechanism is that it ties income to *institutions*. This means that the power-law distribution of top incomes is given an explicit institutional basis—something that is important when it comes to policy discussions about how to reduce inequality. The disadvantage of the hierarchy mechanism is that relatively little is known about firm hierarchical structure. This means that the distributional consequences of hierarchy are little understood. This paper aims to remedy this situation by using the available empirical data to build a large-scale hierarchical model

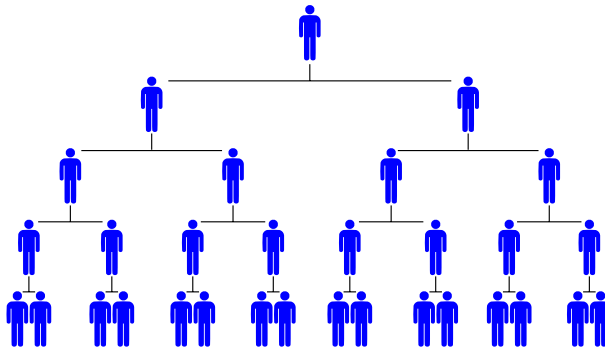


Fig. 1 A branching hierarchy. This figure shows an idealized branching hierarchy in which each superior has two subordinates. This superior/subordinate ratio—often called the span of control—can be used to mathematically describe the hierarchy. Starting from the bottom rank, each consecutive rank *decreases* in size by a factor of the span of control. Unlike employment, we expect income to *increase* with hierarchical rank

of the US private sector. This model generalizes the sparse firm hierarchy empirical data to allow the first quantitative study of the distributional effects of hierarchy.

A firm hierarchy model

The firm hierarchy model (herein the ‘hierarchy model’) is based on the hypothesis that human institutions are hierarchically organized, and that hierarchical rank plays a key role in determining income. The starting point for my approach is the seminal work of Simon [52] and Lydall [13]. In the late 1950s, Simon and Lydall both developed simple models that focused on the branching structure of firm hierarchies. The distinguishing feature of a branching hierarchy is that each superior has control over *multiple* subordinates (see Fig. 1).

Simon and Lydall both showed how branching hierarchical structure could explain regularities in income distribution. Simon used a simple hierarchical model of the firm to explain the observed scaling between CEO pay and firm sales [53]. Lydall showed how firm hierarchy could lead to a power-law distribution of top income (as discussed above). This paper draws on the work of Simon and Lydall, but updates their model in light of recent empirical work.

Both Simon and Lydall assumed a constant span of control within firms. (The span of control is the number of subordinates per superior). However, case-study evidence indicates that the span of control is *not* constant within firms, but instead tends to increase with hierarchical rank (see Online Appendix B). Simon and Lydall also assumed that the average income ratio between adjacent hierarchical ranks was constant. Again, case-study evidence suggests that this is not quite true. Like the span of control, the pay ratio between ranks also tends to increase with rank.

The key difference between my approach and that of Simon and Lydall is that I take full advantage of modern computational power to build a large-scale, stochastic

simulation. In contrast, Simon and Lydall used simple analytic methods. Simulation allows investigation that would otherwise be impossible with a purely analytic approach.

Modeling goals and methods

Unlike the power-law generation models discussed in “[Power-law generation mechanisms](#)”, my hierarchy model is not *designed* to produce a power law. Rather, it is designed to match the available firm-level evidence, with the intention of generalizing this evidence to investigate the distributional effects of firm hierarchy. The *hope* is that the resulting model will produce a power law that matches macro-level data, but there is no guarantee that it will.

In principle, we could directly investigate the income effects of firm hierarchy using empirical data (with no need for a model). However, the available firm-level evidence is too sparse to draw conclusions about the wider distributional role of firm hierarchy. The purpose of the hierarchy model is to investigate what is *implied* by the available firm-level data. The model takes the limited firm-hierarchy evidence that does exist and fits trends and parameterized distributions to it. I then use the model algorithm (outlined in detail in Online Appendices D and E) to extrapolate these trends to create a large-scale simulation of the economy. The resulting model is *entirely* dependent on the input, firm-level data. I do not tune the model to reproduce macro-level results. Therefore, the model output is purely what is implied by generalizing the trends found in input data.

The hierarchy model is built on a tripartite income-dispersion classification scheme that allows for three sources of income dispersion (see Fig. 2):

- **Source 1:** Income dispersion *between* hierarchical levels of each firm (*inter-hierarchical* dispersion);
- **Source 2:** Income dispersion *within* hierarchical levels of each firm (*intra-hierarchical* dispersion);
- **Source 3:** Income dispersion *between* different firms (*inter-firm* dispersion).

Inter-firm and intra-hierarchical level dispersion are not explained by the model. (In the jargon of economic modeling, these dispersion sources are *exogenous*). In contrast, inter-hierarchical dispersion is *partially* explained by the model. It is explained in the sense that it is not *ex nihilo*—this dispersion does not come from nowhere. The model contains firms that have a specific hierarchical structure of employment and pay. However, the reason for this hierarchical structure is *not* explained by the model. Rather, hierarchical structure is determined from regressions on case-study data, in conjunction with firm-level data from the Compustat and Execucomp databases.

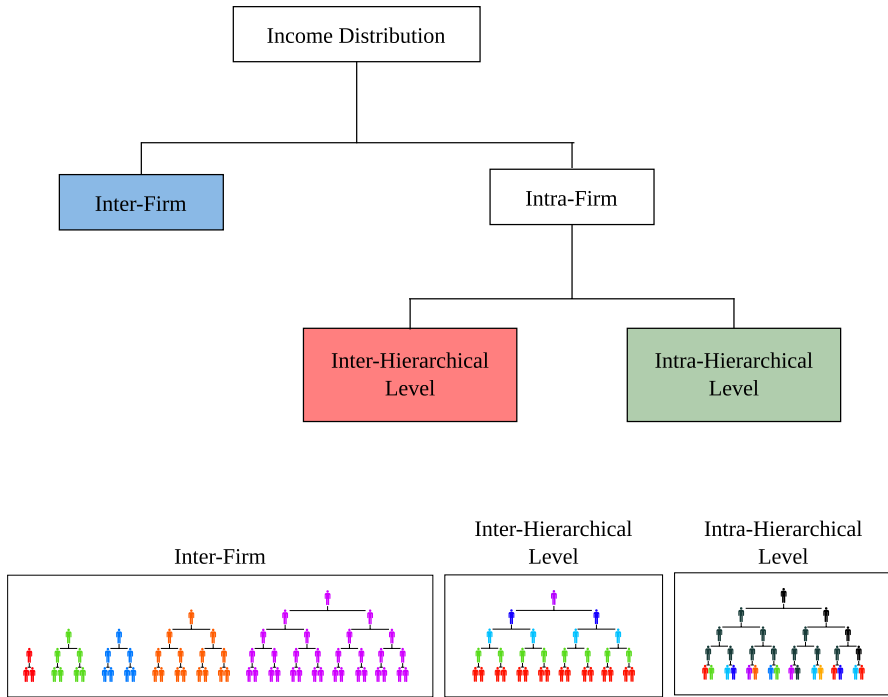


Fig. 2 A tripartite division of income distribution. This figure illustrates the income distribution grouping scheme used by the hierarchy model. The model allows for three sources of income dispersion. *Inter-firm* dispersion consists of differences in (average) pay between firms. Within each firm, there are two further sources of dispersion. *Inter-hierarchical* level dispersion consists of differences in (average) pay between hierarchical levels, while *intra-hierarchical* level dispersion consists of differences in pay within each hierarchical level

Modeling the USA

The model is designed to study the hierarchical structure of the US private sector as it was (on average) over the years 1992–2015. At the highest level of abstraction, the model has three parts. First, the model creates a firm-size distribution that dictates how many firms of a given size will exist. Second, for each firm in this distribution, the model creates a hierarchical structure. This means the model determines how many ranks will exist, and how many individuals will occupy each hierarchical rank. Lastly, the model uses each of the three dispersion sources (outlined above) to stochastically generate an income for every individual in every firm. I review here the most important elements of each step. A technical discussion can be found in the Online Appendices.

Step 1: Create a firm-size distribution The first step of the model is to generate a distribution of firm sizes. The available evidence suggests that national firm-size distributions can be modeled by a power law [54–56]. Under this assumption, the

probability of finding a firm of size x is proportional to $x^{-\alpha}$, where α is a constant. I model the US firm-size distribution with 1 million firms distributed according to a discrete power-law distribution with exponent $\alpha = 2.01$ (see Online Appendix E). This may seem like the model uses one power law (the firm size distribution) to create another (the distribution of income). However, this is not the case. Without hierarchy, the model will not create a power-law distribution (see Fig. 6).

Step 2: Endow firms with hierarchical structure The hierarchy model captures only the *aggregate* hierarchical structure of firms. That is, I model the number of employees in each hierarchical level, not the exact chain of command. I base the model on a number of recent case studies that have documented the aggregate hierarchical structure of firms in various developed countries (see Online Appendix B). From this data, I make generalizations about the hierarchical structure of firms. The evidence suggests that the span of control (the ratio between adjacent hierarchical levels) increases exponentially with hierarchical rank.

For simplicity, all firms in the model have the same hierarchical structure—they are governed by the same span of control function. However, since there is a great deal of uncertainty in this function, I run the model many times. Each different model run uses a slightly different span of control function, determined by resampling from case-study data. The result is that the hierarchical structure of firms varies stochastically between different model runs, allowing us to capture uncertainty in the underlying empirical data. For more details, see Online Appendices D and E.

Step 3: Endow individuals with income After each firm has a hierarchical structure, the model assigns every individual an income. Because the model has three dispersion mechanisms, this step has three components, outlined below.

Step 3A: Generate inter-hierarchical level dispersion In the model, firm hierarchical pay is constructed from the bottom up. Starting from the bottom rank, I define a function that determines the rate at which pay increases by hierarchical rank. This function is informed by case-study data (see Online Appendix B). Unlike hierarchical employment structure, each modeled firm is given a *different* hierarchical pay structure. The process of assigning different hierarchical pay structure to each firm is informed by firm-level data in the Compustat database. (See Online Appendix C for a detailed discussion of the Compustat data).

Before running the full simulation, I fit the hierarchy model to Compustat data for real-world American firms. Compustat (in conjunction with Execucomp) provides data on CEO pay, average pay, and firm employment. Assuming the CEO occupies the top hierarchical level, we can use this information to model the hierarchical pay structure of each Compustat firm. Once this is complete, we have an indication of how hierarchical pay should vary across firms. The model's main simulation is then informed by this variation. The result is a unique hierarchical pay structure for each firm. For more details, see Online Appendices D and E.

Step 3B: Generate inter-firm dispersion I create inter-firm income dispersion by varying (average) pay in the bottom hierarchical level of each firm. This variation

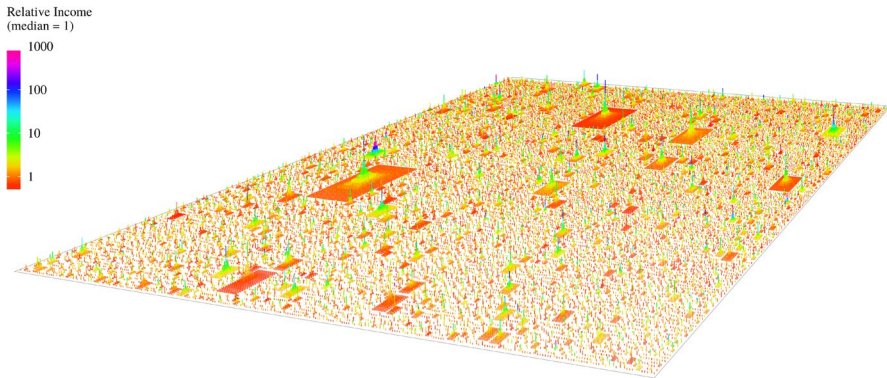


Fig. 3 A Landscape view of the hierarchy model. This figure visualizes the US hierarchy model as a landscape of three-dimensional firms. Each pyramid represents a single firm, with size indicating the number of employees and height corresponding to the number of hierarchical levels. If you look closely, you will see vertical lines corresponding to individuals. Income (relative to the median) is indicated by color. This visualization has 20,000 firms—a small sample of the actual model, which uses 1 million firms

is informed by firm-level data in the Compustat database. As discussed in Step 3A, prior to running a full-scale simulation, I fit the model to firms in the Compustat database. After having fit hierarchical pay, I use this information to estimate how base-level pay varies across these firms. This variation then informs the model's main simulation. For more details, see Online Appendices D and E.

Step 3C: Generate intra-hierarchical level dispersion The last step is to model the income dispersion *within* the hierarchical levels of each firm. The available case-study evidence suggests that income dispersion within hierarchical levels is roughly constant across all hierarchical levels (see Online Appendix B). To simplify the model, I further assume that intra-hierarchical level dispersion is constant across all firms. Informed by case-study data, I use a single parameterized distribution to randomly generate income dispersion within all hierarchical levels of every firm. For more details, see Online Appendices D and E.

Visualizing the US hierarchy model

To give an intuitive understanding of what the hierarchy model 'looks' like, Fig. 3 shows a landscape view of the model's structure. Each pyramid represents a different hierarchically organized firm. The size of each pyramid corresponds to the number of employees, height represents hierarchical level, and color represents relative income.

This figure highlights the main characteristics of the model. The firm power-law distribution is clearly visible. The vast majority of firms are small, but there are a few behemoths. Inter-firm income dispersion and inter-hierarchical level income dispersion are also visible, while intra-hierarchical level income dispersion appears negligible. Lastly, top incomes are concentrated in upper hierarchical levels, and

consequently occur mostly in larger firms. These facts, which are qualitatively visible here, become more clear as we analyze the model results in quantitative terms.

Testing the hierarchy model: macro predictions

The purpose of the hierarchy model is to study the hierarchical structure of the US private sector. The first step, then, is to make sure that the model produces realistic results. To that end, Fig. 4 compares the model's aggregate income distribution to US empirical data. Although the model aims only to capture the private sector (not government), I compare the model's results to macro-level data for the entire USA. I do this because the most reliable income distribution data (from the IRS) does not differentiate between the private and public sector.

Even though the model is an extrapolation from a limited set of data, it does a reasonably accurate job of reproducing the US distribution of income. Note, though, that the model underestimates US income inequality, both in terms of the Gini index (Fig. 4a) and the income share of the top 1% (Fig. 4b). What is the source of this discrepancy? Looking at the income probability density in Fig. 4d, it appears that the US income distribution is more 'bottom heavy' than the model. The model produces too few extremely small incomes, relative to the US. This tendency is also evident in the cumulative distribution (Fig. 4f).

Why does this discrepancy occur? I demonstrate in Online Appendix F that the discrepancy can be removed by increasing the model's inter-firm income dispersion. This suggests that the model's underestimate of US inequality is due to an underestimate of inter-firm income dispersion. My guess is that this occurs because the model is based on Compustat firm data, which is not a representative sample of the US firm population. Compustat contains data for public firms only, and as a result is biased towards large firms. I suspect that a more representative firm sample would give greater inter-firm income dispersion. I include adjusted results in the Online Appendix to show that the model is *capable* of closely reproducing the important features of US income distribution. However, I do *not* use this adjusted data for any of the proceeding analysis. The purpose of the hierarchy model is to extrapolate empirical data, warts and all.

While the model slightly misrepresents the 'body' of US income distribution, it accurately reproduces the *tail*. This is evident in the complementary cumulative distribution (Fig. 4f) in the form of virtually identical model and empirical slopes in the right tail. This is important because it is the tail of the distribution (particularly, its power-law properties) that we are interested in studying. When plotted on a log-log scale (as in Fig. 4f), a power-law tail is visually evident as a straight line in the complementary cumulative distribution.

Dating back to the work of Pareto [1], it has been common to estimate the power-law exponent by means of a linear regression on the complementary cumulative distribution. However, Clauset et al. show that this approach is inaccurate [57]. Instead, I use the more accurate maximum-likelihood method (see Online Appendix A). Estimating the power-law exponent requires making a choice about where the 'tail' of the distribution begins. I define the tail as the top 1% of

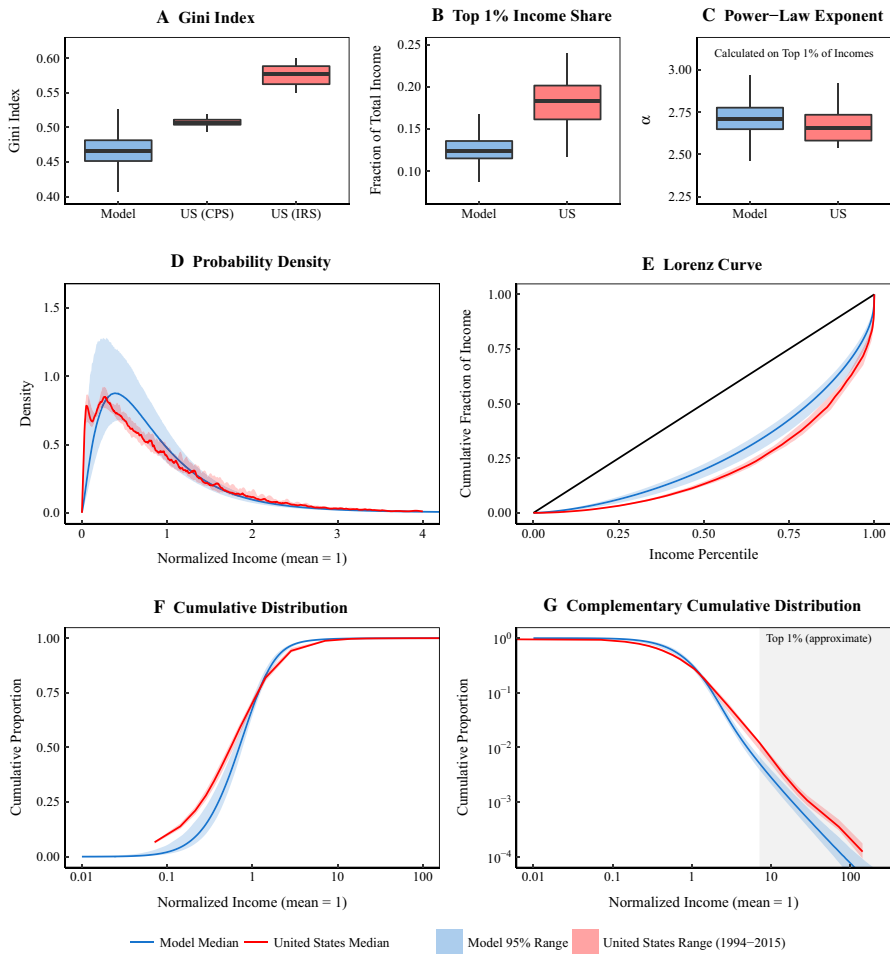


Fig. 4 Modeled income distribution vs. US data. This figure compares various aspects of the model's income distribution to US data over the years 1992–2015. **a** The Gini index, with two different US sources—the Current Population Survey (CPS) and the Internal Revenue Service (IRS). **b** The top 1% income share, using data from 17 different time series. **c** The results of fitting a power-law distribution to the top 1% of incomes (where α is the scaling exponent). **d** plots the income density curve with mean income normalized to 1 (using data from the CPS). **e–g** Use IRS data to construct the Lorenz curve, cumulative distribution, and complementary cumulative distribution (respectively). The cumulative distribution shows the proportion of individuals with income *less* than the given x value. The complementary cumulative distribution shows the proportion of individuals with income *greater* than the given x value. The shaded region shows the approximate threshold for the top 1% of incomes. For sources and methods, see Online Appendix A

incomes, a threshold that has been popularized by Piketty [58]. Figure 4c shows the results of fitting a power law to the top 1% of incomes (for methods, see Online Appendix A). Over many runs, the model produces a distribution tail with fitted power-law exponents that are very close to the exponents fitted to historical US data.

To summarize, the hierarchy model produces an income distribution that is roughly consistent with the US distribution of income. In particular, the model closely reproduces the tail of the US distribution, including its power-law properties.

Testing the hierarchy model: micro predictions

When discussing the model visualization shown in Fig. 3, I noted that top-earning individuals are clustered at the tops of *large* firms. This is a defining feature of the hierarchy model. It occurs because income scales strongly with hierarchical rank. As a result, top earners are found at the tops of large firms, because these firms have the most hierarchical levels. To my knowledge, this prediction is not made by any other model of income distribution. It is important, therefore, that we put it to the test.

To test this prediction, I look at the firm-size distribution associated with top-earning individuals. What does this mean? I take a sample of Americans with top incomes, and then record the firms associated with these individuals. I then look at the size distribution of these firms. I do the same with the model and then compare the results.

I conduct this test using data from the Forbes 400 and Execucomp. The Forbes 400 list is a definitive ranking of the 400 richest Americans, and it provides the institutional source of each individual's wealth. The caveat is that this list is a ranking by *wealth*, not income. I use the Forbes 400 as a proxy for top US incomes, under the assumption that wealth and income are strongly related. I supplement the Forbes 400 data with the 'Execucomp 500', which is composed of the 500 top-paid US executives in the Execucomp database (in each year between 1992 and 2015). The advantage of the Execucomp 500 is that it is a ranking explicitly by income. The disadvantage is that we do not know if these 500 executives are actually the top-paid US individuals.

Before conducting this test, it is instructive to know what a *null effect* would look like. If there is absolutely no relation between income and firm membership, what sort of firm-size distribution should be associated with top incomes? It turns out that for the USA, we should expect a null effect to return a roughly *log-uniform* firm-size distribution (see Online Appendix G for a derivation).

Results for the Fortune 400 and Execucomp 500 firm-size distributions are shown in the main panel of Fig. 5. These density plots represent the size distribution of firms associated with the richest 400 Americans and the 500 top-paid executives in the Execucomp database (respectively). To better visualize the distribution, I plot the density of the *logarithm* of firm size. Under this transformation, the null-effect result will appear as a *uniform* distribution. From the evidence shown in Fig. 5, we can immediately conclude that the null effect is false. There is definitely a relation between top incomes/wealth and firm size. But is it the relation that is predicted by the hierarchy model?

To find out, I conduct the same analysis on the model. I select the model's 500 top-paid individuals and record the size distribution of associated firms. The results are shown in Fig. 5 as the 'Model 500'. The model predicts a relation between top incomes and firm size that is very similar to the US empirical data. To be sure, the

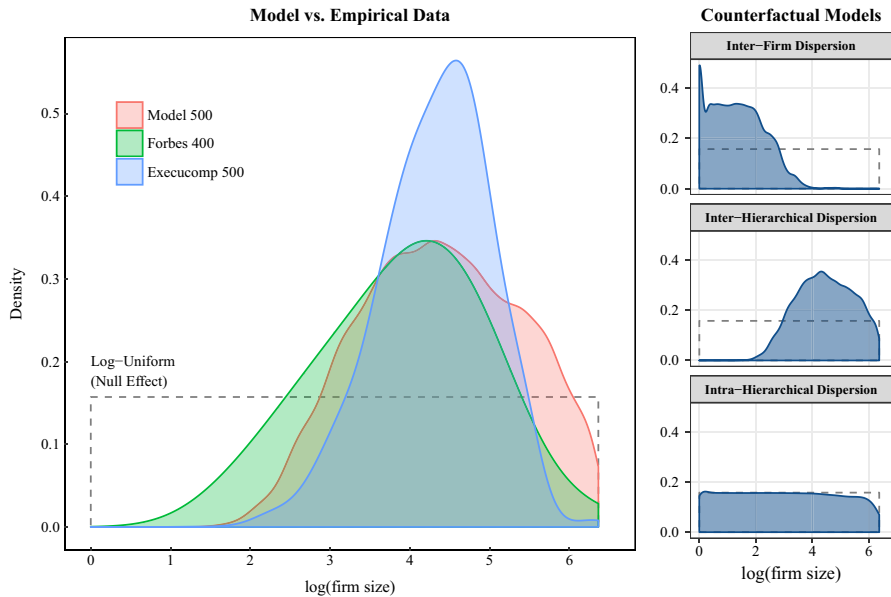


Fig. 5 Firm-size distributions associated with top incomes and wealth. This figure shows the size distribution of firms associated with top-earning individuals in the USA and in the hierarchy model (of the USA). The ‘Forbes 400’ represents the size distribution of firms associated with (owned by) the wealthiest 400 Americans in the year 2014. The ‘Execucomp 500’ represents the size distribution of firms associated with the 500 top-earning American executives (in each year from 1992 to 2015) in the Execucomp database. The ‘Model 500’ represents the size distribution of firms associated with the 500 top-earning individuals in the hierarchy model. Results for counterfactual models are shown on the right. Each counterfactual model isolates a single source of income dispersion. The top panel shows a model with inter-firm dispersion only, the middle shows a model with inter-hierarchical dispersion only, and the bottom shows a model with intra-hierarchical level dispersion only. In all plots, I also show the log-uniform distribution (dotted line), which is the null-effect prediction (i.e., no relation between firm membership and income). For sources and methods, see Online Appendix A

model results are not identical to either the Forbes 400 or the Execucomp 500 distributions. But, given the paucity of data on which the model is based (as well as the general uncertainty in the empirical analysis of top incomes), I count this result as a success. The model produces results that are roughly consistent with the US data.

Since the hierarchy model has three sources of income dispersion, we naturally want to know which of these sources is responsible for producing the results in Fig. 5. To answer this question, I use a counterfactual analysis. I create three different counterfactual models to supplement the original. Each counterfactual model isolates a single source of dispersion as it appears in the original model. The results of this counterfactual analysis are shown in the right-hand panels in Fig. 5. This analysis indicates that it is *exclusively* inter-hierarchical income dispersion that is responsible for associating top incomes with large institutions. The inter-hierarchical dispersion model produces results that are virtually identical to the original model. At the same time, inter-firm dispersion only and intra-hierarchical level dispersion only models produce *drastically* different results. (Note that with intra-hierarchical

dispersion only, we recover the null effect. Why? In this model, firms play no part in determining income).

To summarize, the hierarchy model correctly predicts that top-paid individuals should be associated with firms that are far larger than those of the general population. Moreover, the model indicates that this effect is purely a result of inter-hierarchical pay dispersion.

Isolating the distributional role of firm hierarchy

Having established that the hierarchy model gives credible results, I now use it to isolate the distributional effects of firm hierarchy. In particular, I am interested in determining whether or not it is hierarchy that shapes the income distribution tail. As in Fig. 5, I isolate the effects of hierarchy using a counterfactual analysis. I create three different counterfactual models of the USA, each containing only one source of income dispersion. By comparing these counterfactual models to the original model, we can determine how each dispersion source affects income distribution.

Figure 6 shows the results of this analysis. Here, I plot the income distribution (the probability density) of the original and counterfactual models. I use a log–log transformation to more clearly illustrate the distribution tail. This visualization allows us to see how each factor contributes to the original model’s distribution of income. To interpret this plot, look at the vertical distance between the original and counterfactual models. The closer a particular counterfactual model comes to the original model, the more important that dispersion factor is for shaping income distribution at the point in question.

The results of this analysis are unambiguous. A clear division exists between the *body* and *tail* of the distribution. The body of the distribution is almost completely determined by *inter-firm* dispersion, while the tail of the distribution is almost completely determined by *inter-hierarchical* dispersion. Intra-hierarchical dispersion amounts to negligible noise. The inset panel in Fig. 6 shows the fitted power-law exponent for the top 1% of incomes in the original and inter-hierarchical dispersion model. This confirms what is visually obvious in the main plot—the tail of the inter-hierarchical dispersion model is virtually identical to that of the original.

To summarize, the counterfactual analysis indicates that it is inter-hierarchical pay-scaling (alone) that is responsible for generating the model’s income distribution tail. This suggests that it is hierarchy that is responsible for generating the power-law tail, and that the effects of hierarchy become important in the top 1% of incomes.

How hierarchy generates the power-law tail

How does hierarchy create the (approximate) power-law distribution of top incomes? The basic mechanism was theorized by Lydall [13]. It relies on the following contrapuntal exponential tendencies of hierarchical organization:

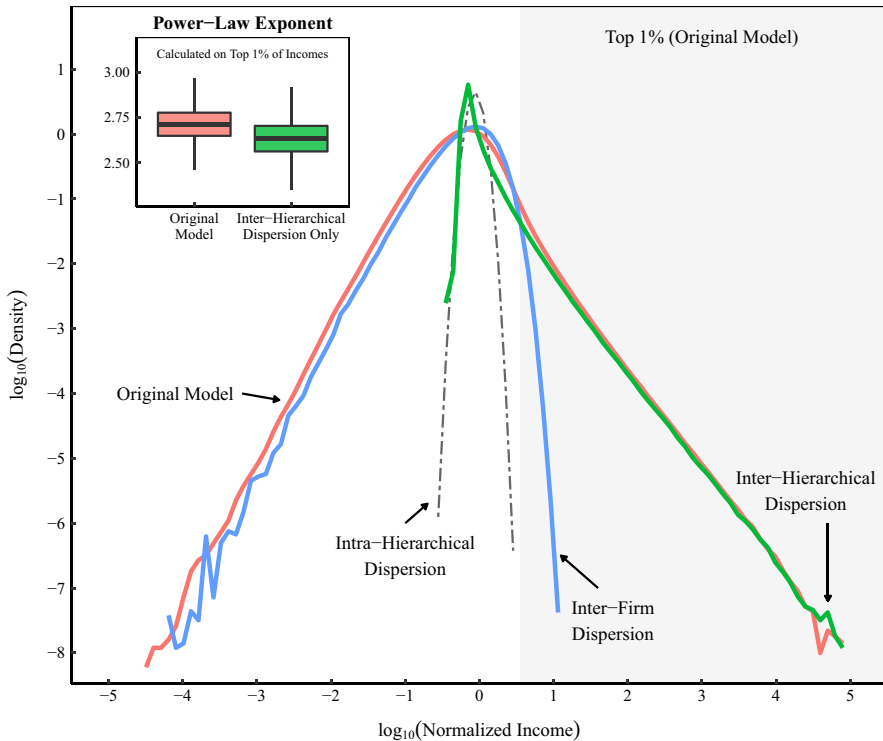


Fig. 6 Isolating the effects of hierarchy with counterfactual models. This figure compares the original hierarchy model of the USA to three different counterfactual models. Each counterfactual model contains only one of the three sources of income dispersion. The main plot shows the income probability density of each model, plotted using a log–log transformation (these results show the average distribution over many iterations). To interpret this plot, look at the vertical distance between each counterfactual model’s distribution and that of the original. The smaller the distance, the greater is the distributional role played by that dispersion factor at the point in question. The shaded region indicates the top 1% of incomes (in the original model). The inset panel shows power-law exponents fitted to the top 1% of incomes in the original and inter-hierarchical dispersion model

1. Hierarchical organization causes the share of employment to *decrease* exponentially with hierarchical rank.
2. Hierarchical pay structure causes income to *increase* exponentially with rank.

These two opposing tendencies interact to produce a power-law distribution of income (in the tail). This mechanism is a specific case of a more general method. A power law will be created any time we *exponentially* transform an *exponential* distribution [32].

The proof works as follows. Suppose we have some quantity y that is exponentially distributed (here a is a negative constant):

$$p(y) \sim e^{ay}. \quad (2)$$

In the case of hierarchical class structure, this would be the probability of finding someone with a hierarchical rank y . Suppose that we have another variable, x , that is *also* exponentially related to y :

$$x = e^{by}. \quad (3)$$

In the context of hierarchical organization, x would be income, which *increases* exponentially with rank. We want to know how income (x) is distributed. To find out, we use the change of variable formula to get f_x , the density function of x :

$$f_x = f_y(y(x)) \cdot |y'(x)|. \quad (4)$$

We let $f_y = e^{ay}$. Since $x = e^{by}$, we note that $y(x) = \frac{1}{b} \ln x$ and $y'(x) = 1/bx$. Substituting into the change of variable formula gives:

$$f_x = e^{\frac{a}{b} \ln x} \cdot \frac{1}{bx} = \frac{1}{b} x^{a/b-1}. \quad (5)$$

Thus, the variate x (income) has a power-law distribution with exponent $\alpha = a/b - 1$.

To reiterate, hierarchical organization creates a power-law distribution because of two contrapuntal, exponential tendencies: (1) employment tends to *decrease* exponentially with rank; and (2) income tends to *increase* exponentially with rank. Figure 7 illustrates this contrapuntal behavior in the hierarchy model. Figure 7a shows the aggregate hierarchical employment structure of the model. As expected, the hierarchical employment distribution has a bottom-heavy pyramid shape. The vast majority of people occupy low ranks and only a tiny elite have high rank. The inset panel highlights the *exponential* nature of this distribution. Figure 7b shows the model's aggregate hierarchical *pay* structure. As expected, hierarchical pay has an inverted pyramid shape. The average income at the top of the hierarchy dwarfs that at the bottom. Again, the inset plot highlights the exponential nature of this relation.

Note that neither employment nor pay has a *purely* exponential relation with rank. This is a design feature of the model, stemming from case-study evidence. In the case-study data, income tends to *increase* supra-exponentially (faster than an exponential) with rank. Conversely, employment tends to *decrease* supra-exponentially with rank (see Online Appendix B for details). In any case, when we combine these two supra-exponential tendencies, the result still seems to be (roughly) a power-law distribution of income in the model's tail.

While the above derivation highlights the basic power-law generation mechanism, the hierarchy model's inner workings involve some added complexity. First, the above derivation assumes that rank (y) is a continuous variable. In the model, rank is a *discrete* variable, which would result in a discontinuous distribution of pay (x) in Eq. 5. Lydall noted this in his original derivation, and posited that a process of 'blurring' would occur (due to stochastic differences in pay between firms) that would make the resulting distribution continuous [13]. In this regard, Lydall's intuition appears to be correct.

Figure 8 shows how the various discrete hierarchical ranks contribute to produce the continuous power-law tail. Each panel shows the distribution of income

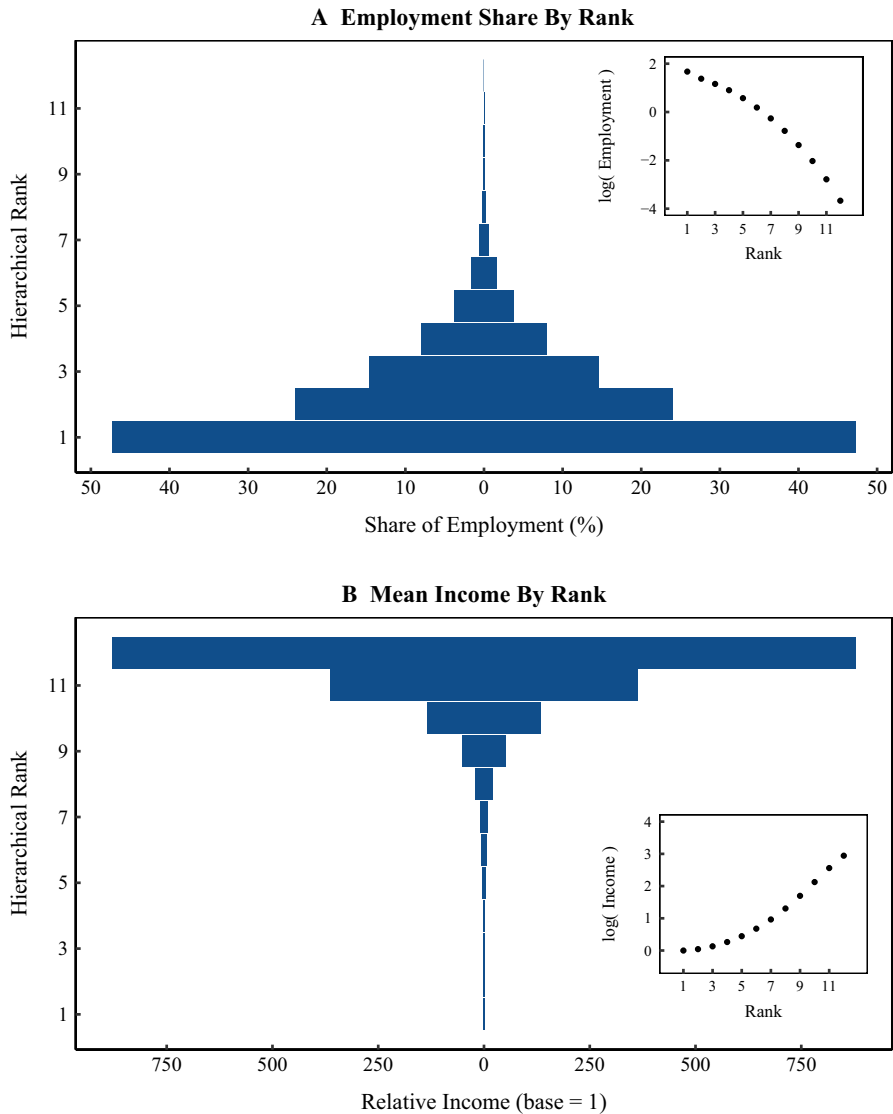


Fig. 7 The hierarchy model's contrapuntal exponential tendencies. This figure shows the two contrapuntal exponential tendencies associated with the hierarchy model's class structure. **a** The model's aggregate distribution of employment by hierarchical rank. The bottom-heavy shape results from firm hierarchical structure (in conjunction with the firm-size distribution). The inset graph shows the logarithm of employment share, plotted against rank. A pure exponential function would appear as a straight line. The curve in this relation indicates that employment declines with rank slightly faster than an exponential function. **b** The model's mean pay by hierarchical rank (normalized so that the base level = 1). The inset graph shows the logarithm of income plotted against rank. The curve in this relation indicates that income increases with rank slightly faster than an exponential function

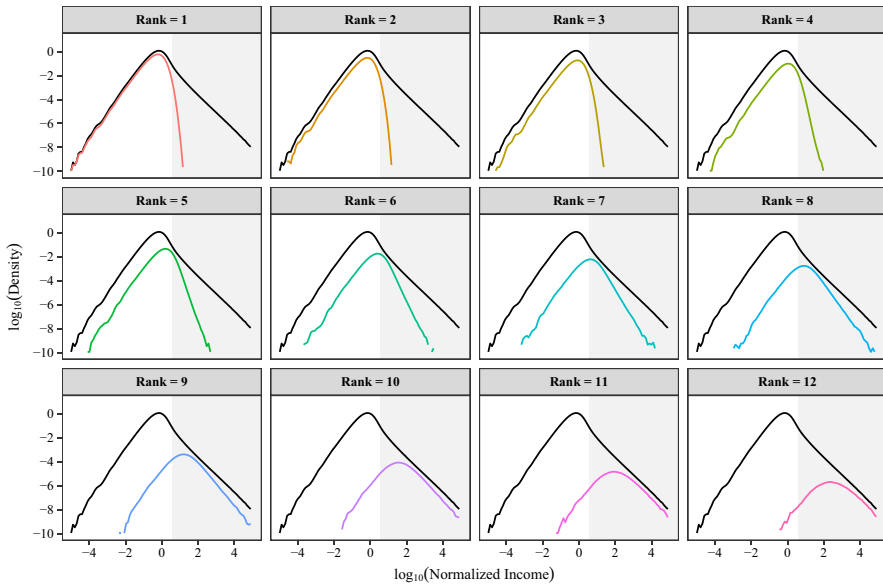


Fig. 8 The model's distribution of income by hierarchal rank. This figure shows the distribution of income for each hierarchical rank in the hierarchy model. In each panel, a rank-specific income distribution (color) is compared to the model's aggregate income distribution (black). The rank-specific distributions are normalized so that cumulative density of all ranks sums to one. The shaded region indicates the top 1% of incomes (in the aggregate model distribution). To interpret this plot, look at how closely each rank-specific distribution comes to the aggregate distribution. The closer the two are, the greater is the rank's contribution to income distribution at that point. The power-law right tail (evident as the straight line in the aggregate distribution) is jointly created by ranks five and up

of a *specific* hierarchical rank in relation to the model's aggregate income distribution. (The rank-specific distributions are normalized so that the cumulative density of all ranks sums to one.) In this plot, the exponential growth of income with rank appears as a horizontal shift in the income distribution of each rank. At the same time, each successive rank has exponentially fewer members, which appears as a downward shift in the income distribution. When the contributions of all ranks are summed, the result is an approximate power-law distribution of top incomes. As Lydall suspected, a complex blurring process occurs (between ranks) that smooths out what would otherwise be a discontinuous distribution.

Discussion

Whenever two or more theories describe the same phenomenon, we need to determine if they are consistent with one another, or if they are mutually inconsistent. Thus, we should ask—is the hierarchy model's explanation of the power-law distribution of top incomes at odds with the stochastic growth models described

in “Power-law generation mechanisms”? Or are the two approaches mutually consistent?

The primary difference between the two approaches is that the hierarchy model is *static*, while the stochastic models are *dynamic*. The hierarchy model begins with the observation that firms have a hierarchical structure, and that this (static) structure could explain the distribution of income at a point in time. The hierarchy model says nothing about the dynamics of individual income, but instead focuses on institutional structure. In contrast, the stochastic approach begins with the observation that individual incomes *change* over time. Since income distribution represents a snapshot of these changing incomes, it must be possible to explain income distribution in terms of the dynamics of individual income. The static and dynamic approaches explain the power-law distribution of top incomes from very different angles. Therefore, I see no fundamental clash between the hierarchy model and exogenous stochastic growth models in the tradition of Champervowne [34]. More research is needed to determine how the two approaches are related.

That being said, the stochastic growth and firm hierarchy models each have very different implications for how we should study (and potentially alleviate) inequality. Stochastic growth models put the focus on isolated individuals. This makes it difficult to connect inequality to the wider political and socioeconomic setting (the search for such a connection is a major goal of many economists and sociologists [59–71]). In contrast, the hierarchy model suggests that the income distribution power-law tail is an outcome of the internal compensation policies of firms. This puts the focus squarely on firms and how they remunerate their employees as a function of hierarchical rank. This perspective opens the door to future research that connects the internal pay policies of firms to the wider distribution of income (and potentially to government policy).

Conclusions

In 1959, when Lydall [13] first proposed that firm hierarchy could create a power-law distribution of income, his hypothesis was largely speculative. At the time, little was known about the internal pay structure of firms. Nearly 60 years later, data on firm hierarchical structure is still scarce, but enough evidence exists that we can begin to investigate the distributional effects of firm hierarchy. This paper has presented a first attempt at doing so.

The key finding is that the empirically informed hierarchy model is capable of reproducing the power-law scaling of top US incomes, while at the same time accurately connecting top-earning individuals to large firms. Importantly, the model indicates that it is hierarchical pay-scaling alone that is responsible for these results. Of course, the hierarchy model’s results are contingent on the input data, which is limited. While I have made every effort to incorporate uncertainty in the underlying case-study data, results may change when new data comes along. Further research is needed to verify these results and see if they can be replicated in other countries.

Uncertainty aside, the hierarchy model suggests that the ubiquitous power-law scaling of top incomes may be a result of the ubiquitous use of hierarchical

organization in human societies. This implies that when we study the tail of the distribution of income, we may be studying the effects of social hierarchy.

Acknowledgements I would like to thank the Social Sciences and Humanities Resource Council of Canada (Grant no. 767-2015-1015) for its support. I thank also Jonathan Nitzan, who has offered feedback on aspects of this paper.

References

1. Pareto, V. (1897). *Cours d'economie politique* (Vol. 1). Geneva: Librairie Droz.
2. Atkinson, A. B. (2017). Pareto and the upper tail of the income distribution in the UK: 1799 to the present. *Economica*, 84(334), 129.
3. Aoyama, H., Souma, W., Nagahara, Y., Okazaki, M. P., Takayasu, H., & Takayasu, M. (2000). Pareto's law for income of individuals and debt of bankrupt companies. *Fractals*, 8(03), 293.
4. Coelho, R., Richmond, P., Barry, J., & Hutzler, S. (2008). Double power laws in income and wealth distributions. *Physica A: Statistical Mechanics and its Applications*, 387(15), 3847.
5. Clementi, F., & Gallegati, M. (2005). Pareto's law of income distribution: Evidence for Germany, the United Kingdom, and the United States. In A. Chatterjee, S. Yarlagadda, & B. K. Chakrabarti (Eds.), *Econophysics of wealth distributions* (pp. 3–14). Berlin: Springer.
6. Clementi, F., & Gallegati, M. (2005). Power law tails in the Italian personal income distribution. *Physica A: Statistical Mechanics and Its Applications*, 350(2–4), 427.
7. Di Guilmi, C., Gaffeo, E., & Gallegati, M. (2003). Power law scaling in world income distribution. *Economics Bulletin*, 15, 1–7.
8. Drgulescu, A., & Yakovenko, V. M. (2001). Exponential and power-law probability distributions of wealth and income in the United Kingdom and the United States. *Physica A: Statistical Mechanics and its Applications*, 299(1–2), 213.
9. Nirei, M. (2009). Pareto distributions in economic growth models. IIR working paper WP#09-05
10. Toda, A. A. (2012). The double power law in income distribution: Explanations and evidence. *Journal of Economic Behavior & Organization*, 84(1), 364.
11. Silva, A. C., & Yakovenko, V. M. (2004). Temporal evolution of the thermal and superthermal income classes in the USA during 19832001. *EPL (Europhysics Letters)*, 69(2), 304.
12. Souma, W. (2001). Universal structure of the personal income distribution. *Fractals*, 9(04), 463.
13. Lydall, H. F. (1959). The distribution of employment incomes. *Econometrica: Journal of the Econometric Society*, 27(1), 110.
14. Audas, R., Barmby, T., & Treble, J. (2004). Luck, effort, and reward in an organizational hierarchy. *Journal of Labor Economics*, 22(2), 379.
15. Baker, G., Gibbs, M., & Holmstrom, B. (1993). Hierarchies and compensation: A case study. *European Economic Review*, 37(2–3), 366.
16. Dohmen, T. J., Krieche, B., & Pfann, G. A. (2004). Monkey bars and ladders: The importance of lateral and vertical job mobility in internal labor market careers. *Journal of Population Economics*, 17(2), 193.
17. Grund, C. (2005). The wage policy of firms: Comparative evidence for the US and Germany from personnel data. *The International Journal of Human Resource Management*, 16(1), 104.
18. Lima, F. (2000). Internal labor markets: A case study. FEUNL working paper, vol. 378.
19. Morais, F., & Kakabadse, N. K. (2014). The corporate gini index (CGI) determinants and advantages: Lessons from a multinational retail company case study. *International Journal of Disclosure and Governance*, 11(4), 380.
20. Treble, J., Van Gameren, E., Bridges, S., & Barmby, T. (2001). The internal economics of the firm: Further evidence from personnel data. *Labour Economics*, 8(5), 531.
21. Ariga, K., Brunello, G., Ohkusa, Y., & Nishiyama, Y. (1992). Corporate hierarchy, promotion, and firm growth: Japanese internal labor market in transition. *Journal of the Japanese and International Economies*, 6(4), 440.
22. Bell, B., & Van Reenen, J. (2012). Firm performance and wages: Evidence from across the corporate hierarchy. CEP discussion paper, vol. 1088.

23. Eriksson, T. (1999). Executive compensation and tournament theory: Empirical tests on Danish data. *Journal of Labor Economics*, 17(2), 262.
24. Heyman, F. (2005). Pay inequality and firm performance: Evidence from matched employer–employee data. *Applied Economics*, 37(11), 1313.
25. Leonard, J. S. (1990). Executive pay and firm performance. *Industrial and Labor Relations Review*, 43(3), 13.
26. Main, B. G., O'Reilly, C. A. I. I., & Wade, J. (1993). Top executive pay: Tournament or team-work? *Journal of Labor Economics*, 11(4), 606.
27. Mueller, H. M., Ouimet, P. P., & Simintzi, E. (2016). Within-firm pay inequality. SSRN working paper.
28. Rajan, R. G., & Wulf, J. (2006). The flattening firm: Evidence from panel data on the changing nature of corporate hierarchies. *The Review of Economics and Statistics*, 88(4), 759.
29. Tao, H. L., & Chen, I. T. (2009). The level of technology employed and the internal hierarchical wage structure. *Applied Economics Letters*, 16(7), 739.
30. Kumamoto, S. I., & Kamihigashi, T. (2018). Power laws in stochastic processes for social phenomena: An introductory review. *Frontiers in Physics*, 6, 20.
31. Mitzenmacher, M. (2004). A brief history of generative models for power law and lognormal distributions. *Internet mathematics*, 1(2), 226.
32. Newman, M. E. (2005). Power laws, Pareto distributions and Zipf's law. *Contemporary Physics*, 46(5), 323.
33. Gibrat, R. (1931). *Les inegalites economiques*. Paris: Recueil Sirey.
34. Champernowne, D. G. (1953). A model of income distribution. *The Economic Journal*, 63(250), 318.
35. Gabaix, X., Lasry, J. M., Lions, P. L., & Moll, B. (2016). The dynamics of inequality. *Econometrica*, 84(6), 2071.
36. Nirei, M., & Souma, W. (2007). A two factor model of income distribution dynamics. *Review of Income and Wealth*, 53(3), 440.
37. Rutherford, R. (1955). Income distributions: A new model. *Econometrica: Journal of the Econometric Society*, 23(3), 277.
38. Wold, H. O., & Whittle, P. (1957). A model explaining the Pareto distribution of wealth. *Econometrica, Journal of the Econometric Society*, 25(4), 591.
39. Yule, G. U. I. I. (1925). A mathematical theory of evolution, based on the conclusions of Dr. JC Willis, FR S. *Philosophical Transactions of the Royal Society B*, 213(402–410), 21.
40. Simon, H. A. (1955). On a class of skew distribution functions. *Biometrika*, 42(3/4), 425.
41. Price, D. D. S. (1976). A general theory of bibliometric and other cumulative advantage processes. *Journal of the Association for Information Science and Technology*, 27(5), 292.
42. Barabasi, A., & Albert, R. (1999). Emergence of scaling in random networks. *Science*, 286(5439), 509.
43. Angle, J. (1986). The surplus theory of social stratification and the size distribution of personal wealth. *Social Forces*, 65(2), 293.
44. Angle, J. (2006). The inequality process as a wealth maximizing process. *Physica A: Statistical Mechanics and Its Applications*, 367(15), 388.
45. Chatterjee, A., Chakrabarti, B. K., & Chakraborti, A. (2007). *Econophysics and sociophysics: Trends and perspectives*. Milan: Wiley.
46. Chatterjee, A., & Chakrabarti, B. K. (2007). Kinetic exchange models for income and wealth distributions. *The European Physical Journal B*, 60(2), 135.
47. Hegyi, G., Neda, Z., & Santos, M. A. (2007). Wealth distribution and Pareto's law in the Hungarian medieval society. *Physica A: Statistical Mechanics and its Applications*, 380(1), 271.
48. Kitov, I. O. (2009). Mechanical model of personal income distribution (pp. 1–220). arXiv preprint [arXiv:0903.0203](https://arxiv.org/abs/0903.0203)
49. Pianegonda, S., Iglesias, J. R., Abramson, G., & Vega, J. L. (2003). Wealth redistribution with conservative exchanges. *Physica A: Statistical Mechanics and its Applications*, 322, 667.
50. Scheffer, M., Bavel, B. V., Leemput, I. A. V. D., Nes, E. H. V. (2017) Inequality in nature and society. *Proceedings of the National Academy of Sciences*, 114, 201706412. <https://doi.org/10.1073/pnas.1706412114>
51. Yakovenko, V. M., & Rosser, J. B, Jr. (2009). Colloquium: Statistical mechanics of money, wealth, and income. *Reviews of Modern Physics*, 81(4), 1703.
52. Simon, H. A. (1957). The compensation of executives. *Sociometry*, 20(1), 32.

53. Roberts, D. R. (1956). A general theory of executive compensation based on statistically tested propositions. *The Quarterly Journal of Economics*, 70(2), 270.
54. Axtell, R. L. (2001). Zipf distribution of US firm sizes. *Science*, 293, 1818.
55. Fix, B. (2017). Energy and institution size. *PLoS ONE*, 12(2), e0171823. <https://doi.org/10.1371/journal.pone.0171823>.
56. Gaffeo, E., Gallegati, M., & Palestini, A. (2003). On the size distribution of firms: Additional evidence from the G7 countries. *Physica A: Statistical Mechanics and its Applications*, 324(12), 117. [https://doi.org/10.1016/S0378-4371\(02\)01890-3](https://doi.org/10.1016/S0378-4371(02)01890-3).
57. Clauset, A., Shalizi, C. R., & Newman, M. E. (2009). Power-law distributions in empirical data. *SIAM Review*, 51(4), 661.
58. Piketty, T. (2014). *Capital in the twenty-first century*. Cambridge: Harvard University Press.
59. Brown, C. (1988). In Mangum, G. & P. Philips (Eds), *Three worlds of labor economics* (Vol. 51 pp. 515–530).
60. Brown, C. (2005). Is there an institutional theory of distribution? *Journal of Economic Issues*, 39(4), 915.
61. Dahrendorf, R. (1959). *Class and class conflict in industrial society*. Stanford: Stanford University Press.
62. Huber, E., Huo, J., & Stephens, J. D. (2017). Power, policy, and top income shares. *Socio-Economic Review*. <https://doi.org/10.1093/ser/mwx027>.
63. Kalecki, M. (1971). *Selected essays on the dynamics of the capitalist economy 1933–1970*. Cambridge: Cambridge University Press.
64. Lenski, G. E. (1966). *Power and privilege: A theory of social stratification*. Chapel Hill: UNC Press Books.
65. Marx, K. (1867). *Capital* (Vol. I). Penguin/New Left Review: Harmondsworth.
66. Mills, C. W. (1956). *The power elite*. Oxford: Oxford University Press.
67. Nitzan, J., & Bichler, S. (2009). *Capital as power: A study of order and creorder*. New York: Routledge.
68. Peach, J. T. (1987). Distribution and economic progress. *Journal of Economic Issues*, 21(4), 1495.
69. Sidanius, J., & Pratto, F. (2001). *Social dominance: An intergroup theory of social hierarchy and oppression*. Cambridge: Cambridge University Press.
70. Weber, M. (1978). *Economy and society: An outline of interpretive sociology*. Berkeley: University of California Press.
71. Wright, E. O. (1979). *Class structure and income determination* (Vol. 2). New York: Academic Press.

Appendices for ‘Hierarchy and the Power-Law Income Distribution Tail’

Blair Fix

Supplementary materials for this paper are available at the Open Science Framework repository:

<https://osf.io/mb3ah/>

The supplementary materials include:

1. Raw source data;
2. R code for all analysis;
3. Hierarchy model code.

Contents

A Sources and Methods	2
B Case-Study Firms	7
C Compustat Data	13
D Hierarchy Model Equations	16
E Restricting Parameters	22
F The Adjusted Hierarchy Model	35
G A Null-Effect Model for US Top Incomes and Firm Size	37
H The Effect of Hierarchy on Inequality	39

A Sources and Methods

Sources are listed by the figure in which they appear.

Sources for Figure 4 (Modeled Income Distribution vs. US Data)

Complementary Cumulative Distribution

The US complementary cumulative distribution is calculated from data in the IRS *Individual Complete Report* (Publication 1304), [Table 1.1](#), from 1996 to 2015.

Cumulative Distribution

The US cumulative distribution is calculated from data in the IRS *Individual Complete Report* (Publication 1304), [Table 1.1](#), from 1996 to 2015.

Gini Index

I use two sources for the US Gini index. The first source is the US Current Population Survey, Table PINC-08 (available from the [US Census](#)) over the years 1994 to 2015. The second source is the IRS *Individual Complete Report* (Publication 1304), [Table 1.1](#), from 1996 to 2015. I estimate the Gini index by constructing a Lorenz curve from the reported cumulative frequency data. R code implementing this method is available in the Supplementary Material.

The Census and IRS data are *not* mutually consistent. IRS data is based on tax units, not individuals. The advantage of the IRS data is that it is an administrative record. Current Population Survey (CPS) data, on the other hand, is obtained by *interview*. The advantage of the CPS data is that it explicitly counts individuals. The disadvantage is that “there is a tendency in household surveys for respondents to *under report* their income” [1].

Lorenz Curve

The US Lorenz curve is calculated from data in the IRS *Individual Complete Report* (Publication 1304), [Table 1.1](#), from 1996 to 2015.

Power Law Exponents

I estimate the power law exponent of the income distribution tail using the maximum likelihood method. US empirical data comes from the IRS *Individual*

Table 1: Power Law Cutoff Boundaries in US Data

Year	Percentile	α
1996	0.987	2.92
1997	0.985	2.89
1998	0.996	2.58
1999	0.996	2.58
2000	0.995	2.54
2001	0.996	2.63
2002	0.996	2.67
2003	0.996	2.65
2004	0.995	2.59
2005	0.994	2.54
2006	0.993	2.54
2007	0.993	2.54
2008	0.994	2.66
2009	0.995	2.78
2010	0.994	2.73
2011	0.994	2.74
2012	0.992	2.64
2013	0.993	2.74
2014	0.992	2.70
2015	0.991	2.72

Complete Report (Publication 1304), [Table 1.1](#). Since this data is reported in *binned* form, I use the binned log-likelihood equation developed by Virkar and Clauset [2]:

$$\mathcal{L} = n(\alpha - 1) \cdot \ln b_{\min} + \sum_{i=\min}^k h_i \ln \left[b_i^{(1-\alpha)} - b_{i+1}^{(1-\alpha)} \right] \quad (1)$$

Here α is the power law exponent, b_i and b_{i+1} are consecutive bin boundaries, h_i and h_{i+1} are consecutive bin counts, k is the number of bins, and n is the sum of bin counts above $b_{\min}$ (the cutoff point for the power law). The best-fit exponent α is the value that maximizes the log-likelihood function ($\mathcal{L}$). Since there is no closed-form solution to this maximization problem, I solve for α numerically. To determine the power law exponent for the top 1% of incomes in each year, I set the power law cutoff boundary ($b_{\min}$) to the empirical bin that is closest to the 99th percentile. Results are shown in [Table 1](#).

To find the power law exponent in modeled data, I use the following maximum likelihood estimator:

$$\hat{\alpha} = 1 + n \left[\sum_i^n \ln \frac{x_i}{x_{\min}} \right]^{-1} \quad (2)$$

Here $\hat{\alpha}$ is the best-fit power law exponent, x_i is the i th data point, $x_{\min}$ is the lower bound of the power law, and n is the number of data points above $x_{\min}$. To ensure compatibility with empirical power law estimates, I estimate the model's power law exponent using the *empirical* cutoff values. For each model run, I set $x_{\min}$ by randomly selecting a percentile value from Table 1.

All data and code are available in the Supplementary Material.

Probability Density Function

I estimate the normalized probability density function for US income using data from Current Population Survey Table PINC-08 (available from the [US Census](#)) over the years 1994 to 2015. This table reports binned data.

To estimate the normalized probability density function in each year, I first create a simulated income distribution (**I**) using bin midpoints. Each midpoint income M_i is repeated F_i times, where F_i is the frequency count for the i th bin. I then normalize **I** by dividing all elements by the mean income $\bar{I}$.

$$\mathbf{I} = \frac{(M_1 \cdot \dots \cdot \times F_1, M_2 \cdot \dots \cdot \times F_2, \dots, M_i \cdot \dots \cdot \times F_i)}{\bar{I}} \quad (3)$$

Lastly, I fit the simulated income distribution (**I**) with a numerical density function. R code implementing this method is available in the Supplementary Material.

Top 1% Income Share

Sources for top 1% income share data are shown in Table 2.

Table 2: US Top 1% Income Share Sources

Series	Info	Source
sfainc992j	Pre-tax factor income equal-split adults Share Adults share of total (ratio)	[3]
sfainc996i	Pre-tax factor income individuals Share 20 to 64 share of total (ratio)	[3]
sfainc999i	Pre-tax factor income individuals Share All Ages share of total (ratio)	[3]
sfainc999t	Pre-tax factor income tax unit Share All Ages share of total (ratio)	[3]
sfiinc992j	Fiscal income equal-split adults Share Adults share of total (ratio)	[3]
sfiinc992t	Fiscal income tax unit Share Adults share of total (ratio)	[3]
sfiinc996i	Fiscal income individuals Share 20 to 64 share of total (ratio)	[3]
sfiinc999i	Fiscal income individuals Share All Ages share of total (ratio)	[3]
sfiinc999t	Fiscal income tax unit Share All Ages share of total (ratio)	[3]
sptinc992j	Pre-tax national income equal-split adults Share Adults share of total (ratio)	[3]
sptinc996i	Pre-tax national income individuals Share 20 to 64 share of total (ratio)	[3]
sptinc999i	Pre-tax national income individuals Share All Ages share of total (ratio)	[3]
sptinc999t	Pre-tax national income tax unit Share All Ages share of total (ratio)	[3]
sfiinc_z_US	World Top Incomes Legacy Series	[4]
lakner	Calculated from micro data	[5]
piketty_book_no_kgains	Legacy data from Capital in the 21st Century	[6]
piketty_book_with_kgains	Legacy data from Capital in the 21st Century	[6]

Sources for Figure 5 (Firm Size Distributions Associated With Top Incomes and Wealth)

Forbes 400 data is from the year 2014. Firm size data was collected by the author. For public companies, firm size data comes from Compustat. For private companies, data comes from firm websites and annual reports. The Execucomp 500 consists of the 500 top paid US executives in the Execucomp database in each year from 1992 to 2015.

B Case-Study Firms

In this section I review the case-study evidence that informs the hierarchy model. Table 3 summarizes the source data, while Figure 1 shows the hierarchical employment and pay structure of these firms. The firms remain anonymous, and are named after the authors of the case-study papers. Although the exact shapes vary, all the firms in this sample have a roughly pyramidal employment structure and inverse pyramid pay structure.

Figure 2 dissects these trends to allow further analysis. Figure 2A shows how the span of control (the employment ratio between adjacent ranks) changes as a function of hierarchical level. In these firms, the span of control is not constant, but instead tends to *increase* with hierarchical level. Similarly, Figure 2B shows the ratio of mean pay between adjacent levels. Like the span of control, the pay ratio tends to increase with hierarchical level. Lastly, Figure 2C shows income dispersion within hierarchical ranks of each firm (measured with the Gini index). Note that income dispersion within levels is quite low and there is no evidence of a trend.

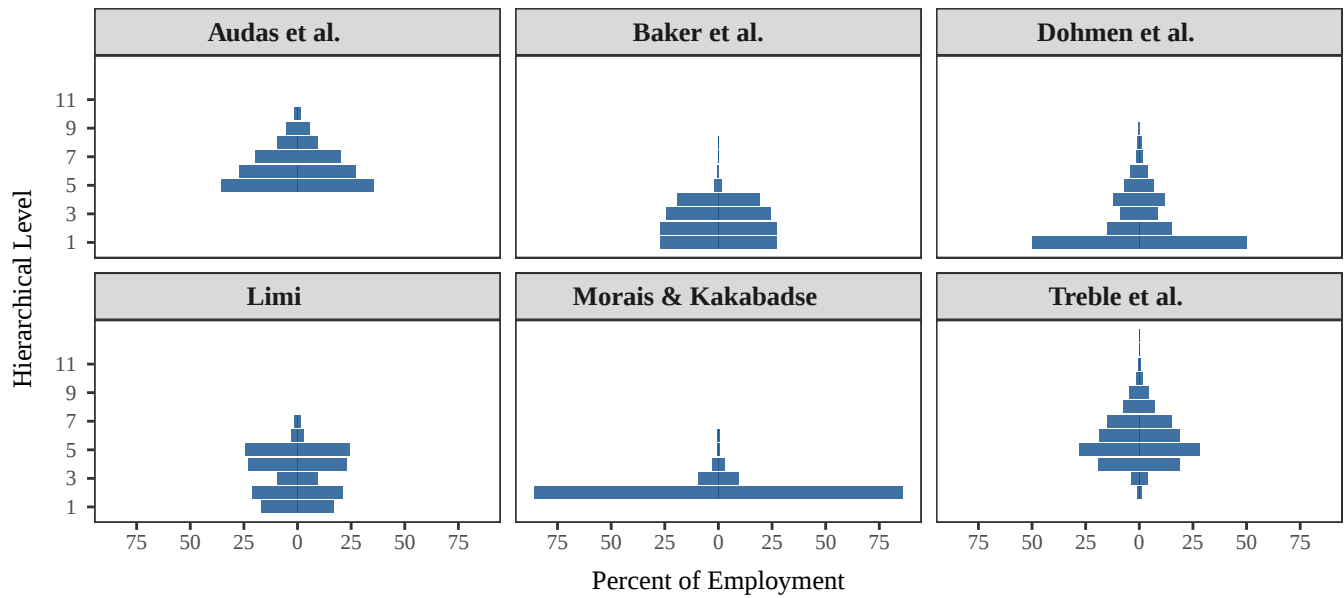
In addition to case-study data of single firms, several studies have reported the aggregate hierarchical structure of a sample of firms (see Table 4 and Figure 4). The data from these firms reveals the same general trends as the case studies. However, the aggregate data is less useful because these studies capture only the top few hierarchical ranks within firms.

The case-study data plays a central role in the hierarchical model developed in this paper. From the case-study evidence, I propose the following ‘stylized’ facts about firm employment and pay structure:

1. The span of control tends to *increase* with hierarchical level.
2. The inter-level pay ratio tends to *increase* with hierarchical level.
3. Intra-level income inequality is approximately *constant* across all hierarchical levels.

The case-study evidence informs the basic structure of the model, and also some of its key parameters. The ‘shape’ of modeled firm hierarchies is determined from the fitted span-of-control trend shown in Figure 2A. Figure 3 shows the idealized employment hierarchy that is implied by case-study data. Error bars indicate uncertainty, calculated using the bootstrap resampling method. Parameters for intra-level income dispersion are determined from the mean of data in Figure 2C. For a detailed discussion of the model algorithm and parameter-fitting procedure, see Sections D and E.

A. Firm Hierarchical Employment Structure



B. Firm Hierarchical Pay Structure

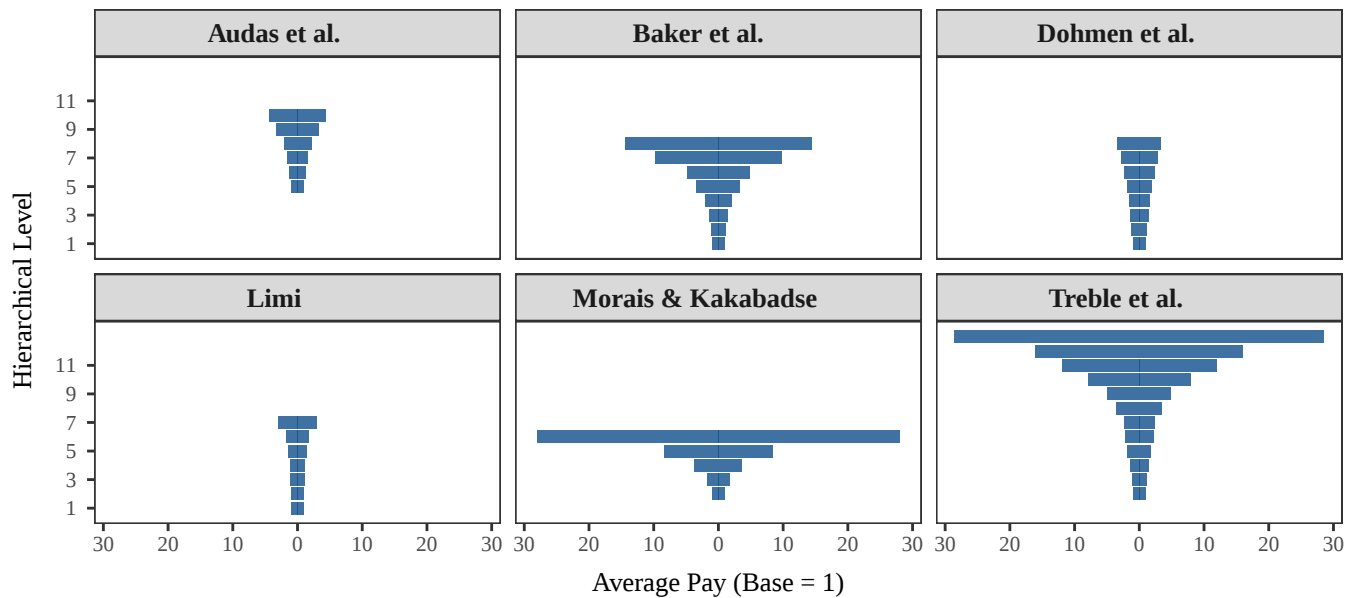


Figure 1: The Hierarchical ‘Shape’ of Six Different Case-Study Firms

This figure shows the hierarchical employment and pay structure of six different case-study firms. Panel A shows the hierarchical structure of employment, while Panel B shows the hierarchical pay structure.

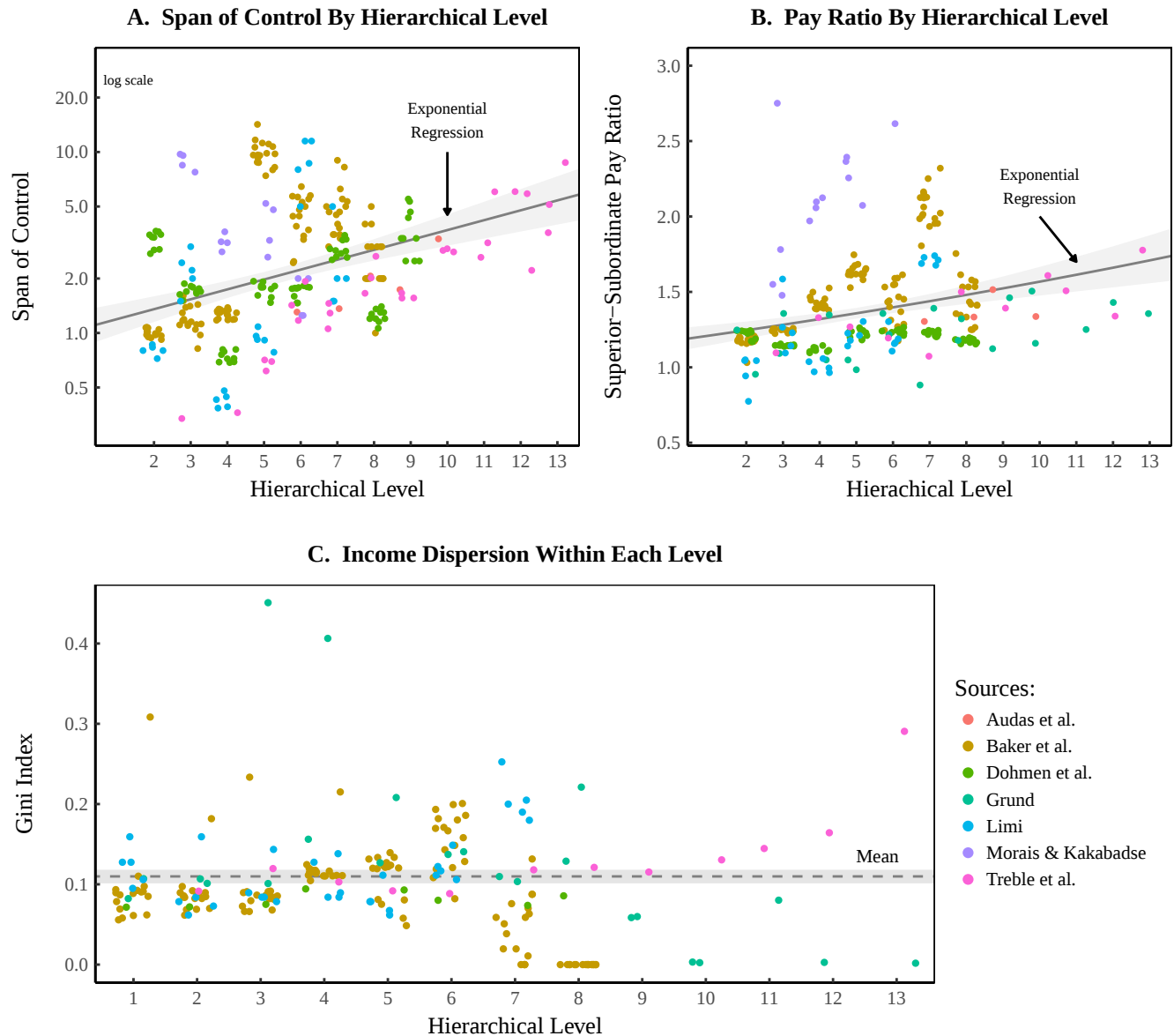


Figure 2: Analyzing the Hierarchical Structure of Case-Study Firms

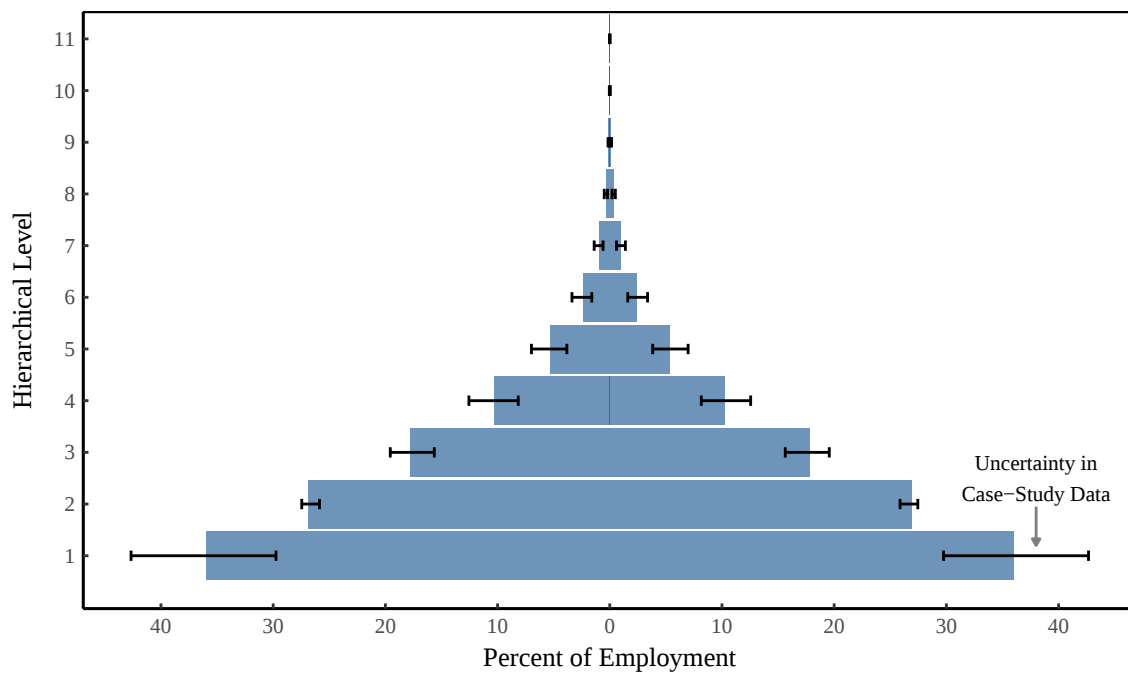
This figure shows data from 7 case-study firms. Panel A shows how the span of control (the subordinate-to-superior employment ratio between adjacent levels) varies with hierarchical level. Note the log scale on the y-axis. Panel B shows how the superior-to-subordinate pay ratio varies with hierarchical level. In Panels A and B, the x-axis corresponds to the *upper* hierarchical level in each corresponding ratio. Panel C shows the Gini index of income inequality within each hierarchical level. Different case-study firms are indicated by color, with names indicating the study author. Note that horizontal ‘jitter’ has been introduced in all three plots in order to better visualize the data (hierarchical level is a discrete variable). The lines in Panels A and B indicate exponential regressions, while the line in Panel C shows the average Gini index. Grey regions correspond to the 95% confidence intervals.

Table 3: Summary of Firm Case Studies

Source		Years	Country	Firm Levels	Span of Control	Level Income	Level Income Dispersion
Audas	[7]	1992	Britain	All	✓	✓	
Baker	[8]	1969-1985	United States	Management	✓	✓	✓
Dohmen	[9]	1987-1996	Netherlands	All	✓	✓	✓
Grund	[10]	1995 & 1998	US and Germany	All		✓	✓
Lima	[11]	1991-1995	Portugal	All	✓	✓	✓
Morais*	[12]	2007-2010	Undisclosed	All	✓	✓	
Treble	[13]	1989-1994	Britain	All	✓	✓	✓

Notes: This table shows metadata for the firm case studies displayed in Fig. 2. The ‘Firm Levels’ column refers to the portion of the firm that is included in the study. ‘Management’ indicates that only management levels were studied.

*For the analysis conducted in this paper I discard (as an outlier) the bottom hierarchical level in Morais and Kakabadse’s data.

**Figure 3: Idealized Firm Employment Hierarchy Implied by Case Studies**

This figure shows the idealized firm hierarchy that is implied by fitting trends to case-study data (Fig. 2A). Error bars show the uncertainty in the hierarchical shape, calculated using a bootstrap resample of case-study data.

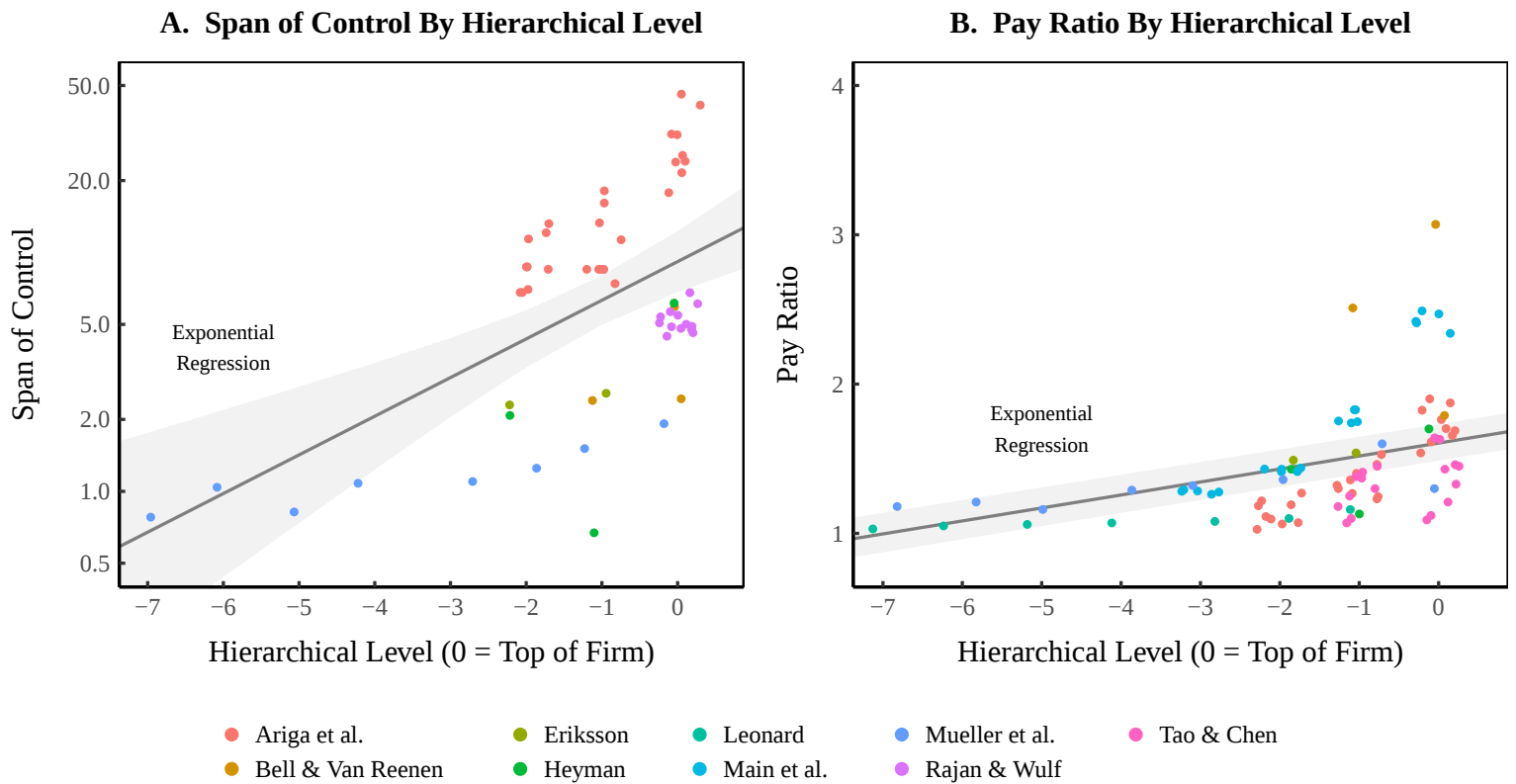


Figure 4: Aggregate Studies of Firm Hierarchical Structure

This figure shows data from 9 different aggregate firm studies. Most of these studies only survey the top several hierarchical levels in each firm. Because of this, I order hierarchical levels from the top down, where the CEO is level 0, the level below is -1, etc. Panel A shows how the span of control (the employment ratio between adjacent levels) relates to hierarchical level. Panel B shows how the pay ratio between adjacent levels varies with hierarchical level. In both plots, horizontal 'jitter' has been introduced in order to better visualize the data (hierarchical level is a discrete variable). Grey regions correspond to the 95% confidence interval for regressions.

Table 4: Summary of Firm Aggregate Studies

Source		Years	Number of Firms	Country	Firm Levels	Span of Control	Level Income
Ariga	[14]	1981-1989	unknown	Japan	All	✓	✓
Bell	[15]	2001-2010	552	United Kingdom	Top 3	✓	✓
Eriksson	[16]	1992-1995	210	Denmark	Management	✓	✓
Heyman	[17]	1991,1995	560	Sweden	Management	✓	✓
Leonard	[18]	1981-1985	439	United States	Top 9		✓
Main	[19]	1980–1984	200	United States	Top 4		✓
Mueller	[20]	2004-2013	880	United Kingdom	All	✓	✓
Rajan	[21]	1986-1998	261	United States	Top 2	✓	
Tao	[22]	1986-1998	8101	Taiwan	Top 2		✓

Notes: This table shows metadata for the aggregate studies displayed in Fig. 4. The ‘Firm Levels’ column refers to the portion of the firm that is included in the study. ‘Top 2’, ‘Top 3’, etc. indicates that only the top n levels were included in the study (where the top level is the CEO).

C Compustat Data

This paper makes extensive use of the *Compustat* and *Execucomp* databases. Compustat contains data for most publicly traded US companies, while Execucomp contains data for executive compensation. Three key statistics used throughout this paper are calculated from this data: *firm mean income*, the *CEO-to-average-employee pay ratio*, and the *capitalist income fraction of executives*. I discuss the data and methods used for these calculations in the following sections.

C.1 Firm Mean Income

Firm mean income is calculated by dividing total staff expenses (Compustat Series XLR) by total employment (Compustat Series EMP):

$$\text{Firm Mean Income} = \frac{\text{Total Staff Expenses}}{\text{Total Employment}} \quad (4)$$

C.2 CEO Pay Ratio

Throughout this paper, I use the term ‘CEO’ to refer to the executive at the top of the corporate hierarchy. I identify CEOs using the titles contained in the Execucomp series TITLEANN. Because titles vary greatly by company, identifying the top executive is not always a simple task. While a manual search would be most accurate, this is unrealistic given that the Execucomp database contains over 275 000 entries. Instead, I use the following three-step algorithm to identify the ‘CEO’:

1. Find all executives whose title contains one or more of the words in the ‘CEO Titles’ list (Table 5).
2. Of these executives, take the subset whose title does *not* contain any of the words in the ‘Subordinate Titles’ list (Table 5).
3. If this search returns more than one executive per firm per year, chose the executive with the highest pay.

After identifying the CEO (and matching CEO pay data with firm data contained in the Compustat database), I calculate the CEO pay ratio using the following equation:

$$\text{CEO Pay Ratio} = \frac{\text{CEO Pay}}{\text{Firm Mean Income}} \quad (5)$$

Table 5: Titles Used to Identify the ‘CEO’

CEO Titles:	Subordinate Titles
president	vp
chairman	v-p
CEO	cfo
Chief Executive Officer	vice
chmn	chief finance officer
	president of
	coo
	division
	div
	president-
	group president
	chairmain-
	co-president
	deputy chairman
	pres.-
	Chief Financial Officer

Notes: This table shows the Execucomp titles used to identify the CEO of each company. CEOs are deemed to be those whose title contains words in the left column, but not those in the right column. Titles such as ‘president-’ and ‘president of’ are included in the subordinate list because they typically refer to a president of a division with the company: i.e. ‘president of western division’ or ‘president-western hemisphere’.

CEO pay ratio and firm mean income data are collectively available for roughly 6000 firm-year observations over the period 1992-2016. I use this data to ‘tune’ my hierarchical model of the firm (see Section E) . Figure 5 shows selected summary statistics of this dataset.

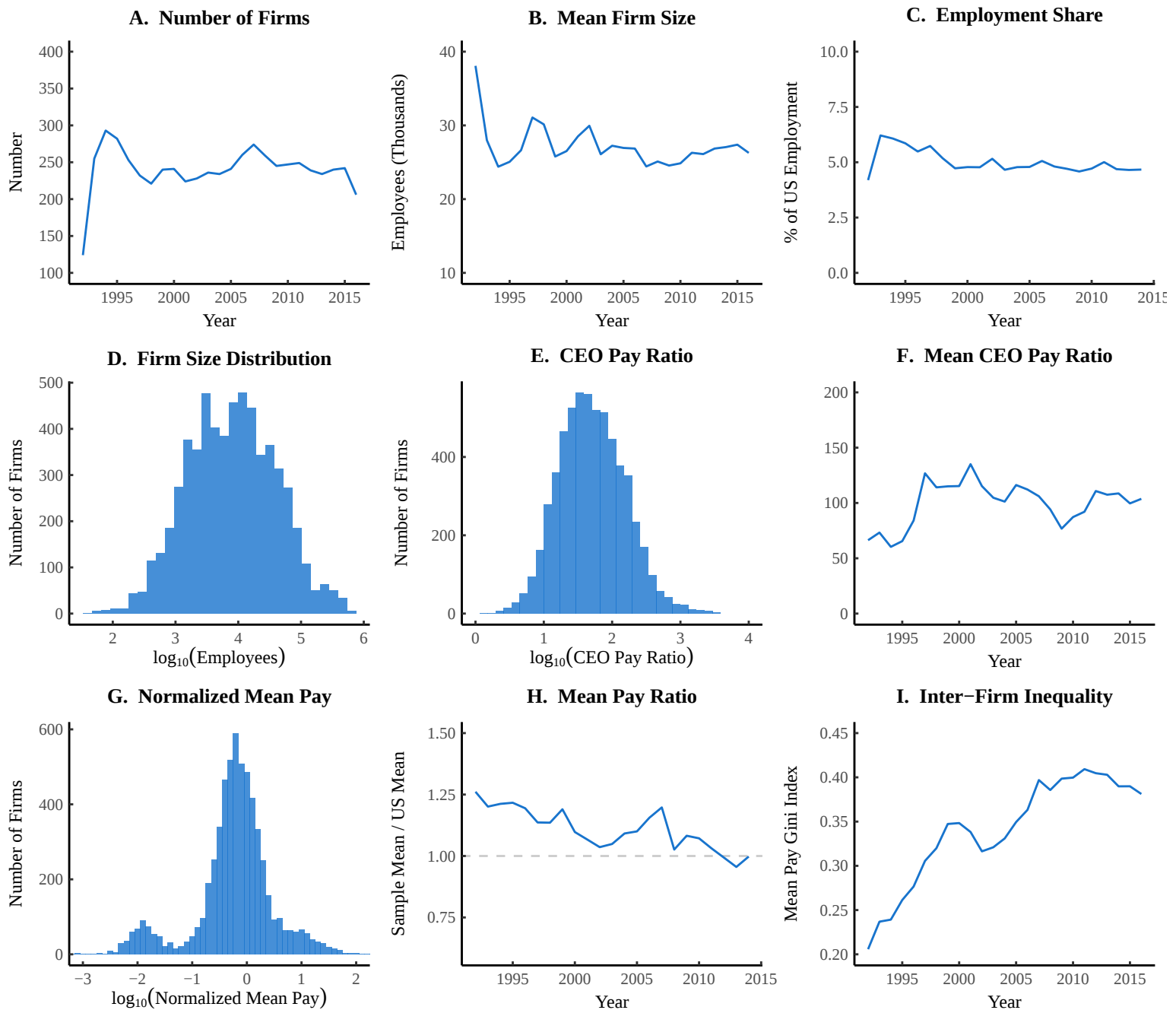


Figure 5: Selected Statistics from the Firm Sample Used for Model Tuning

This figure shows statistics for the Compustat firm sample used to tune my hierarchical model. Panel A shows the number of firms in the sample over time, Panel B the average firm size, and Panel C the share of US employment held by these firms. Panel D shows the logarithmic distribution of firm size, and Panel E shows the logarithmic distribution of the CEO pay ratio. Panel F shows the mean CEO pay ratio of all firms over time. Panel G shows the logarithmic distribution of normalized mean pay (mean pay divided by the average pay of the firm sample in each year). Panel H shows the ratio of mean pay in the Compustat sample relative to the US average (calculated from BEA Table 1.12 by dividing the sum of employee and proprietor income by the number of workers in BEA Table 6.8C-D). Panel I shows the Gini index of firm mean pay over time.

D Hierarchy Model Equations

In this section, I outline the mathematics underlying my hierarchical model of the firm. The model assumptions, outlined below, are based on the stylized facts gleaned from the real-world firm data in section B.

1. Firms are hierarchically structured, with a span of control that increases *exponentially* with hierarchical level.
2. The ratio of mean pay between adjacent hierarchical levels increases *exponentially* with hierarchical level.
3. Intra-hierarchical-level income is lognormally distributed and constant across all levels.

Using these assumptions, I first develop an algorithm that describes the hierarchical employment within a model firm, followed by an algorithm that describes the hierarchical pay structure.

Table 6: Notation

Symbol	Definition
a	span of control parameter 1
b	span of control parameter 2
C	CEO to average employee pay ratio
E	employment
F	cumulative distribution function
G	Gini index of inequality
h	hierarchical level
$\bar{I}$	average income
μ	lognormal location parameter
n	number of hierarchical levels in a firm
p	pay ratio between adjacent hierarchical levels
r	pay-scaling parameter
s	span of control
σ	lognormal scale parameter
T	total for firm
$\downarrow$	round down to nearest integer
$\prod$	product of a sequence of numbers
$\sum$	sum of a sequence of numbers

D.1 Generating the Employment Hierarchy

To generate the hierarchical structure of a firm, we begin by defining the span of control (s) as the ratio of employment (E) between two consecutive hierarchical levels (h), where $h = 1$ is the *bottom* hierarchical level. It simplifies later calculations if we define the span of control in level 1 as $s = 1$. This leads to the following piecewise function:

$$s_h \equiv \begin{cases} 1 & \text{if } h = 1 \\ \frac{E_{h-1}}{E_h} & \text{if } h \geq 2 \end{cases} \quad (6)$$

Based on our empirical findings in Section B, we assume that the span of control is *not* constant; rather it increases *exponentially* with hierarchical level. I model the span of control as a function of hierarchical level (s_h) with a simple exponential function, where a and b are free parameters:

$$s_h = \begin{cases} 1 & \text{if } h = 1 \\ a \cdot e^{bh} & \text{if } h \geq 2 \end{cases} \quad (7)$$

As one moves up the hierarchy, employment in each consecutive level (E_h) *decreases* by $1/s_h$. This yields Eq. 8, a recursive method for calculating E_h . Since we want employment to be *whole* numbers, we round down to the nearest integer (notated by $\downarrow$). By repeatedly substituting Eq. 8 into itself, we can obtain a non-recursive formula (Eq. 9). In product notation, Eq. 9 can be written as Eq. 10.

$$E_h = \downarrow \frac{E_{h-1}}{s_h} \quad \text{for } h > 1 \quad (8)$$

$$E_h = \downarrow E_1 \cdot \frac{1}{s_2} \cdot \frac{1}{s_3} \cdot \dots \cdot \frac{1}{s_h} \quad (9)$$

$$E_h = \downarrow E_1 \prod_{i=1}^h \frac{1}{s_i} \quad (10)$$

Total employment in the whole firm (E_T) is the sum of employment in all hierarchical levels. Defining n as the total number of hierarchical levels, we get Eq. 11, which in summation notation, becomes Eq. 12.

$$E_T = E_1 + E_2 + \dots + E_n \quad (11)$$

$$E_T = \sum_{h=1}^n E_h \quad (12)$$

In practice, n is not known beforehand, so we define it using Eq. 10. We progressively increase h until we reach a level of zero employment. The highest level n will be the hierarchical level directly *below* the first hierarchical level with zero employment:

$$n = \{h \mid E_h \geq 1 \text{ and } E_{h+1} = 0\} \quad (13)$$

To summarize, the hierarchical employment structure of our model firm is determined by 3 free parameters: the span of control parameters a and b , and base-level employment E_1 . Code for this hierarchy generation algorithm can be found in the C++ header files `hierarchy.h` and `exponents.h`, located in the Supplementary Material.

D.2 Generating Hierarchical Pay

To model the hierarchical pay structure of a firm, we begin by defining the inter-hierarchical pay-ratio (p_h) as the ratio of mean income ($\bar{I}$) between adjacent hierarchical levels. Again, it is helpful to use a piecewise function so that we can define a pay-ratio for hierarchical level 1:

$$p_h \equiv \begin{cases} 1 & \text{if } h = 1 \\ \frac{\bar{I}_h}{\bar{I}_{h-1}} & \text{if } h \geq 2 \end{cases} \quad (14)$$

Based on our empirical findings in Section B, we assume that the pay ratio increases *exponentially* with hierarchical level. I model this relation with the following function, where r is a free parameter:

$$p_h = \begin{cases} 1 & \text{if } h = 1 \\ r^h & \text{if } h \geq 2 \end{cases} \quad (15)$$

Using the same logic as with employment (shown above), the mean income I_h in any hierarchical level is defined recursively by Eq. 16 and non-recursively by Eq. 17.

$$\bar{I}_h = \frac{\bar{I}_{h-1}}{p_h} \quad (16)$$

$$\bar{I}_h = \bar{I}_1 \prod_{i=1}^h p_i \quad (17)$$

To summarize, the hierarchical pay structure of our model firm is determined by 2 free parameters: the pay-scaling parameter r , and mean pay in the base level ($\bar{I}_1$). Code for generating hierarchical pay can be found in the C++ header files `model.h`, located in the Supplementary Material.

D.2.1 Useful Statistics

Two statistics are used repeatedly within the model: mean firm pay, and the CEO-to-average-employee pay ratio.

Mean income for all employees ($\bar{I}_T$) is equal to the average of hierarchical level mean incomes ($\bar{I}_h$) weighted by the respective hierarchical level employment (E_h):

$$\bar{I}_T = \sum_{h=1}^n \bar{I}_h \cdot \frac{E_h}{E_T} \quad (18)$$

To calculate the CEO pay ratio, we define the CEO as the person(s) in the top hierarchical level. Therefore, CEO pay is simply $\bar{I}_n$, average income in the top hierarchical level. The CEO pay ratio (C) is then equal to CEO pay divided by average pay:

$$C = \frac{\bar{I}_n}{\bar{I}_T} \quad (19)$$

D.3 Adding Intra-Level Pay Dispersion

Up to this point, we have modeled only the *mean* income within each hierarchical level of a firm. The last step in the modeling process is to add pay *dispersion* within each hierarchical level.

I assume that pay dispersion within hierarchical levels is *lognormally* distributed. The lognormal distribution is defined by location parameter μ and scale parameter σ . Our empirical investigation of firm case studies indicated that pay dispersion with hierarchical levels is relatively constant (see Fig. 2C). Given this finding, I assume *identical inequality* within all hierarchical levels. This means that the lognormal scale parameter σ is the same for all hierarchical levels.

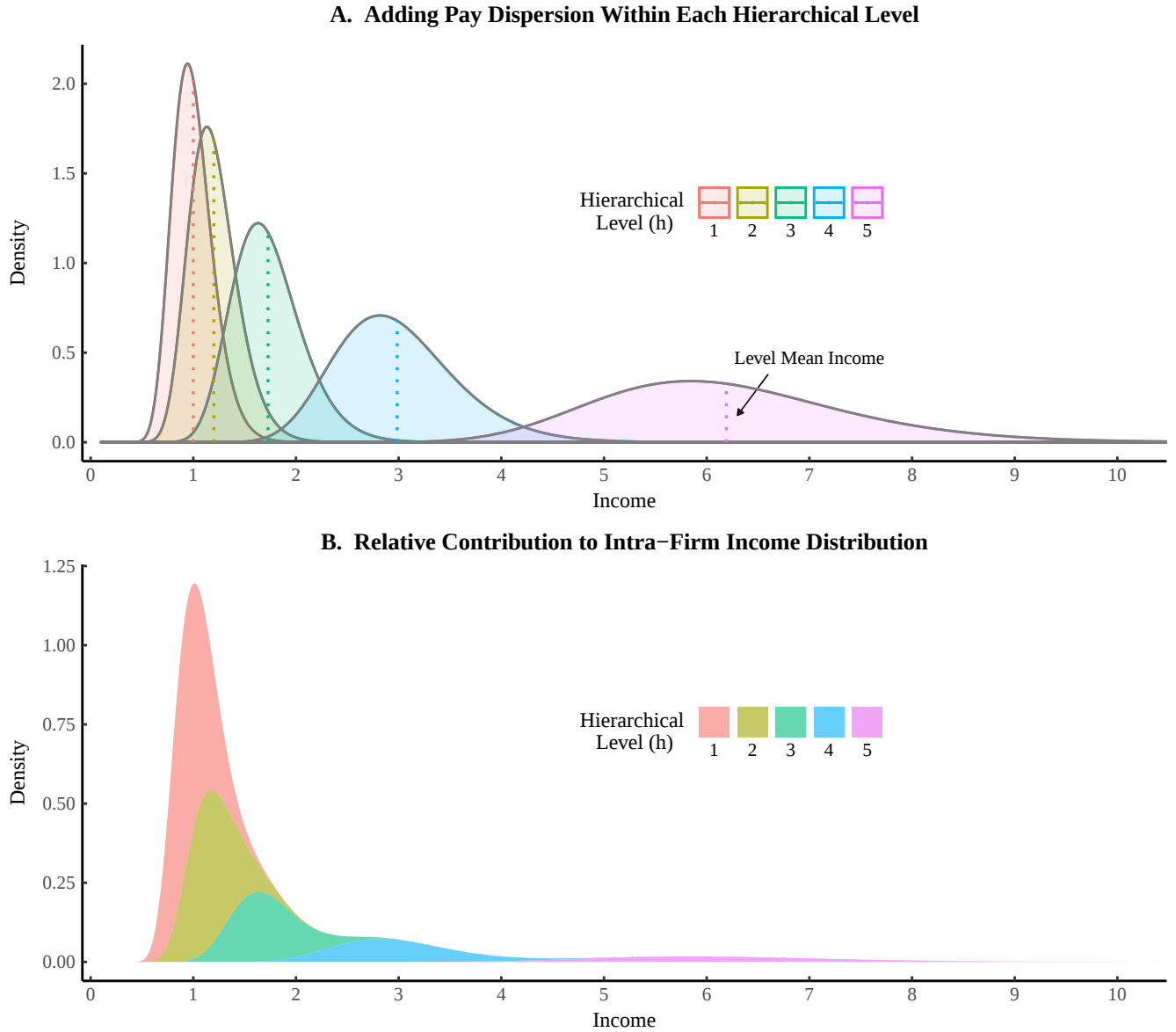


Figure 6: Adding Intra-Level Pay Dispersion to a Model Firm

This illustrates a model firm with lognormal pay dispersion in each hierarchical level. The model firm has a pay-scaling parameter of $r = 1.2$ and an intra-level Gini index of 0.13. Panel A shows the separate distributions for each level, with mean income indicated by a dashed vertical line. Panel B shows contribution of each hierarchical level to the resulting income distribution for the whole firm (income density functions are summed while weighting for their respective employment).

In order to add dispersion within each hierarchical level, I multiply mean pay $\bar{I}_h$ by a lognormal random variate with an expected mean of one. Formally, this is represented by Eq. 20. Since the mean of a lognormal distribution is equal to $e^{\mu + \frac{1}{2}\sigma^2}$, I leave it to the reader to show that a mean of one requires that μ be defined by Eq. 21.

$$I_h = \bar{I}_h \cdot \ln \mathcal{N}(\mu, \sigma) \quad (20)$$

$$\mu = -\frac{1}{2}\sigma^2 \quad (21)$$

Given a value for σ (which is a free parameter), we can define the pay distribution within any hierarchical level of a firm. This process is shown graphically in Figure 6. Figure 6A shows the lognormal income distributions for each hierarchical level of a 5-level firm. Figure 6B shows the size-adjusted contribution of each hierarchical level to the overall intra-firm income distribution. Lower levels have more members, and thus dominate the overall distribution. The code implementing this method can be found in the C++ header file `model.h`, located in the Supplementary Material.

Table 7: Model Parameters

Parameter	Definition	Action	Scope
α	Firm size distribution exponent	Determines the skewness of the firm size distribution	—
a, b	Span of control parameters	Determines the shape of the firm hierarchy.	Identical for all firms.
E_1	Employment in base hierarchical level	Used to build the employment hierarchy from the bottom up. Determines total employment.	Specific to each firm.
r	Pay-scaling parameter	Determines the rate at which mean income (within a firm) increases by hierarchical level.	Specific to each firm.
$\bar{I}_h$	Mean pay in base hierarchical level	Sets the base level income of the firm, which determines firm average pay.	Specific to each firm.
σ	Intra-hierarchical level pay dispersion parameter	Determines the level of inequality within hierarchical levels of a firm.	Identical for all firms.

E Restricting Parameters

As discussed in section D, the hierarchy model has many ‘free’ parameters. Table 7 summarizes all of the parameters used in this model. While free to take on any value, I restrict these parameters exclusively using empirical data. In the following sections, I outline the methods used for this restriction.

E.1 Firm Size Distribution

Recent studies have found that firm size distributions in the United States [23] and other G7 countries [24] can be modeled accurately with a power law. A power law has the simple form shown in Eq. 22, where the probability of observation x is inversely proportional to x raised to some exponent α :

$$p(x) \propto \frac{1}{x^\alpha} \quad (22)$$

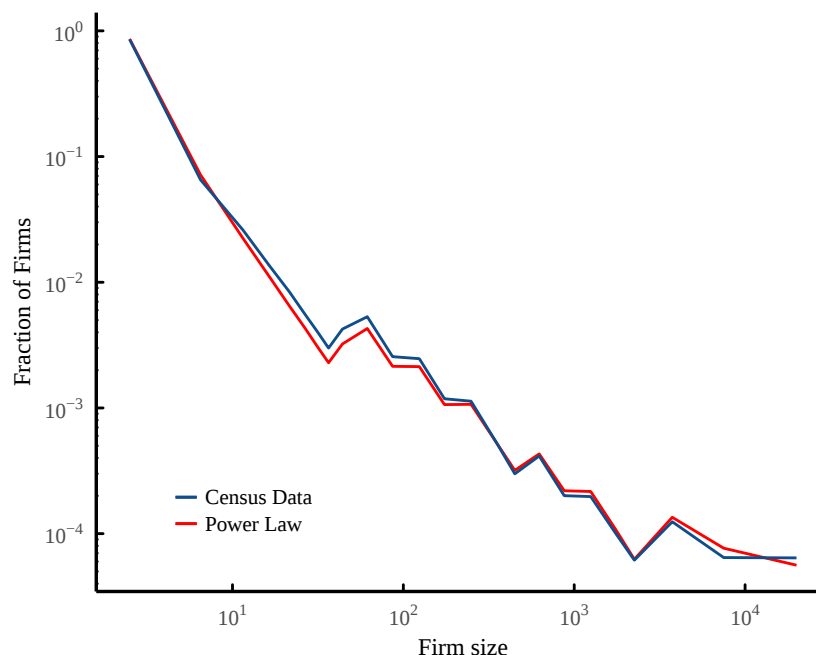


Figure 7: The United States Firm Size Distribution

This figure shows the US firm size distribution compared to a power law distribution with exponent $\alpha = 2.01$ (a simulation with 15 million firms). The US histogram combines data for ‘employer’ firms with data for unincorporated self-employed workers. Data for ‘employer’ firms is from the US Census Bureau, Statistics of U.S. Businesses (using data for 2013). This data is augmented with Bureau of Labor Statistics data for unincorporated self-employed workers (series LNU02032185 and LNU02032192). The histogram preserves Census firm-size bins, with self-employed data added to the first bin. The last point on the histogram consists of all firms with more than 10,000 employees.

Figure 7 compares the US firm size distribution with a power law of exponent $\alpha = 2.01$. Although not perfect, the fit is good enough for modeling purposes. I assume that the firm sizes can be modeled with a discrete power law random variate. I model the US firm size distribution with $\alpha = 2.01$.

A characteristic property of power law distributions is that as α approaches 2, the mean becomes *undefined*. In the present context, this means that the model can produce firm sizes that are extremely large — far beyond anything that exists in the real world. To deal with this difficulty, I *truncate* the power law distribution at a maximum firm size of 2.3 million. This happens to be the present size of Walmart, the largest US firm in existence.

Code for the discrete power law random number generator can be found in

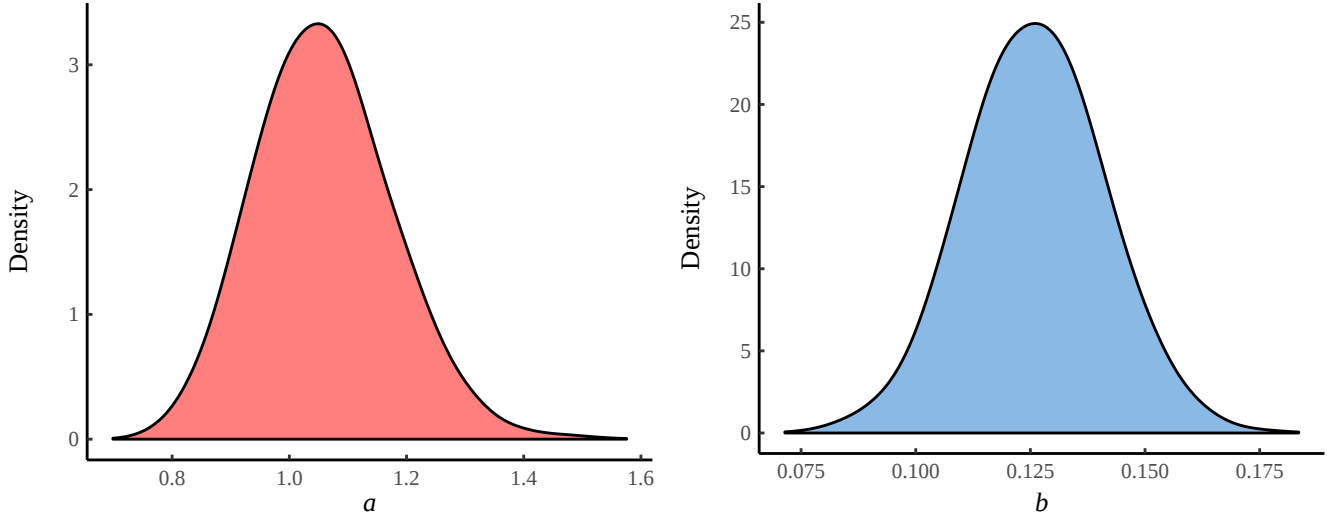


Figure 8: Density Estimates for Span of Control Parameters

This figure shows density estimates for the parameters a and b , which together determine the ‘shape’ of the firm hierarchy. These parameters are determined from regressions on firm case-study data (Fig. 2). The density functions are estimated using a bootstrap analysis, which involves resampling (with replacement) the case study data many times, and calculating the parameters a and b for each resample.

the C++ header file `rp1d.h`, located in the Supplementary Material. This code is an adaption of Collin Gillespie’s [25] discrete power law generator found in the R `powerLaw` package (which is, in turn, an adaption of the algorithm outline by Clauset [26]).

E.2 Span of Control Parameters

The parameters a and b together determine the shape of firm employment hierarchy. These parameters are estimated from an exponential regression on case study data (Fig. 2A). The model proceeds on the assumption that these parameters are *constant* across all firms.

Because the case-study sample size is small, there is considerable uncertainty in these values. I incorporate this uncertainty into the model using the *bootstrap* method [27], which involves repeatedly resampling the case-study data (with replacement) and then estimating the parameters a and b from this resample. Figure 8 shows the probability density distribution resulting from this bootstrap analysis. I run the model many times, each time with a and b determined by a bootstrap resample of case-study data.

Code implementing this bootstrap can be found in the C++ header file `boot_span.h`.

E.3 Base Level Employment

Given span of control parameters a and b , each firm hierarchy is constructed from the bottom hierarchical level up. Thus, we must know base level employment. In practice, however, we don't know this value — instead we are given *total* employment for a particular firm. While it may be possible to use the equations in section D to define an analytic function relating total employment to base level employment, this is beyond my mathematical abilities.

Instead, I use the model to reverse engineer the problem. I input a range of different base employment values into equations 7, 10, and 12 and calculate total employment for each value. The result is a discrete mapping relating base-level employment to total employment. I then use the C++ Armadillo interpolation function to linearly interpolate between these discrete values. This allows us to predict base level E_1 , given total employment E_T . Code implementing this method can be found in the C++ header file `base_fit.h`, located in the Supplementary Material.

E.4 Pay-Scaling Parameter

The pay-scaling ratio r determines the rate at which mean pay increases by hierarchical level. Unlike the span of control parameters, the pay-scaling parameter is allowed to vary between firms. But how should it vary? I restrict the variation of this parameter in a two-step process. I first 'tune' the model to Compustat data. This results in a distribution of pay-scaling parameters specific to Compustat firms. I then fit this data with a parameterized distribution, from which simulation parameters are randomly chosen.

E.4.1 Fitting Compustat Pay-Scaling Parameters

I fit the pay-scaling parameter r to Compustat firms using the CEO-to-average-employee pay ratio (C). The first step of this process is to build the employment hierarchy for each Compustat firm using parameters a , b , and E_1 (the latter is determined from total employment). Given this hierarchical employment structure, the CEO pay ratio in the modeled firm is uniquely determined by the parameter r . Thus, we simply choose r such that the model produces a CEO pay ratio that is equivalent to the empirical ratio.

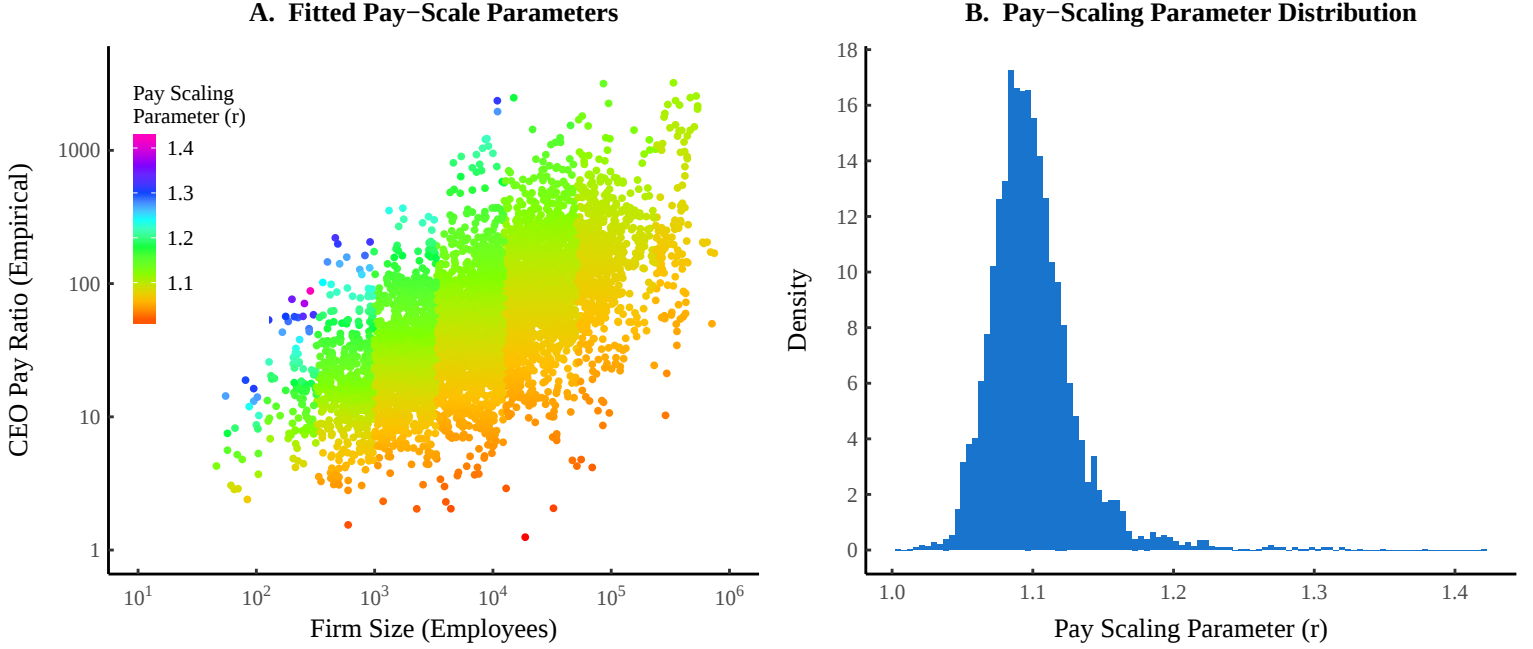


Figure 9: Fitting Compustat Firms with a Pay-Scaling Parameter

This figure shows the fitted pay-scaling parameters (r) for all Compustat firms. Panel A shows the relation between the CEO pay ratio and firm size, with the fitted pay-scaling parameter indicated by color. The discrete changes in color (evident as vertical lines) correspond to changes in the number of hierarchical levels within firms. The pay-scaling parameter distribution for all firms (and years) is shown in panel B.

To solve for this r value, I use numerical optimization (the bisection method) to minimize the error function shown in Eq. 23. Here $C_{\text{Compustat}}$ and C_{model} are Compustat and modeled CEO pay ratios, respectively.

$$\epsilon(r) = |C_{\text{model}} - C_{\text{Compustat}}| \quad (23)$$

For each firm, the fitted value of r minimizes this error function. To ensure that there are no large errors, I discard Compustat firms for which the best-fit r parameter produces an error that is larger than $\epsilon = 0.01$. Fitted results for r are shown in Figure 9. Code implementing this method can be found in the C++ header file `fit_model.h`, located in the Supplementary Material.

E.4.2 Generating a Pay Scaling Distribution

Once we have generated r parameters for every Compustat firm, the next step is to fit a parameterized distribution to this data. For Compustat firms, the dis-

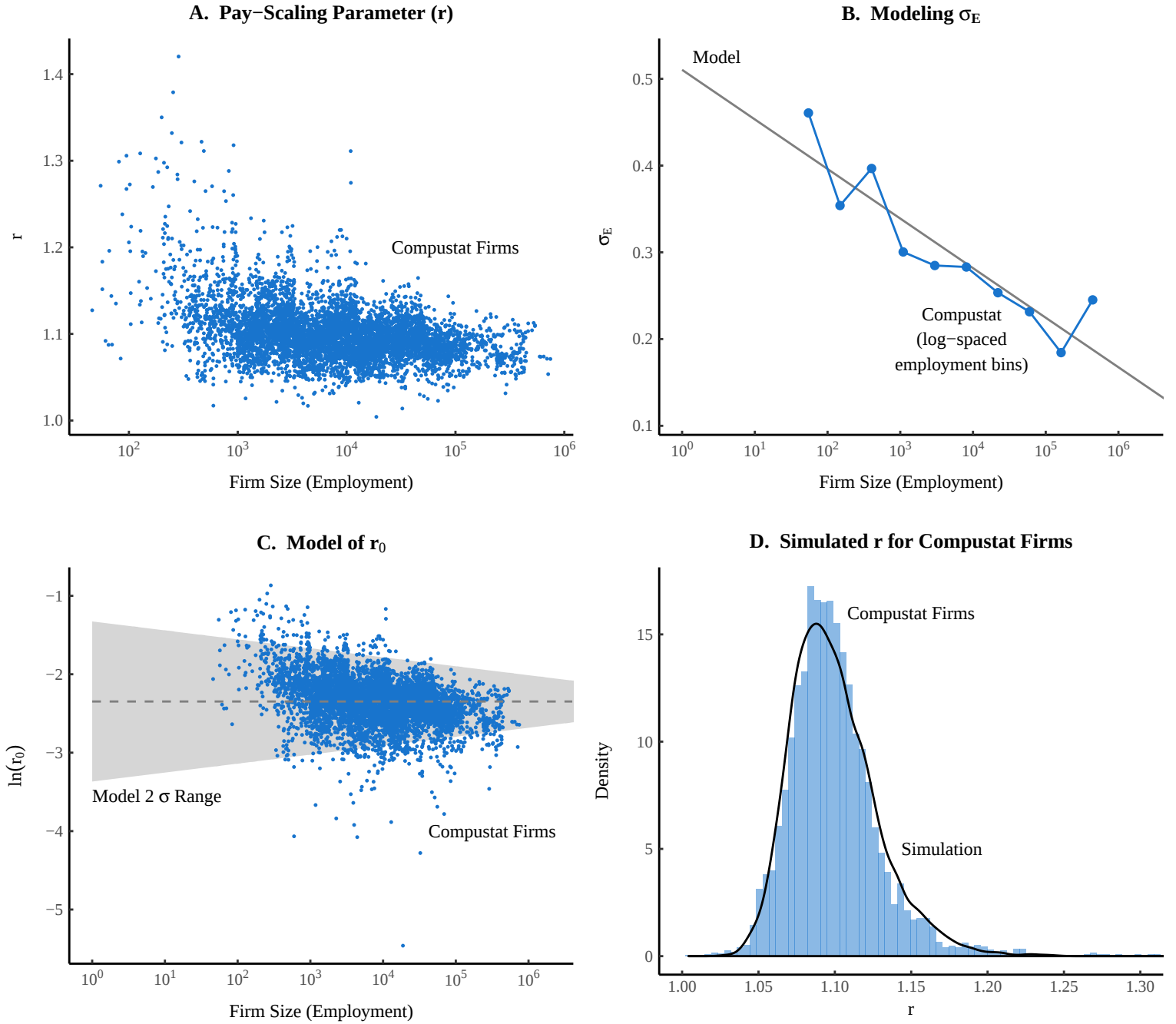


Figure 10: Modeling the Firm Pay Scaling Distribution

This figure visualizes the model used to simulate firm pay-scaling parameters (r). Panel A shows the relation between r and firm employment for Compustat firms. For the simulation, the distribution of r is modeled with the lognormal variate r_0 . Panel B shows how the lognormal scale parameter σ_E (defined by Eq. 28) changes with firm size. The straight line indicates the modeled relation. Panel C shows how the modeled dispersion of $\ln(r_0)$ declines with firm size, and how this relates to Compustat r_0 data. The 2σ range indicates 2 standard deviations from the mean (on log-transformed data). Panel D shows how the distribution of r for Compustat firms compares to the *simulated* distribution achieved by applying the model to the same Compustat firms.

person of r is approximately lognormal, and tends to decline with firm size (see Figure 10A). I model r as a shifted function of the lognormal variate r_0 :

$$r = 1 + \ln \mathcal{N}(r_0) \quad (24)$$

The lognormal variate r_0 is defined by location parameter μ and scale parameter σ . While μ is assumed to be constant for all firms, σ is a function of firm size E :

$$r_0(E) = \ln \mathcal{N}(r_0; \mu, \sigma_E) \quad (25)$$

I use the tuned Compustat data to solve for the parameters μ and σ . We first transform Compustat r values using Eq. 26 to get the Compustat distribution of r_0 :

$$r_0 = r - 1 \quad (26)$$

The best-fit value for μ is defined by taking the mean of $\ln(r_0)$:

$$\mu = \overline{\ln(r_0)} \quad (27)$$

Similarly, we can solve for the best-fit value for σ by taking the standard deviation of $\ln(r_0)$. However, unlike μ , the value σ will depend on the size range of firms (E):

$$\sigma_E = \text{SD} [\ln(r_0)]_E \quad (28)$$

Figure 10B plots σ_E vs. E for logarithmically spaced size groupings of Compustat firms. I model this relation using a log-linear regression. Figure 10C shows how the modeled dispersion in r_0 varies with firm size, and how this compares to Compustat data.

Once we have fitted the parameters μ and σ to the tuned Compustat data, we can generate r values for simulated firms using equations 24 and 25. Although the model is simple, it produces reasonably accurate results. To test this accuracy, we can apply the model to the same Compustat firms for which it is ‘tuned’. For each Compustat firm, we use the method outlined above to stochastically generate a pay-scaling value r . As Figure 10D shows, the resulting simulated distribution of r fairly accurately reproduces the original data.

When we move from simulating Compustat firms to a real-world distribution of firms, this model involves significant extrapolations for small firms. Why?

The Compustat firm sample has very few observations for firms smaller than 100. And those small firms that are included in the sample are likely *not* representative of the wider population, since they are small *public* firms. In the real world, virtually all small firms are *private*. As with all extrapolations, we simply do the best with the data that is available, while noting that better data might render the extrapolation moot. The code implementing this model can be found in the C++ header file `r_sim.h`, located in the Supplementary Material.

E.5 Base-Level Mean Pay

As with the pay-scaling parameter, base level mean pay varies across firms. How should it vary? Again, I restrict the variation of this parameter in a two-step process. I first ‘tune’ the model to Compustat data. This results in a distribution of base pay specific to Compustat firms. I then fit this data with a parameterized distribution, from which simulation parameters are randomly chosen.

E.5.1 Fitting Compustat Base Level Pay

Having already fitted a hierarchical pay structure to each Compustat firm (in the process of finding r), we can use this data to estimate base pay for each firm. To do this, we set up a ratio between base level pay ($\bar{I}_1$) and firm mean pay ($\bar{I}_T$) for both the model and Compustat data:

$$\frac{\bar{I}_1^{\text{Compustat}}}{\bar{I}_T^{\text{Compustat}}} = \frac{\bar{I}_1^{\text{model}}}{\bar{I}_T^{\text{model}}} \quad (29)$$

The modeled ratio between base pay and firm mean pay ($\bar{I}_1^{\text{model}}/\bar{I}_T^{\text{model}}$) is *independent* of the choice of base pay. This is because the modeled firm mean pay is actually a *function* of base pay (see Eq. 17 and 18). If we run the model with $\bar{I}_1^{\text{model}} = 1$, then Eq. 29 reduces to:

$$\frac{\bar{I}_1^{\text{Compustat}}}{\bar{I}_T^{\text{Compustat}}} = \frac{1}{\bar{I}_T^{\text{model}}} \quad (30)$$

We can then rearrange Eq. 30 to solve for an estimated base pay for each Compustat firm ($\bar{I}_1^{\text{Compustat}}$):

$$\bar{I}_1^{\text{Compustat}} = \frac{\bar{I}_T^{\text{Compustat}}}{\bar{I}_T^{\text{model}}} \quad (31)$$

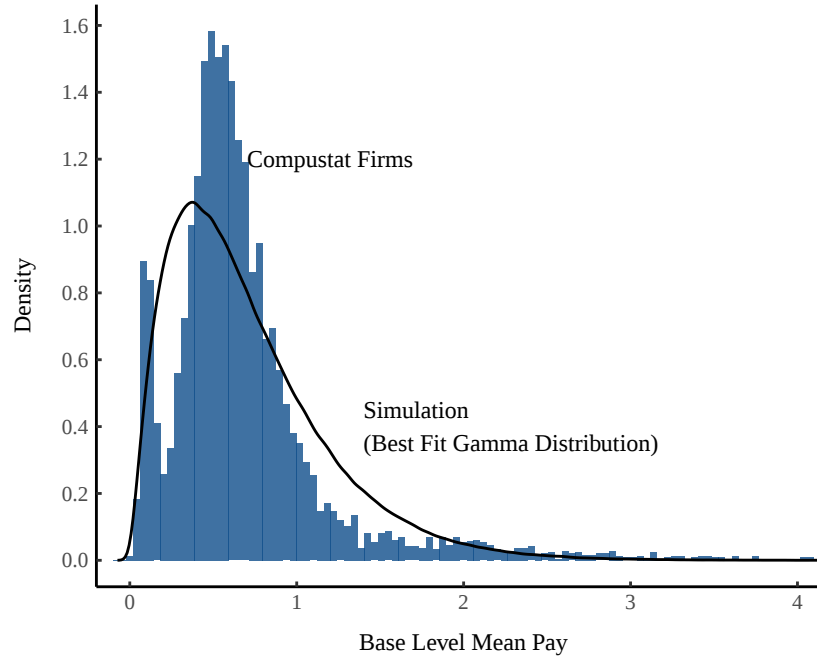


Figure 11: Modeling Firm Base Level Mean Pay

This figure shows the distribution of fitted base-level mean pay for Compustat firms. I model this data with a gamma distribution, from which simulated firm base-level mean pay is randomly drawn. Note that fitting the unimodal gamma distribution to the bimodal Compustat data means that the fit is not great. (The gamma distribution does fit the data better than other skewed distributions such as the Weibull or lognorma). The lower mode in the Compustat data is likely not representative of the general firm population. This lower mode is made up almost entirely of chain restaurants, which seem to be over-represented in this sample.

Code implementing this method is found in the C++ header file `fit_model.h`, located in the Supplementary Material.

E.5.2 Generating a Base Pay Distribution

Once each Compustat firm has a fitted value for base-level mean pay, we fit this data with a parametric distribution which is then used to stochastically generate base-level mean pay for the simulation. Since Compustat data is comprised of observations over *multiple* years, in order to aggregate this data into a single distribution, we must account for inflation. Rather than use a price index like the GDP deflator, I divide all firm mean pay data by the average Compustat mean pay in the appropriate year. Since our simulation is concerned only with relative

incomes (rather than absolute incomes) no pertinent information is lost in this process.

I model the Compustat firm base pay distribution with a gamma distribution (Fig. 11). Note that because the Compustat data has a bimodal structure (that I do not aim to replicate), the gamma distribution is not a particularly strong fit. Nonetheless the gamma model closely replicates the inequality of firm base pay (which has a Gini index of roughly 0.35). Code implementing this model can be found in the C++ header file `base_pay_sim.h` (in the Supplementary Material).

E.6 Intra-Hierarchical Level Income Dispersion

Intra-hierarchical level income dispersion is modeled with a lognormal distribution, with the amount of inequality determined by the scale parameter σ . I estimate σ from the case-study data shown in Figure 2C. This data uses the Gini index as the metric for dispersion.

To estimate σ , we first calculate the mean Gini index of all data ($\bar{G}$). We then use Eq. 32 to calculate the value σ , which corresponds to the lognormal scale parameter that would produce a lognormal distribution with an equivalent Gini index. This equation is derived from the definition of the Gini index of a lognormal distribution: $G = \text{erf}(\sigma/2)$.

$$\sigma = 2 \cdot \text{erf}^{-1}(\bar{G}) \quad (32)$$

The model proceeds on the assumption that σ is constant for all hierarchical levels within all firms. Because the case-study sample size is small, there is considerable uncertainty in these values. I quantify this uncertainty using the *bootstrap* method [27], which involves repeatedly resampling the case-study data (with replacement) and then estimating the parameter σ from this resampled data.

Figure 12 shows the probability density distribution resulting from this bootstrap analysis. In order to incorporate this uncertainty, I run the model many times, with each run using a different bootstrapped value for σ . Code implementing this method can be found in the C++ header file `boot_sigma.h`, located in the Supplementary Material.

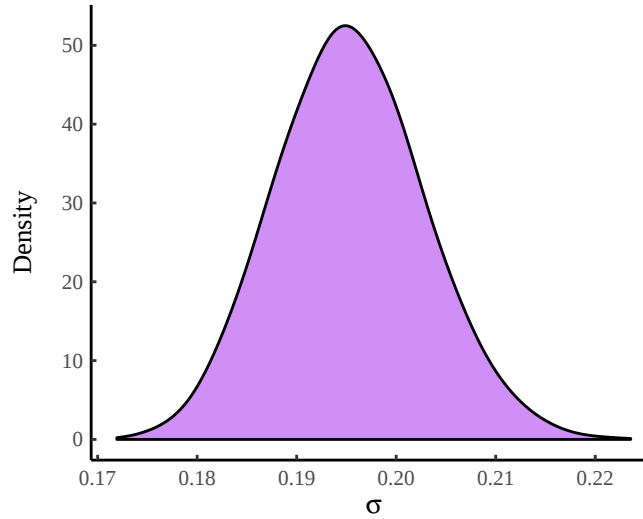


Figure 12: Density estimates for Intra-Hierarchical Level Pay Dispersion Parameter σ

This figure shows the distribution of the lognormal scale parameter σ , which determines pay dispersion within all hierarchical levels of all firms. The distribution is calculated using the bootstrap method.

E.7 Counterfactual Models

To isolate the distributional effects of hierarchy, I create three counterfactual models, each with only one income-dispersion source. This is achieved as follows:

Inter-firm dispersion only: To create this model, I set the hierarchical pay-scaling parameter (r) to 1 for all firms (removing hierarchical pay-scaling) and set the intra-hierarchical dispersion parameter (σ) to zero (removing dispersion within hierarchical levels).

Inter-hierarchical dispersion only: To create this model, I set base-level pay ($\bar{I}_1$) in all firms to an identical constant (removing dispersion between firms), and set the intra-hierarchical dispersion parameter (σ) to zero (removing dispersion within hierarchical levels).

Intra-hierarchical dispersion only: To create this model, I set base-level pay ($\bar{I}_1$) in all firms to an identical constant (removing dispersion between firms),

set the hierarchical pay-scaling parameter (r) to 1 for all firms (removing hierarchical pay-scaling).

E.8 Summary of Model Structure

The model is implemented in C++ using a modular design. Each major task is carried out by a separate function that is defined in a corresponding header file. Table 8 summarizes this structure sequentially in the order that functions are called. In each step, I briefly summarize the action that is performed, giving reference to the section where this action is described in detail.

Table 8: Model High-Level Structure

Step	Action	Reference Section	Parameter(s)	Header File(s)
1	Bootstrap case-study data	E.2, E.6	a, b, σ	boot_span.h boot_sigma.h
2	Get Compustat base-level employment	E.3	E_1	base_fit.h
3	Fit Compustat pay-scaling parameters	E.4.1	r	fit_model.h
4	Get Compustat base-level mean pay	E.5.1	$\bar{I}_1$	fit_model.h
5	Generate power law firm size distribution	E.1	α	rp1d.h
6	Get simulation base-level employment	E.3	E_1	base_fit.h
7	Simulate pay-scaling parameter distribution by fitting Compustat data	E.4.2	r	r_sim.h
8	Simulate base mean pay distribution by fitting Compustat data	E.5.2	$\bar{I}_1$	base_pay_sim.h
9	Run hierarchy model	D	all	model.h

Notes: Model code makes extensive use of Armadillo, an open-source C++ linear algebra library [\[28\]](#).

F The Adjusted Hierarchy Model

The hierarchy model tends to underestimate US income inequality. This may be caused by the model's reliance on Compustat Firm data (see Appendix E), which is biased towards large firms. The result is that the model likely has too little inter-firm income dispersion. Here I present the results of an *adjusted* model in which inter-firm income dispersion is increased so that the model closely reproduces US macro-level data.

As outlined in Appendix E, inter-firm income dispersion is modeled by fitting a gamma distribution to Compustat firm data. The gamma distribution has the following probability density function:

$$p(x) = \frac{1}{\Gamma(k)\theta^k} \cdot x^{k-1} \cdot e^{-x/\theta} \quad (33)$$

In the original model, the parameters k and θ are both determined by empirical data. In the adjusted model, I introduce a fudge-factor c that allows me to adjust the fitted k parameter by a constant amount:

$$k_{\text{adjust}} = c \cdot k_{\text{fit}} \quad (34)$$

The adjusted model then uses the parameter k_{adjust} instead of k_{fit} . All of the model's other parameters remain constant. Note that for $c > 1$, inter-firm dispersion is *decreased* (relative to the original model). For $c < 1$, inter-firm dispersion is *increased*. I choose the value c so that the adjusted model produces the best match to US data. Model results for $c = 0.5$ are shown in Figure 13. By increasing inter-firm dispersion, we significantly improve model's fit to the body of the US distribution of income. Note that the adjusted model's Gini index is significantly higher than in the original model, and now better matches US data. Results in the tail remain virtually unchanged. (This is expected, since hierarchy shapes the tail).

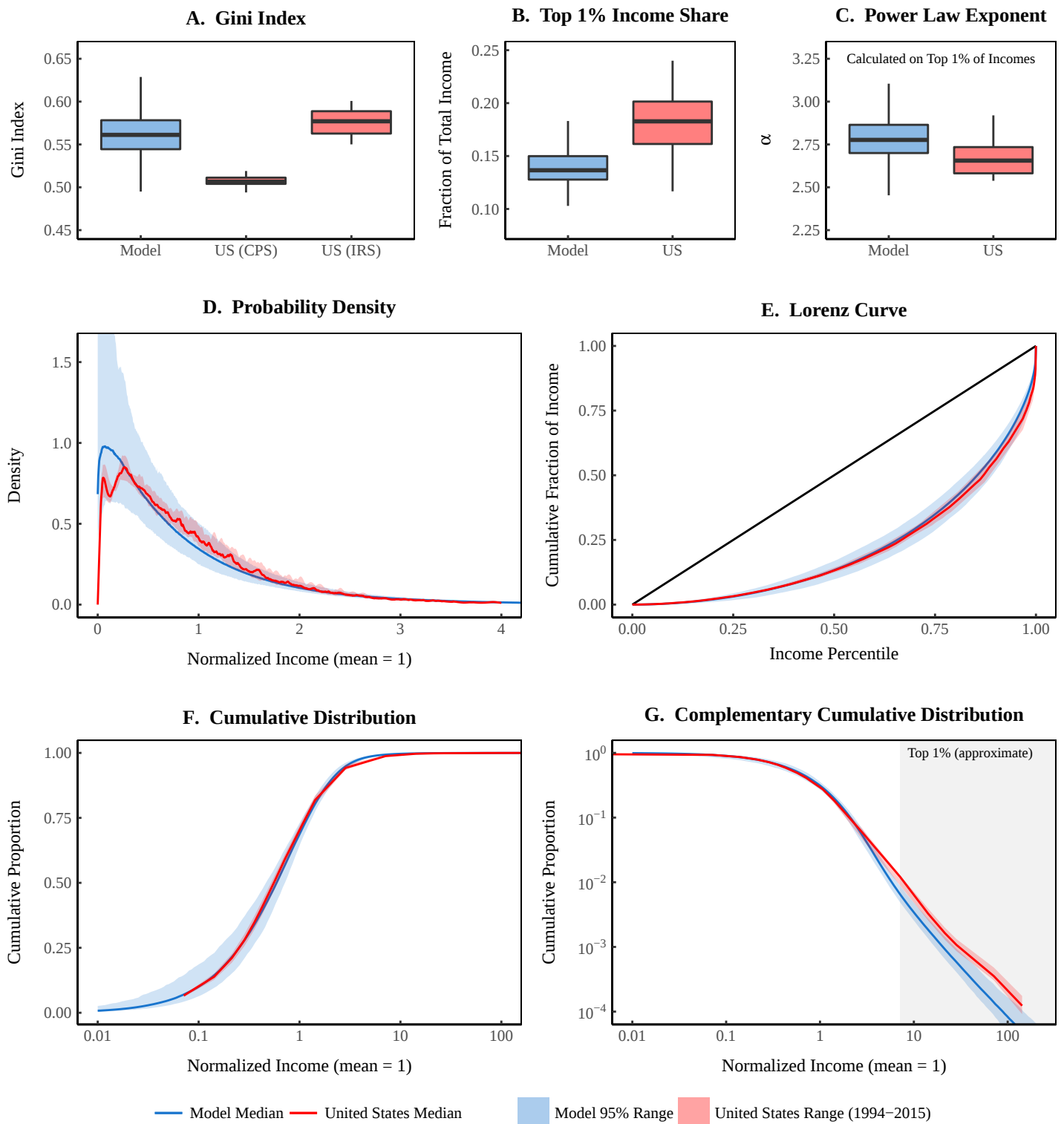


Figure 13: Adjusted Model Income Distribution vs. US Data

This figure compares various aspects of the *adjusted* model's income distribution to US data over the years 1992-2015. The adjusted model has increased inter-firm income dispersion relative to the original model. Panel A shows the Gini index, with two different US sources — the Current Population Survey (CPS) and the Internal Revenue Service (IRS). Panel B shows the top 1% income share, using data from 17 different time series. Panel C shows the results of fitting a power law distribution to the top 1% of incomes (where α is the scaling exponent). Panel D plots the income density

curve with mean income normalized to 1 (using data from the CPS). Panels E, F, and G use IRS data to construct the Lorenz curve, cumulative distribution, and complementary cumulative distribution (respectively). The cumulative distribution shows the proportion of individuals with income *less* than the given x value. The complementary cumulative distribution shows the proportion of individuals with income *greater* than the given x value. Note the log scale on the x -axis for these last two plots. For sources and methods, see Appendix A.

G A Null-Effect Model for US Top Incomes and Firm Size

A key prediction of the hierarchy model is that top incomes should be concentrated at the top of large institutions. To test this prediction, I look at the size distribution of firms associated with top incomes. Here I develop a null-effect model, which is what we would expect to find if there is absolutely no relation between firm membership and income. In the null-effect case, we should find that the size distribution of firms associated with *top earners* is exactly the same as the size distribution of firms associated with the general population.

To determine the null-effect we must find the size distribution of firms associated with the general population. Before doing so, some clarification is in order. What we are talking about is the size distribution of firms associated with *individuals*. As shown in Figure 14, this is quite different from the firm-size distribution. To determine the firm-size distribution, each firm is counted *once*. However, when we map firm size to individuals, each firm is weighted by the number of individuals within it. When we do this, we are really looking at the distribution of *employment* by firm size. So what is this distribution?

If we randomly select an individual from the private sector population, let $p(i_x)$ be the probability that this individual is associated with a firm of size x . This probability will determine the size distribution of firms associated with a random sample of individuals. Let $p(x)$ be the probability of randomly selecting a firm of size x from the firm population. Using Figure 14 for guidance, we can see that $p(i_x)$ is given by:

$$p(i_x) \sim x \cdot p(x) \quad (35)$$

If we know $p(x)$ — the probability distribution of firms — we can use Eq. 35 to predict the firm-size distribution associated with a random sample of individuals. Let's do so for the United States. The US firm-size distribution can be approximated by the power-law distribution $p(x) \sim x^{-2}$ (see Appendix E). Substituting this into Eq. 35 gives:

$$p(i_x) \sim x^{-1} \quad (36)$$

Because firm sizes generally span many orders of magnitude, it is more convenient to look at the log transformation of Eq. 36. Therefore, we want to know

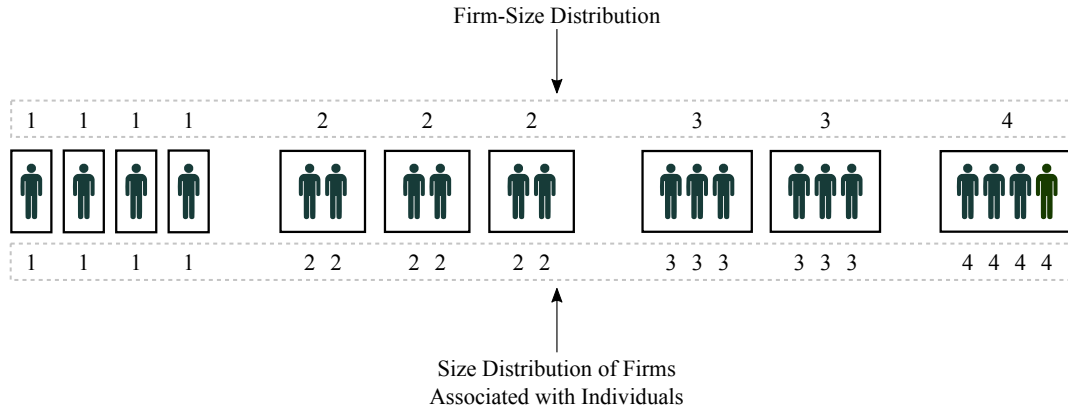


Figure 14: Mapping Firm Sizes to Individuals

This figure illustrates the mapping of firm size to individuals. Each box represents a firm, with size indicated above. The mapping of firm size to individuals appears below each firm. Let $p(x)$ be the probability of randomly selecting a *firm* of size x from the firm population. Let $p(i_x)$ be the probability of randomly selecting an *individual* associated with a firm of size x (from the individual population). Noting that each firm size x appears x times in the individual-to-firm mapping, we can state that $p(i_x) \propto x \cdot p(x)$.

the probability density for $p(\ln i_x)$. To find this, we use the standard change-of-variable function for a probability density:

$$f_y = f_x(x(y)) \cdot |x'(y)| \quad (37)$$

We let $f_y = p(\ln i_x)$ and $f_x = c \cdot x^{-1}$ (where c is constant). The transformation function is $y = \ln x$. We then note that $x(y) = e^y$ and $x'(y) = e^y$. Substituting into Eq. 37 gives:

$$f_y = c \cdot (e^y)^{-1} \cdot e^y = c \quad (38)$$

Since $f_y = p(\ln i_x)$, we can state that $p(\ln i_x) = c$, the uniform distribution. If we randomly draw a sample of individuals from the US private sector, we predict that their associated firm-size distribution will be *log-uniform*. This is the null-effect. If there is absolutely no relation between income and firm membership, we should find that the size distribution of firms associated with top incomes (in the US) is log-uniformly distributed.

H The Effect of Hierarchy on Inequality

An interesting question to ask is — what effect does hierarchy have on income inequality? In this section, I isolate the inequality effects of hierarchy using the counterfactual models of the United States. Each model contains only *one* of the three sources of income dispersion used in the original model. By comparing these counterfactual models to the original model, we can determine how each dispersion source affects income inequality.

The results in Figure 15 indicate that hierarchy’s effect on inequality depends on *how* we measure inequality. When using the Gini index (Figure 15A), we see that the model with inter-firm dispersion has inequality that is closest to the original model. (The model with inter-hierarchical dispersion comes a distant second). This suggests that hierarchy does not have a particularly strong effect on inequality. However, things change drastically when we switch to measuring inequality in terms of the income share of the top 1% (Fig. 15B). Now we find that the model with inter-hierarchical dispersion has inequality that is nearly identical to the original model. The other two sources of dispersion are inconsequential. How can this be?¹

To understand this apparent contradiction, we can look at the Lorenz curves for each model (Fig. 15C). The Lorenz curve offers a convenient way to visualize the ‘shape’ of inequality. The curve traces the cumulative fraction of income held by all individuals below a given income percentile. The Gini index and the top 1% income share are both intimately related to the Lorenz curve. The Gini index is proportional to the area between the Lorenz curve and the line of perfect equality (the black line in Fig. 15C). The income share of the top 1% is

¹ Some readers may note that I am using non-decomposable metrics to measure inequality. Since neither the Gini index nor the top 1% income share is decomposable, the inequality of the counterfactual models will *not* sum to the inequality of the original model. Thus we cannot quantify exactly ‘how much’ each factor contributes to income inequality. Although there are inequality metrics that are decomposable (such as the Theil index, or simply the variance), I choose not to use them here. For starters, such measures are generally far less intuitive than the Gini index or top income shares. Second, decomposable measures merely give *a* decomposition of inequality — not *the* decomposition. Decomposition requires deciding how to weight the number of incomes of a given size against the size of the income. Since there are many ways to do this, there are many equally valid decompositions of inequality. Anthony Shorrocks [29] summarizes the problem nicely: “Inequality comparisons are invariably sensitive to the choice of inequality index used since alternative measures tend to emphasize inequality at different points in the distribution. Replacing one index by another will therefore almost always change the relative significance of the between- and with-group terms”.

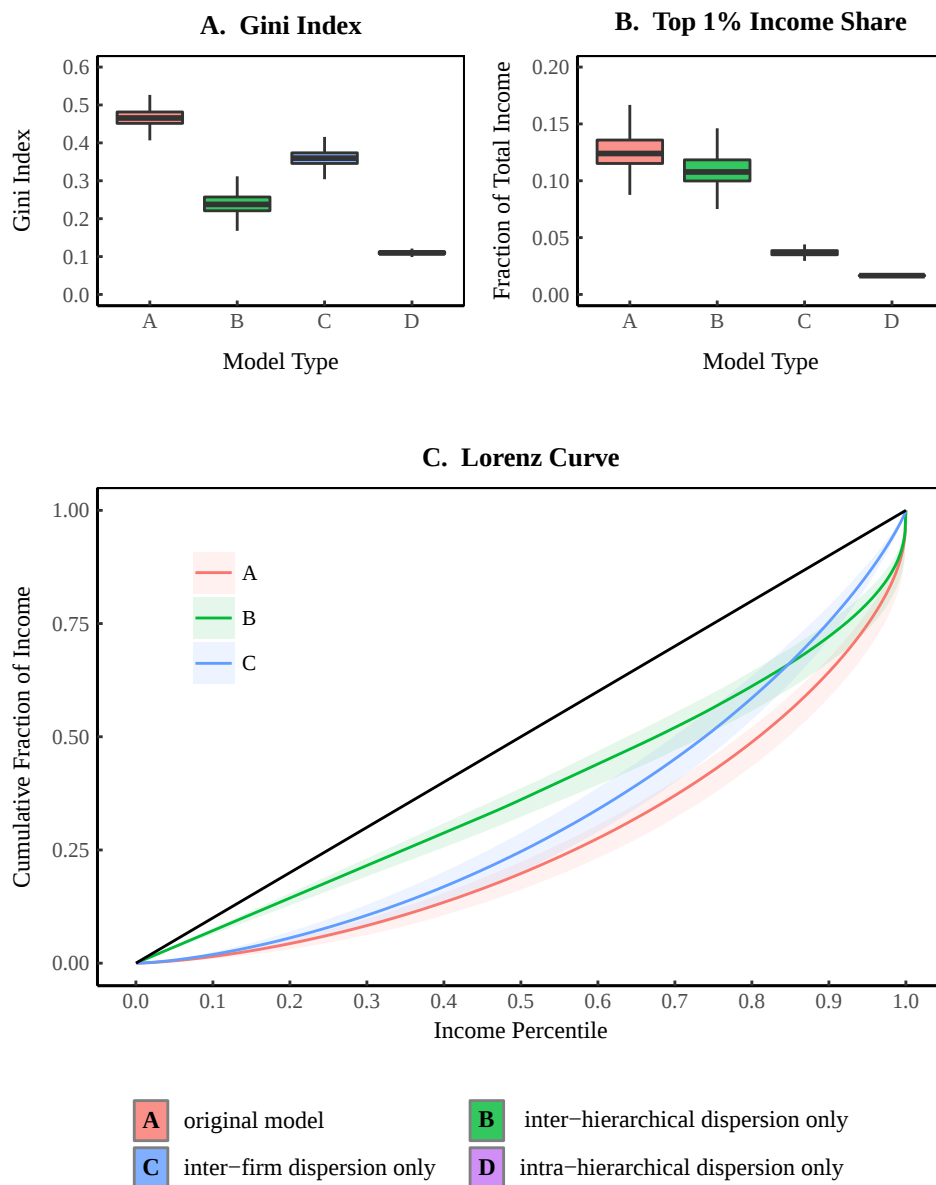


Figure 15: How Hierarchy Affects Inequality

This figure compares the original hierarchy model of the United States to three different counterfactual models. Each counterfactual model contains only one of the three sources of income dispersion. Panel A compares the Gini index of each model, while Panel B compares the top 1% income share. Note that since both of these inequality metrics are not additive, the inequality in the counterfactual models will not sum to the inequality in the original model. Panel C shows the Lorenz curve for each model, with shaded regions indicating the 95% range. For clarity (and because it plays a negligible role determining income distribution), the intra-hierarchical dispersion model is not shown in Panels C.

equal to the vertical distance between the Lorenz curve and $y = 1$, at the point $x = 0.99$.

The apparent contradiction between the Gini and top 1% results is now easy to understand. It is caused by an *intersection* between the inter-firm Lorenz curve and the inter-hierarchical level Lorenz curve. For incomes *below* this intersection, inter-firm dispersion plays the most important role in shaping inequality. However, for incomes *above* the intersection, hierarchy plays the most important role in shaping inequality.

These results reinforce those in the main paper. Hierarchy is important for shaping the tail of the distribution (the top 1%), while dispersion between firms shapes the rest of the distribution. These results also demonstrate the pitfalls of using a single metric to quantify inequality. No single metric can capture all of the information in a Lorenz curve. The Gini index places an emphasis on the body of the distribution, while top income fraction metrics capture the dynamics of the tail. The hierarchy model suggest that when we study top income shares, we are studying the effects of firm hierarchy.

References

1. Current Population Survey. Annual Social and Economic Supplement; 2017.
2. Virkar Y, Clauset A. Power-law distributions in binned empirical data. *The Annals of Applied Statistics*. 2014;8(1):89–119.
3. Alvaredo F, Atkinson AB, Piketty T, Saez E, Zucman G. The world wealth and income database. Website: <http://www.wid.world>. 2016;.
4. Alvaredo F, Atkinson T, Piketty T, Saez E. The World Top Incomes Database; 2013. Available from: <http://topincomes.parisschoolofeconomics.eu/#>.
5. Atkinson A, Lakner C. Capital and labor: the factor income composition of top incomes in the United States, 1962-2006. World Bank Policy Research Working Paper. 2017;8268.
6. Piketty T. *Capital in the Twenty-first Century*. Cambridge: Harvard University Press; 2014.
7. Audas R, Barmby T, Treble J. Luck, effort, and reward in an organizational hierarchy. *Journal of Labor Economics*. 2004;22(2):379–395.
8. Baker G, Gibbs M, Holmstrom B. Hierarchies and compensation: A case study. *European Economic Review*. 1993;37(2-3):366–378.
9. Dohmen TJ, Kriechel B, Pfann GA. Monkey bars and ladders: The importance of lateral and vertical job mobility in internal labor market careers. *Journal of Population Economics*. 2004;17(2):193–228.
10. Grund C. The wage policy of firms: comparative evidence for the US and Germany from personnel data. *The International Journal of Human Resource Management*. 2005;16(1):104–119.
11. Lima F. Internal Labor Markets: A Case Study. FEUNL Working Paper. 2000;378.
12. Morais F, Kakabadse NK. The Corporate Gini Index (CGI) determinants and advantages: Lessons from a multinational retail company case study. *International Journal of Disclosure and Governance*. 2014;11(4):380–397.
13. Treble J, Van Gameren E, Bridges S, Barmby T. The internal economics of the firm: further evidence from personnel data. *Labour Economics*. 2001;8(5):531–552.

14. Ariga K, Brunello G, Ohkusa Y, Nishiyama Y. Corporate hierarchy, promotion, and firm growth: Japanese internal labor market in transition. *Journal of the Japanese and International Economies*. 1992;6(4):440–471.
15. Bell B, Van Reenen J. Firm performance and wages: evidence from across the corporate hierarchy. *CEP Discussion paper*. 2012;1088.
16. Eriksson T. Executive compensation and tournament theory: Empirical tests on Danish data. *Journal of Labor Economics*. 1999;17(2):262–280.
17. Heyman F. Pay inequality and firm performance: evidence from matched employer–employee data. *Applied Economics*. 2005;37(11):1313–1327.
18. Leonard JS. Executive pay and firm performance. *Industrial and Labor Relations Review*. 1990;43(3):13–28.
19. Main BG, O'Reilly III CA, Wade J. Top executive pay: Tournament or teamwork? *Journal of Labor Economics*. 1993;11(4):606–628.
20. Mueller HM, Ouimet PP, Simintzi E. Within-Firm Pay Inequality. *SSRN Working Paper*. 2016;.
21. Rajan RG, Wulf J. The flattening firm: Evidence from panel data on the changing nature of corporate hierarchies. *The Review of Economics and Statistics*. 2006;88(4):759–773.
22. Tao HL, Chen IT. The level of technology employed and the internal hierarchical wage structure. *Applied Economics Letters*. 2009;16(7):739–744.
23. Axtell RL. Zipf distribution of US firm sizes. *Science*. 2001;293:1818–1820.
24. Gaffeo E, Gallegati M, Palestrini A. On the size distribution of firms: additional evidence from the G7 countries. *Physica A: Statistical Mechanics and its Applications*. 2003;324(1–2):117–123. doi:10.1016/S0378-4371(02)01890-3.
25. Gillespie CS. Fitting heavy tailed distributions: the powerLaw package. *arXiv preprint arXiv:14073492*. 2014;.
26. Clauset A, Shalizi CR, Newman ME. Power-law distributions in empirical data. *SIAM review*. 2009;51(4):661–703.
27. Efron B, Tibshirani RJ. *An introduction to the bootstrap*. London: CRC press; 1994.

-
28. Sanderson C, Curtin R. Armadillo: a template-based C++ library for linear algebra. *Journal of Open Source Software*. 2016;.
 29. Shorrocks AF. The class of additively decomposable inequality measures. *Econometrica: Journal of the Econometric Society*. 1980;48(3):613–625.