

The Rise of Meritocracy and the Inheritance of Advantage

David Comerford

University of Strathclyde

José V. Rodríguez Mora*

University of Edinburgh & CEPR

Michael J. Watts

University of Edinburgh

April, 2017

Abstract

We present a model human capital accumulation with statistical discrimination on the background of agents. Firms do not observe the productivity of workers, they determine their wage after observing two signals. A direct signal on the productivity of the workers and a second signal on the income of the parents of the agent. The signal on background is useful (insofar the direct signal is not perfectly accurate) because parents invest in their children education, increasing their human capital. Knowing that richer parents invest more in education and that their children will be in average more productive, firms rationally discriminate favoring those with perceived good backgrounds. Thus, the income of an agent depends on both the objective information that society has on its human capital and the existing perceptions on her background. Our main result is that the accuracy of both signals have very similar effects in both inequality and intergenerational mobility. More accurate information on the background of individuals facilitates statistical discrimination increasing inequality and its persistence. More accurate information on the actual capabilities of workers leads to exactly the same result via an interesting feed back mechanism. It allows to discern good workers from the able ones, increasing their wage differential and inequality. But more inequality increases the uncertainty that firms have on workers characteristics, which itself increases the weight that firms give to signals, which further increases inequality, and so on. There is, though, a big difference in both types of signals in the manner that they affect investment. Accurate direct information on productivity is directly linked to the return to investment in human capital, as it is what connects human capital with income. The effects of information on parental background on investment are much more indirect, as it implies that your income affects your children independently from the investment you make in them. Using our model to interpret the data suggests that a country like the US (with relatively little intergenerational mobility, large degree of inequality and very large investment in education) might be characterized by having very precise direct signals on ability, resulting in high rewards to merit, and (conditioning on it) little rewards to have a good background. The US might be still a land of opportunities *conditional* in having the right ability. Of course, those abilities are to a large extent inherited, and that is the reason why the intergenerational persistence of income is so large. But notice that the reason why they are so inheritable might precisely be that merit is highly rewarded; that in a deep sense, the US may still be a very meritocratic society.

*Corresponding author. School of Economics, University of Edinburgh, 31 Buccleuch Place, Edinburgh EH8 9JT, United Kingdom. Email: sevimora@gmail.com

1 Introduction

In this paper we present a theory of human capital accumulation stressing the role of statistical discrimination benefiting those from privileged backgrounds, and the effects of having imperfect direct information on an individual's productive abilities.

In our model the income of an individual reflects the perception that society has on her productive abilities, summarizing all the information that the market has on the individual. It partly reflects what it is known about her abilities themselves, but also reflects any independent information available on the agent's background. This is because the market is aware that some backgrounds are more conducive to human capital accumulation. In equilibrium, richer parents tend to invest more on their children, consequently their children tend to be more productive. Thus, knowing the background of an individual is informative on her abilities, and this information will be used by the market for determining her income *even conditioning in the direct information available on her abilities*.

We present two main results. The first concerns the effects on inequality and mobility of having more information on the background and abilities of the individuals. It is not surprising that more accurate information on background translates into a larger degree of inequality and lowers intergenerational social mobility. This is because more information on background allows firms to positively discriminate in favor of those known to have a better education. Moreover, and perhaps more surprisingly, we show that the same effects (more inequality and less mobility) appear as a consequence of the market having better information on the agent's productive ability – giving more “*advantages*” to the rich (by the means of differentiating them more clearly from the poor and thus favoring them via positive statistical discrimination) has qualitatively the same effects on inequality and mobility than increasing the degree of “*meritocracy*” (in the sense of allowing firms to judge the abilities of workers with more accuracy independently of their backgrounds).

The intuition for this result is as follows. By definition in a “*meritocratic*” society the effect of parental income should be very small when conditioning on the productivity of the agent, but these two variables (background and productivity) are not independent. More information on the productive abilities of agents increases the differences in rewards that those agents receive as the market is more efficient at differentiating the very productive from the less so. These differences in income translate into differences in investment in the human capital of children, which results in larger differences in their productive abilities and a larger correlation between the income

of parents and their children. In equilibrium this gets amplified by a feed back mechanism, as larger inequality increases the value that the market assigns to any available information on an individual's talent.

Our second result concerns the effects of information (again, on background and on productivity) on the incentives of parents to invest in their children.

In a very meritocratic society (where firms will know the productivity of your children and pay them accordingly), the return to educational investment will be large, and thus your investment in your children will be larger than in a less meritocratic one.

In contrast, an increase in the information than the market has on the background of individuals does not have the same effect. As we will see, it helps firms better guess the productivity of agents (via statistical discrimination), but the effects on the incentives to educate children are minimal for the obvious reason that your income level does not depend on the education that your children have. Firms employing your children care more about what income you had, as it helps them guess the education that your children received, but you cannot change their beliefs about your children's background by educating them more or less.

Thus, in societies with better information on productive ability, we should expect more investment in education than in societies where "*advantages*" are more prevalent (i.e., where statistical discrimination based on background). But the former and the latter do not need to differ in the degree of inequality and mobility.

Thus, our contribution is to stress that a low level of intergenerational mobility and a high degree of inequality are perfectly compatible with a high degree of "*meritocracy*", with people being judged and rewarded according to their productive abilities. Moreover, two societies with the same degree of mobility and inequality may have arrived at such a situation from very different routes. One may be characterized by people being judged and rewarded by their abilities, while the other by the prevalence of advantages, the children of the rich being *presumed* to be more productive than the children of the poor. From the point of view of mobility and inequality both societies may look identical, but they differ radically in the effort they make in education. In the society where people is rewarded for their productive ability there are strong forces conducive to high effort in educational investment, as this investment is likely to be rewarded with future income. These forces are absent in a society where advantages prevail.

Our contribution is theoretical, but we perform an exercise that aims to be an example of the type of effects our mechanism could produce. We do the intellectual exercise of looking at

what the model implies for the degree of “*meritocracy*” in different societies. In particular, we are interested in knowing if the implied degree of “meritocracy” in the US is large relative to other countries.

When compared with other developed economies the US shows very large degree of inequality, a low degree of intergenerational mobility and very large investments in education.¹ Following the logic of the model, it seems as plausible that the US could have a large degree of inequality and low mobility precisely because it has a high degree of “*meritocracy*”. We show that that is indeed the manner in which the model reads the data.

This is, the model is able to replicate the data only if the degree of implied “*meritocracy*” in the US is much larger than in other OECD countries. If this were true, the US would still be a land of opportunity *provided that you have high productive ability*. This fulfillment of this restriction is what it would produce the dismal performance in mobility and inequality. In order to solve the model we do implausible assumptions on the extent of redistribution and public education that are bound to overestimate the weight of our mechanism. Thus, we are reluctant to interpret these results as evidence of the distribution of meritocracy across societies. It is nevertheless a nice illustration of the mechanism that we model.

The paper is organized as follows. Section 2 reviews the existing literature and discusses semantics. Section 3 describes and sets up the model. Section 4 will look at the agent’s problem and human capital investment. Section 5 will consider the firm’s problem and pricing the signals on human capital and background. It will also consider some comparative statics exercises on inequality and mobility. Section 6 will describe the equilibrium and comparative statics on educational investment. Section 7 shows the way in which the model filters existing data. Section 8 summarizes and concludes. All proofs are relegated to the appendix.

2 Semantics and Related Literature

We are by no means the first to think along these lines. The word “*meritocracy*” itself was coined (as recently as 1958) by sociologist Michael Young in his book “*The Rise of Meritocracy*”² in order to put forward some of the ideas that we model. Young’s book narrates an imaginary history of the UK up to 2034 (the book is supposedly written in 2033). In the meritocratic society

¹It is well known that intergenerational mobility correlates negatively with inequality. See Krueger (2012). This correlation has been documented across countries (Corak (2013)), across US commuting zones (Chetty et al. (2014)) and across Italian provinces (Güell et al. (2015)). The US performs particularly poorly on this relationship.

²Young (1958)

it describes, positions are allocated based solely on merit, not on birth, but the book describes a dystopia, not an utopia. The reasons are in essence the ones we describe: the distribution of talent is endogenous to the workings of society. The high reward for talent that meritocracy entails fosters inequality in the distribution of income. The education system then tailors to the rich, who can afford more and better education for their children. Thus, meritocracy fosters both inequality and persistence in income across generations. Young's book finishes with the death of the supposed author at the hands of an anti-meritocratic revolution when those left behind by history rebel against what they perceive as the unbearable inequalities that meritocracy has created.

Given that the word "*meritocracy*" has today a positive meaning, it is surprising to notice that not even 60 years ago it was coined as a warning. It was recognized that it would give rise to a feeling of justice and would foster short-run (perhaps even long-run) efficiency, but it was doomed to increase inequality to the point of generating social instability.

We do not want to make a big deal about semantics, but it is necessary to clarify that when we use the word "meritocracy" we will stick to its original meaning: the ability of society to recognize the productive ability of individuals and reward them accordingly. For us it is a statement about the availability of direct information on a worker's productive ability. We like to think of this information as the result of the education system signaling with more or less accuracy the abilities of individuals. This itself could be a consequence of a more or less widespread use of ability tests and the extension of systematic performance measures across the schooling system. Alternatively, it could be the result of stratifying the education in school-rankings, which select and signal their students abilities.³ In any case, for us this information is an exogenous parameter characterizing society, and we call it the degree of "*meritocracy*" of a society.

The important thing is the concept, though, not the word by which we name it. We recognize that other people may have different conceptions on what "*meritocracy*" means, and some people have strong opinions on this semantic issue. Clearly, many people associate the word with a reduction of the privileges that some members of society may enjoy. These privileges (which we might call "*cronyism*") mean that some people may be rewarded far in excess to their

³Young's book was a reaction to the extension of the tripartite educational system in England, Wales and Northern Ireland at the end of WWII. By this system children were examined at the age of 11 and selected into one of three school types based on their performance. At the top, Grammar schools would feed into university (themselves selecting into further categories according to rank and student quality). At the bottom, students would be trained in "practical skills". Young's book became paramount in the abolition of the system in the 70's, with the introduction of Comprehensive School System that, at least in theory, offers the same education to all.

contribution to society. Examples would be the corruption in the allocation of public positions via friendship or connections, or the impossibility of accessing higher education and positions of power and substance without the benefit of parental wealth and connections.

Notice that the reduction of those privileges is not what we mean by “*meritocracy*”. It could be another perfectly reasonable definition, but is not the one that we use, and it is not what Young meant when he *invented* the word. Since Becker (1957) we know that irrational discrimination has negative effects not only on those discriminated against, but also on the discriminator. Competition and market forces should work against the extent of those privileges: a firm that hires a person because he has an aristocratic name rather than talent, is a firm that is losing money and will be driven out of a competitive market. Thus, our approach is to model the advantages that the rich enjoy as a result of rational statistical discrimination of firms with limited information.⁴ These firms know that *in equilibrium* the children of the rich will have received a better education and are likely to be more productive than the children of the poor. Discrimination is, thus, the rational response to the available information, as in many statistical discrimination models such as Arrow (1973), Coate & Loury (1993), Norman (2003) and Moro & Norman (2004).

Inefficient capital markets are another hurdle that the children of the poor endure. The study of this in economics started with Galor & Zeira (1993) and has since created a large literature. Nevertheless few papers, if any, study the effects of this in connection with the degree of meritocracy. We do not take the route of exploiting credit market imperfections as generators of inequality (our agents cannot access capital markets, but as the expected return to education is homogeneous for all agents, they end up investing in their children’s education a fixed proportion of their income), but there is no reason why the effects of “*meritocracy*” would be different, in such a case, from the ones we find.

A final related stream in the literature refers to the study of inequality and intergenerational mobility pioneered by Becker & Tomes (1979). Several of those papers have dealt with understanding the negative relationship between inequality and mobility, a number of which hint at the process of human capital accumulation as the possible cause. Hassler et al. (2007), for instance, show that if the bulk of the difference across economies resides in the workings of the labor market (skill bias or institutions), then the correlation between inequality and mobility

⁴In any case, in appendix K we include “*cronyism*” (privileges not backed by reason) as an extension of our analysis, and we show that their reduction does indeed decrease the intergenerational persistence of income and (in most cases) the degree of inequality. Still, even in their presence the effects of improving the technology to determine one’s ability (“*meritocracy*” in our sense) are qualitatively the same as in our main analysis.

across economies would be positive. For the correlation to be negative (as the data shows it is) the main differences across economies needs to be focused in the education system and the process of education acquisition. Likewise, Solon (2004) shows that differences in the degree of progressivity of the educational system generate the negative correlation observed.

A paper to which we are particularly related is Abbott & Gallipoli (forthcoming). Like us, they present a model that aims to explain the geographical differences in Intergenerational Mobility. They present extend the Becker-Tomes model introducing a production sector in which workers human capital inputs are complements. This leads to a negative correlation between progressive public policy and IGE. Computing the model they show that geographic differences in skill complementarity directly account for roughly one fifth of cross-country variation in Intergenerational Mobility.

la

3 Model Description

We present an overlapping generations model with no capital markets. There are two different set of actors in the model: families and firms.

Agents work and receive a wage. The wage (the only state variable) is different for different agents, and it is a function of their observable characteristics. They care about their own consumption and the well being of their children (with a certain discount rate). They have no access to capital markets. The only decision they need to make is how much of their income to consume and how much to invest in the education of their children. The only way of moving resources forward in time is by investing in children's education. There is no way of bringing resources back from the future. This investment decision depends crucially on how much education will affect the income of the child, which it is an endogenous variable of the model but which each individual agent takes as given.

Firms produce output competitively using only labor. The productive ability of each worker (her "*human capital*") is an stochastic function that depends positively on the investment that her parents made on her. The production function is linear in the productive ability of workers: output equals human capital one to one. Firms, though, cannot observe the productivity of workers. They observe a set of characteristics of the workers, and need to infer the expected human capital of each worker as a function of these characteristics. Firms, thus, pay each worker her expected "human capital" and make zero profits on average.

The exact manner by which human capital affects income is endogenous to the model. Actually, our main contribution is to articulate the mechanisms by which this process takes place. If firms were able to observe human capital exactly, income would be exactly equal to human capital; but it is observed with noise. That is, the market has two pieces of information about an individual: a noisy signal on her ability and another on her background. Firms determine rationally how much to pay an individual as a function of her two signals.

The quality of each of these signals is exogenous. We say that a society is more meritocratic if firms have better and more accurate information about the human capital of the individual: if the signal on human capital is more precise. In addition to the signal on human capital, firms use the available information on an agent's background because they know that in equilibrium richer parents invest more in education than poorer ones. Thus, there is statistical discrimination favoring children from better backgrounds. We refer to a society with better information on an individual's background as being more "aristocratic" as the circumstances of birth play a larger role in income determination *even when conditioning on individual productive abilities*.

The core of the paper consists in solving a dynamic general equilibrium model with these three components (educational investment, meritocracy and inherited advantages), and to show the comparative statics of the endogenous variables to exogenous changes in the precision of the information available on the human capital of the agent (changes in "meritocracy") and the precision of the information on parental income (changes in "aristocracy") in the steady state. We show that better information (irrespective of whether it is on merit or background) always has the effect of increasing inequality and decreasing mobility. Meritocracy does not differ from aristocracy in these dimensions.

In equilibrium parents use the resultant relationship between investment and income to infer how much they should invest in their children and how much to consume. Thus, closing the model.

We proceed by first exploring in section 4 the decisions of families given a certain stochastic map from investment to income. We then look at how much the market values the different characteristics of workers in order to determine wages (in section 5), and we put both things together in section 6.

4 Human Capital Investment

We start by setting up the problem of families, which have to choose how much to consume and how much to invest in their children's education.

At each point in time the family consists of a father and a son. After each period the father dies and the son becomes a father.

There are no capital markets, so the resources of the family are exclusively the income generated by the father in the labor market. There is no way of borrowing in exchange for future income, and there is no way of leaving resources for future generations except by investing in one's children. Thus, the only state variable is the income of the family.

The father needs to allocate his resources between current consumption and investment in his child's human capital. This investment, *and the income that the father earned*, will determine the observable characteristics of the child.

Firms produce with human capital only, but they do not observe the human capital of workers, only some informative characteristics about it. Consequently, their role is to decide how to translate those observable characteristics into expected human capital and income.

Thus, the problem which parents confront is:

$$W(Y_t^i) = \max_{X_t^i} \left\{ \ln C_t^i + \frac{1}{1+\delta} EW(Y_{t+1}^i | Y_t^i, X_t^i) \right\} \quad (1)$$

s.t.

$$C_t^i = Y_t^i - X_t^i; \quad X_t^i \geq 0 \quad (2)$$

$$Y_{t+1}^i \sim F(H_{t+1}^i, Y_t^i) \quad (3)$$

$$H_{t+1}^i \sim G(X_t^i) \quad (4)$$

where Y_t^i is the (lifetime) income of family i at generation t , C_t^i is their consumption, X_t^i is their investment in the human capital of their offspring, and H_{t+1}^i is the human capital that the offspring actually achieves. Equation 2 is the budget constraint that parents face given the absence of capital markets. Equation 4 states that human capital is a stochastic function of the investment, without a role for other forms of inheritance (genetic or otherwise).

The investment problem is quite standard. As in Becker & Tomes (1979, 1986), agents decide how much to invest in the human capital of their children. We have assumed that children's innate abilities are stochastic, and parents invest without knowing the realization of

their children's abilities. That is, there is a random component to the determination of human capital and, when making the investment decision, parents do not know how much human capital their investment will translate into.

The novelty of our theory refers to equation 3. In the next sections we will develop a theory such that in equilibrium the income of an individual is a stochastic function of her human capital and her parental income, which has effects *even controlling for the human capital of the child*. Moreover, the magnitude of the effect of both human capital and parental income is a function of the degree of inequality of the distribution of income, which is itself an endogenous object of the model.

Equations 3 and 4 together imply a relationship between the income of a child and the income and investment made by their parent. We will begin, in this section, by taking as given the following functional form of this relationship:

$$Y_{t+1}^i = e^{\gamma_0} (Y_t^i)^{\gamma_1} (X_t^i)^{\gamma_2} e^{\varepsilon_{t+1}^i} \quad (5)$$

where γ_0 , γ_1 and γ_2 are endogenous parameters of this *income determination function* and ε_{t+1}^i is a stochastic error term with a certain mean $\bar{\varepsilon}$ and variance V_{ε} . At this stage we are guessing that this is the correct functional form, but in subsequent sections we will show that this is, in fact, the correct form of the income determination function and solve for the equilibrium values of γ_0 , γ_1 and γ_2 , and the stochastic structure of the random variable ε_{t+1}^i .

The last ingredient which we need to solve the parent's maximisation problem and find the equation for the accumulation of human capital is a functional form for equation 4. We assume that it is:

$$H_{t+1}^i = Z (X_t^i)^{\alpha} e^{\tilde{\omega}_{t+1}^i}; \quad \tilde{\omega}_{t+1}^i \sim N\left(-\frac{V_{\omega}}{2}, V_{\omega}\right) \quad (6)$$

where Z is a constant akin to total factor productivity, and $\tilde{\omega}_{t+1}^i$ is an iid shock normally distributed with a mean such that $E\left(e^{\tilde{\omega}_{t+1}^i}\right) = 1$. We assume that $\alpha \in (0, 1)$. Notice that an agent may have a very large human capital either because her parents invested a large amount in her (large X_t^i), or because she is very good at creating human capital (high $\tilde{\omega}_{t+1}^i$). This second route to excellence may be thought as luck or talent. The only imposition that we put is that it is not inheritable. It is independent of who the parent is.

The following result⁵ characterizes the optimal investment decisions of dynasties:

⁵Proof in appendix A

Result 1. *The solution of the maximization problem in equation 1 requires that investment in education is a fixed proportion of the individual's income: $X_t^i = \lambda Y_t^i$, with:*

$$\lambda = \frac{\gamma_2}{1 + \delta - \gamma_1} \quad (7)$$

The value function of agents is $W(Y) = A + B \ln Y$, with

$$A = \frac{\bar{\varepsilon} + \gamma_0 + \ln[(1 + \delta) - (\gamma_1 + \gamma_2)]^{[(1+\delta)-(\gamma_1+\gamma_2)]} + \ln \frac{\gamma_2^{\gamma_2}}{[(1+\delta)-\gamma_1]^{[(1+\delta)-\gamma_1]}}}{[(1 + \delta) - (\gamma_1 + \gamma_2)] \frac{\delta}{1+\delta}} \quad (8)$$

$$B = \frac{(1 + \delta)}{(1 + \delta) - (\gamma_1 + \gamma_2)} \quad (9)$$

Since parents invest a fixed fraction of their income, λ , in their child's education, it follows that, in absolute terms, the rich invest more heavily than the poor. Since this is a known feature of the economy, the children of the rich will be perceived to have more human capital than the children of the poor when human capital is unobservable. This leads to statistical discrimination in favour of the children of the rich.

Notice also that given investment is a fixed proportion of income (result 1) and that human capital has a constant elasticity to investment (equation 6), it follows that:

Result 2. *Log human capital is a linear function of log parental income. The elasticity is an exogenous parameter α , while the constant depends on the investment rate λ :*

$$h_{t+1}^i = \ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} + \alpha y_t^i + \omega_{t+1}^i \quad (10)$$

where $\omega_{t+1}^i \sim N(0, V_\omega)$ is iid noise and h_{t+1}^i and y_t^i represent the log of the child's human capital and the log of parental income respectively.

This is the equation for the accumulation of human capital. Human capital depends positively on parental income and the fraction of income which parents invest in their children. This equation is known to firms but the level of human capital of a particular worker is not, nor is the parental income of any particular worker (at least not perfectly).

5 Wage Determination

We now turn to how the wage of an individual is determined.

Agents can only produce output when working within firms. These firms are competitive in the labor and product markets. The only input is human capital, which produces output in a one to one basis.⁶

Our most salient assumption is that firms cannot observe the productivity of the workers. They observe it with noise along with other characteristics of the agent. We call Ω_{t+1}^i to the set of information that the market has available on agent i . We discuss the composition of this set below. So far we just want to notice that it is public information referring to worker i . Thus, all firms have the same beliefs on any particular worker. All heterogeneity resides in how different workers are, not in how they are perceived by society.

Competition insures that firms will make zero expected profits in equilibrium and that the wage of a worker with observable characteristics Ω_{t+1}^i equals the conditional expectation of her productivity:

$$Y_{t+1}^i = E [\exp \{h_{t+1}^i\} | \Omega_{t+1}^i] \quad (11)$$

5.1 Information Available to Firms.

The set of information available to firms when determining the wage of worker i is composed of 4 elements, two specific to the agent, and two more summarizing the state of the economy at time t :

$$\Omega_{t+1}^i = \{a_{t+1}^i, m_{t+1}^i, \mu_{y_t}, V_{y_t}\} \quad (12)$$

$a_{t+1}^i = y_t^i + \varepsilon_{t+1}^{ia}$ is a public signal on the parental income of agent i , where $\varepsilon_{t+1}^{ia} \sim N(0, V_a)$ is iid noise.

$m_{t+1}^i = h_{t+1}^i + \varepsilon_{t+1}^{im}$ is a public signal on the human capital of agent i , where $\varepsilon_{t+1}^{im} \sim N(0, V_m)$ is iid noise.

μ_{y_t} and V_{y_t} are the mean and variance of the distribution of log-income in the parents' generation. respectively.

In addition, all firms know how human capital relates to parental income (equation 10).

⁶ This view of production omits a role for misallocation in the sense that the productivity of a worker is unaffected by the degree of error in firms' beliefs about their human capital. This is done for simplicity. Presumably though, if firms' beliefs about each worker's human capital were very different to the reality, agents would be assigned to the wrong task and would not realise their full productive potential. In appendix M we consider a production function which includes an allocative cost. This would seem to be more realistic, but the added complexity does not change any of the results presented in the body of the paper.

The two most important parameters for our model are V_a and V_m . They determine the amount of information that firms receive on the background of the agents (V_a) and their productive ability itself (V_m).

a_{t+1}^i is the signal on background, and V_a determines how much information this signal contains. If V_a is very low, firms know with almost certainty the income of the parents, and thus how much was invested in the education of the agent. Notice that if V_a is larger, the degree of uncertainty that they have on the amount of education that the agent received depends on the income distribution of the parents: the more inequality among the parents, the more uncertain are the priors and, consequently, the posteriors. This fact will have important consequences.

m_{t+1}^i is the signal on the productivity of the agent, and V_m determines how much *direct* information the market receives on it – how much firms know about the productivity of a worker *independently* of her background and circumstances. It is the amount of “hard” information available, unaffected by prejudice of any sort.

If V_m is very low, the firm knows almost exactly what the productivity of the worker is, and she will be paid accordingly *independently of her background*. The firm will pay the productivity, and it does not matter if she has a high h because the investment of the parents was high or because she was particularly lucky in the drawing of $\tilde{\omega}$ in equation (6).

If V_m is larger, the firm cares more about the background, as it has predictive value on h . In this sense we call “*meritocracy*” the precision of the signal m_t^i . If it is high, only the productive ability of agents matters for their wage. If it is low, firms are interested in the background of the agent, which will affect her wage.

Notice that our model is a model of the first years of the productive life of individuals. We are implicitly assuming that during those years there are critical events in the personal history of agents that determine their income thereafter. It is not an outlandish assumption, as there is a substantial body of evidence suggesting that the first jobs an individual takes have long lasting effects on her working life.

Starting with ? a literature has studied the mechanisms of human capital accumulation and promotion, indicating that tracks are critical in the wage process. Thus, ? and ? (see also ?) show that initial tracks do matter. There are different tracks of jobs. If you start in a low wage track you tend to promote within the track, not towards high wage tracks, the income differentials across workers with the same tenure being somewhat persistent. Likewise, promotions have persistent effects: if you get promoted first, your wage tends to be persistently

higher.

5.2 Updating Beliefs

Firms process Ω_{t+1}^i in order to generate rational beliefs on agent's i human capital, but before knowing the specifics of any agent, firms have a prior on the distribution of human capital. This prior is generated by (i) the rule with which families invest in their children (equation 10), including the distribution of “luck” in learning (ω) and (ii) the distribution of income among the parents. Clearly, if y_t is normally distributed, the prior on h_{t+1}^i will also be normally distributed, and the larger the variance of parental income, the less precise the prior of firms on the human capital of their children.

Once they receive the signals a_{t+1}^i and m_{t+1}^i on a specific agent i , the market will update its prior on this agent using Bayes rule. We summarize the outcome of this process in the following result:⁷

Result 3. *Given the process of human capital accumulation in equation 10 and the information set given by equation 12, the posterior of the log of human capital is given by:*

$$h_{t+1}^i | \Omega_{t+1}^i \sim N \left(\mu_{h_{t+1}^i | \Omega_{t+1}^i}, V_{h_{t+1}^i | \Omega_{t+1}^i} \right) \quad (13)$$

with

$$\mu_{h_{t+1}^i | \Omega_{t+1}^i} = \beta_{m_{t+1}} m_{t+1}^i + (1 - \beta_{m_{t+1}}) \left[\ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} + \alpha \{ \beta_{a_{t+1}} a_{t+1}^i + (1 - \beta_{a_{t+1}}) \mu_{y_t} \} \right] \quad (14)$$

$$V_{h_{t+1}^i | \Omega_{t+1}^i} = \beta_{m_{t+1}} V_m \quad (15)$$

$$\beta_{a_{t+1}} = \frac{V_{y_t}}{V_{y_t} + V_a} \quad (16)$$

$$\beta_{m_{t+1}} = \frac{\alpha^2 \beta_{a_{t+1}} V_a + V_\omega}{\alpha^2 \beta_{a_{t+1}} V_a + V_\omega + V_m} \quad (17)$$

To understand the result it is easier to focus first on expression 16. $\beta_{a_{t+1}}$ is the weight of the signal on parental income when the agent updates the beliefs given by the prior (determined by μ_{y_t} and V_{y_t}) with the information received in the signal a_t^i . The mean of the posterior on parental income is: $\{ \beta_{a_{t+1}} a_{t+1}^i + (1 - \beta_{a_{t+1}}) \mu_{y_t} \}$.

The expected value of the productivity of the agent (equation 14) is a convex combination of

⁷Proof in appendix C

the realization of the direct signal on productivity (m_t^i) and the expected log human capital for an individual with parental income equal to $\{\beta_{a_{t+1}} a_{t+1}^i + (1 - \beta_{a_{t+1}}) \mu_{y_t}\}$, the posterior belief on parental income.

$\beta_{m_{t+1}}$ is the weight given to the signal m_t^i (in equation 17). The expression $[\alpha^2 \beta_{a_{t+1}} V_a + V_\omega]$ is the variance of h conditional on a_{t+1}^i . Notice that the variance of the *posterior* of parental income is $\beta_{a_{t+1}} V_a = \frac{V_{y_t} V_a}{V_{y_t} + V_a}$, which is smaller than the unconditional variance V_{y_t} .

If the signal on “merit” is very accurate ($V_m \rightarrow 0$), then $\beta_m \rightarrow 1$ and the belief is centered in the real productivity of the agent with zero variance (equation 14). The larger is the variance of the merit signal, the bigger is the weight of the “prior”. That is, the bigger is the weight of the *average* human capital that you would expect from an individual with the perceived background of the person we are evaluating. Thus, V_m is indeed a measure of how much background matters.

We now map these beliefs on h_{t+1}^i into wages. Notice that firms care about expected productivity, not its logarithm (which is normally distributed). We take care of this in the following result:⁸

Result 4. *The logarithm of the wage individual i with signals a_{t+1}^i and m_{t+1}^i is:*

$$y_{t+1}^i = (1 - \beta_{m_{t+1}}) \left[\ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} + \alpha (1 - \beta_{a_{t+1}}) \mu_{y_t} \right] + \frac{\beta_{m_{t+1}} V_m}{2} + \alpha \beta_{a_{t+1}} (1 - \beta_{m_{t+1}}) a_{t+1}^i + \beta_{m_{t+1}} m_{t+1}^i \quad (18)$$

Thus, log income is a linear function with three components:

- A constant that depends on the technology of education acquisition, how much the average person invests in education (remember that λ is the share of income invested and μ_{y_t} the average of the log of income) and a factor that depends on the ex-post variance of the distribution of human capital.⁹

- A term that depends on the information that society has on her specific background (a_{t+1}^i).

We will call $\hat{\beta}_{a_{t+1}} = \alpha \beta_{a_{t+1}} (1 - \beta_{m_{t+1}})$ the “value” that society gives to background.

- A term that depends on the direct (“objective”) information that society has on her productive ability (m_{t+1}^i).

⁸Proof in appendix D

⁹ This is a property of the variance of the log normal: the larger the variance of the logarithm, the exponentially larger the expectation is. It is the variance of the *posterior* of the distribution of log human capital that is relevant here, not its unconditional distribution. Notice that the variance of the posterior is $\beta_{m_{t+1}} V_m$. A property of the normal distribution is that this variance does not depend on the realization of any of the signals. This would not be true with an arbitrary distribution function.

Notice that if the signal on human capital is very accurate ($V_m \rightarrow 0$), then $\beta_{m_{t+1}}$ equals one, and income equals human capital exactly, independently of the distribution within the population, and independently of the background of the individuals. But, in any other case preconceptions matter for wage determination.

Large values of $\beta_{m_{t+1}}$ imply that society rewards highly the observed objective measures of productive ability. Large values of $\hat{\beta}_{a_{t+1}}$, on the other hand, imply that society rewards highly the perception of a favored upbringing *even controlling* for the objective measures of ability. Of course, both are endogenous. They depend not only on the precision of both signals, but also on the distribution of income (and thus investment).

We turn to the determination of these weights next, but before it is convenient to rewrite equation (18) as the stochastic income process, its Becker-Tomes representation. Noticing that a_{t+1}^i and m_{t+1}^i are both stochastic functions of y_t^i we can write the law of motion of the log of income:

$$\begin{aligned} y_{t+1}^i = & \ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} + \alpha (1 - \beta_{a_{t+1}}) (1 - \beta_{m_{t+1}}) \mu_{y_t} + \frac{\beta_{m_{t+1}} V_m}{2} \\ & + \alpha [\beta_{a_{t+1}} (1 - \beta_{m_{t+1}}) + \beta_{m_{t+1}}] y_t^i + \alpha \beta_a (1 - \beta_{m_{t+1}}) \varepsilon_{t+1}^{ia} + \beta_{m_{t+1}} (\varepsilon_{t+1}^{im} + \omega_{t+1}^i) \end{aligned} \quad (19)$$

where the intergenerational income elasticity is then $\rho_{t+1} = \alpha [\beta_{a_{t+1}} (1 - \beta_{m_{t+1}}) + \beta_{m_{t+1}}]$.

From here, after some algebra, we can determine the law of motion of the variance of log income.

$$\begin{aligned} V_{y_{t+1}} &= \alpha^2 [\beta_{a_{t+1}} (1 - \beta_{m_{t+1}}) + \beta_{m_{t+1}}] V_{y_t} + \beta_{m_{t+1}} V_\omega \\ \beta_{a_{t+1}} &= \frac{V_{y_t}}{V_{y_t} + V_a}; \quad \beta_{m_{t+1}} = \frac{\alpha^2 \beta_{a_{t+1}} V_a + V_\omega}{\alpha^2 \beta_{a_{t+1}} V_a + V_\omega + V_m} \end{aligned} \quad (20)$$

This system of differential equations completely describes the dynamics of the model. Notice that V_{y_t} is not stochastic, given initial conditions, as there is no aggregate risk.

Before describing the properties of the steady state of the system we find it useful to make two observations.

1. *In our model there is no role for self-fulfilling expectations. Expectations never drive the dynamics.*

This sets our model apart from the bulk of the literature on race- or gender-based statistical discrimination. In that literature the observed characteristic has no intrinsic value,

as there is nothing inherently good or bad in belonging to a given race or gender. Nevertheless, the fact that a characteristic is observable may make it acquire informative value: if everybody *expects* people of a certain race or gender to act in a certain way (investing little in education, for instance), it might be optimal to privately act according to such an expectation (you will invest little in education if people expect you to have little education and it is your race, not your education, that is observed). Multiple equilibria are thus natural in such an environment, as the informative content of an observable characteristic depends on how people are expected to act and those expectations feed back into actions. Other expectations would lead to other actions, and those actions could feed those different expectations. Most obviously: the race or gender that is expected to have lower education could be changed and nothing of substance would be altered.

Our model is very different in this respect because here it is objectively good to be the son of a rich father, and it is objectively good to have high productive ability. There is no way of sustaining an equilibrium where low income parents invest in their children as much as rich parents do, as marginal utility of consumption of poor parents is larger.

It is easy to see that the market will always discriminate against the children from deprived backgrounds. If it did not, the rich would still invest more in their children than the poor, so it would be irrational not to discriminate.

Likewise, an individual with a large value of m_t is always going to be paid more than another that differs only in having a lower value of m_t . More productivity is more productivity, and there is no other conceivable way of interpreting it.

2. *There is a positive feed-back mechanism by which inequality fosters further inequality.*

On one hand, the more inequality there is (large V_{y_t}), the more heterogeneous agents are in their productive ability, and the more uncertainty firms face. Consequently, the market rewards indications of productive ability, both the direct ones (m_t^i), and the suggestive ones via parental investment (a_t^i). $\beta_{a_{t+1}}$ and $\beta_{m_{t+1}}$ are increasing in V_{y_t} .

On the other hand, the larger the value given to the signals ($\beta_{m_{t+1}}$ and $\beta_{a_{t+1}}$), the larger the amount of inequality the next generation will endure, as the differences among agents are more salient.

Thus, the more that society values the observable differences between agents, the more inequality it creates, and because of that, the more that it values the observable differences

between agents. Or, looking at it from the other side, inequality fosters further inequality via increasing the way in which society differentiates among its members.

This positive feed-back mechanism is ingrained in the process of statistical discrimination. It implies a multiplier effect: the effect of any exogenous change of parameters leading to an increase in inequality will be amplified.¹⁰

Thus, our model does not allow for the existence of different sets of strategies and beliefs which could be mutually compatible for a single set of state variables: equilibrium is unique. Nevertheless, the existence of the positive feed-back mechanism generates the possibility of multiple path dependent steady states. This happens if different values of the state variables lead you to different steady states in the long run, but there is no possibility of moving from one of those steady states once the economy has settled in it.

In our model multiple steady states arise if the elasticity of human capital to investment is large enough. If $\alpha \geq 1$ there are three steady states, of which two are stable. Of these one has infinite variance, with ρ being not smaller than one and huge responses to a and m . The other stable equilibria is qualitatively identical to the one that we present for $\alpha < 1$. We prefer to restrict the parameter space to insure that this possibility does not arise. There are three reasons for this: (1) because we do not know how to interpret an infinite variance, (2) because the restriction necessary ($\alpha < 1$) is eminently reasonable, and (3) because the steady state that results is equivalent with the reasonable (i.e. finite variance) stable steady state if the restriction does not apply. From now on we always assume that $\alpha < 1$.

We characterize the solution of this system of differential equations in the next result.¹¹

Result 5. *If $\alpha < 1$ equations (16), (17) and (20) define a system of differential equations with a unique steady state which is globally stable. In the steady state log income is normally distributed with variance being characterized by the (unique) solution of the following system of equations:*

$$V_y = \frac{\beta_m V_\omega}{1 - \alpha^2 [\beta_a (1 - \beta_m) + \beta_m]} \quad (21)$$

$$\beta_a = \frac{V_y}{V_y + V_a} \quad (22)$$

$$\beta_m = \frac{\alpha^2 \beta_a V_a + V_\omega}{\alpha^2 \beta_a V_a + V_\omega + V_m} \quad (23)$$

Unfortunately, the explicit solution to this system of equations is extremely cumbersome

¹⁰Proof in appendix E

¹¹Proof in appendix F.

and uninformative. We can instead solve analytically the comparative statics on the relevant exogenous variables. That is, we look at how the steady state values of the endogenous variables (V_y , β_a and β_m) move as a consequence of exogenous changes in V_a and V_m .

5.3 Comparative Statics of V_a .

A decrease of V_a means that the market will have more accurate information on the background of agents. This is not good news for equality. Being better able to differentiate who received more education, the market will be more inclined to discriminate in their favor, providing greater advantages to those from affluent backgrounds. But how much the market chooses to discriminate is a function of the degree of inequality in the economy, and this is a state variable whose path is determined endogenously.

In appendix G we prove the following results, which characterize the effects of V_a

Result 6. *An increase in the accuracy of the signal on background (a decrease of V_a) results, in steady state, in more inequality, greater persistence of income across generations, more discrimination based on perceptions of the background of an agent, and a smaller elasticity of income to the signal on ability:*

$$\frac{dV_y}{dV_a} < 0; \quad \frac{d\rho}{dV_a} < 0; \quad \frac{d\beta_a}{dV_a} < 0; \quad \frac{d\hat{\beta}_a}{dV_a} < 0; \quad \frac{d\beta_m}{dV_a} > 0 \quad (24)$$

More accurate information on the background of an individual increases the attention that firms give to this signal, increasing the persistence of income across generations and its variance across individuals.

Notice that this result is far from obvious: Decreasing the amount of noise in the economy (i.e., increasing the information that agents have) *increases* the dispersion of incomes. You reduce noise, but as result you increase dispersion.

The reason lies in the increase in the persistency of the income process. Better information on the background of the individuals allows firms to discriminate more accurately between agents, directly favoring those from better backgrounds. A more persistent income process is bound to have a larger unconditional variance. Thus, inequality increases, which itself increases even further the value of the information on background via the positive feed-back mechanism discussed above.

Notice also that the effect on the human capital signal (β_m) is the opposite. Better infor-

mation on background results in a *lower* elasticity of income to the ability of individuals. This might look surprising, given that inequality is larger for lower values of V_a . More inequality implies that firms are less aware on the abilities of any specific worker, and thus, one could have expected that firms would give more attention to the meritocratic signal of human capital (m_t^i). They do not, and the reason is that, albeit the unconditional variance of income increases, the variance of log income *conditional on the signal a_t^i decreases*. Thus, the dispersion of human capital *conditional on a_t^i* decreases, and there is less of a demand for the meritocratic signal. There is a crowding-out effect, background replacing merit in the determination of one's income and, consequently, a profoundly antipathetic decrease of intergenerational mobility.

5.4 An Exogenous Increase in the Degree of Meritocracy

Next we want to consider the effect of exogenously reducing V_m . A reduction in V_m , all else equal, improves the quality of the human capital signal, providing greater advantages to those with greater human capital. With all the caveats discussed in section 2, we consider it reasonable to say that a society with a lower value of V_m is more “*meritocratic*”.

The first thing to notice is that the same feedback mechanism that is very obvious for information on background, acts in the same manner for the information on merit. The easiest way to see it is by assuming that there is no signal on background ($V_a \rightarrow \infty$). In such a case the law of motion (20) becomes:¹²

$$V_{y_{t+1}} = \alpha^2 \beta_{m_{t+1}} V_{y_t} + \beta_{m_{t+1}} V_\omega; \quad \beta_{m_{t+1}} = \frac{\alpha^2 V_{y_t} + V_\omega}{\alpha^2 V_{y_t} + V_\omega + V_m} \quad (25)$$

More information on merit (lower V_m) makes the market increase the reliance on such information, increasing $\beta_{m_{t+1}}$. This necessarily increases the spread of incomes, as the differences between agents become more salient. Finally, the increase of inequality makes firms of the following generation less sure of the human capital of their workers, increasing the value that they assign to any information on merit, increasing β_m further...

¹²Notice that as V_a approaches infinity, then $\beta_{a_{t+1}} \rightarrow 0$ and $\beta_{a_{t+1}} V_a \rightarrow V_{y_t}$

It is easy to see that the unique steady state of (25) is the unique solution to:

$$V_y = \frac{\beta_m V_\omega}{1 - \alpha^2 \beta_m} \quad (26)$$

$$\beta_m = \frac{\alpha^2 V_y + V_\omega}{\alpha^2 V_y + V_\omega + V_m} \quad (27)$$

and that the steady state values of V_h , V_y and β_m are all decreasing in V_m . The intergenerational income elasticity (now equal to $\alpha\beta_m$) increases and V_m is reduced.

Contrary to what could be thought, meritocracy does not increase the degree of intergenerational mobility. It decreases it. More information on people's ability is bound to decrease intergenerational mobility because ability and background are correlated and, by increasing income dispersion, meritocracy *increases the value of any existing information on people's ability*. The children of the rich, being on average more productive than the children of the poor, benefit from the increase in the accuracy of information on merit, leading to more persistent income shocks and further inequality. Meritocracy has very much the same effects as an increase in the information available on background.

It is now easier to consider the effect of an exogenous improvement in the quality of the human capital signal when both signals are available and useful to the firm (i.e. V_a finite). We do so in the following result.¹³

Result 7. *An increase in the accuracy of the signal on ability (a decrease of V_m) results, in steady state, in more inequality, greater persistence of income across generations, a larger elasticity of income to the signal on ability and more weight given to the signal on background when evaluating an agent's parental income (which is what β_a measures):*

$$\frac{dV_y}{dV_m} < 0; \quad \frac{d\rho}{dV_m} < 0; \quad \frac{d\beta_m}{dV_m} < 0; \quad \frac{d\beta_a}{dV_m} < 0 \quad (28)$$

Moreover, given a set of values for $\alpha \in (0, 1)$ and $V_a \in R^+$ ($V_a < \infty$), there exists a variance of the error in the signal on ability, $\hat{V}_m$, such that $(0 < \hat{V}_m < \infty)$

$$\text{If } V_m < \hat{V}_m, \text{ then } \frac{d\hat{\beta}_a}{dV_m} > 0 \quad (29)$$

$$\text{If } V_m > \hat{V}_m, \text{ then } \frac{d\hat{\beta}_a}{dV_m} < 0 \quad (30)$$

¹³Proof in appendix H

The value of $\hat{\beta}_a$, the weight given to the signal on background when evaluating an agent's human capital, is maximal if $V_m = \hat{V}_m$.

If society is better endowed to judge its members merit, it is doomed to increase the dispersion of their incomes (paying more to those judged to be better). This increased dispersion has effects on both the value assigned to merit, β_m , and the value assigned to “advantages”, $\hat{\beta}_a$.

First it increases the dispersion of the abilities of the children, thus feeding back into increased underlying uncertainty and value of the signal on human capital in the following period. Thus, not surprisingly, better information on human capital increases the market value of that signal.

The effects on the weight given to background when determining income ($\hat{\beta}_a$) are more complicated. First of all, there is a “crowding-out effect” in the opposite direction to that in result 6. Better information on human capital makes you place less weight on background as merit replaces inherited advantages in the determination of human capital. This is clear from the fact that β_m enters negatively in $\hat{\beta}_a = \alpha\beta_a(1 - \beta_m)$. However, there is an effect on the other direction too: as income variance increases, the signal on background becomes more valuable in judging parental income. This increase in β_a acts in the opposite direction to the increase in β_m . The net effect on $\hat{\beta}_a$ depends on the relative size of the effects on β_a and β_m .

We can understand the net effect by doing the following mental exercise. Imagine that V_m were very low (and thus, firms have good information on ability). In that case β_m would be very large (close to one), and the effect of the increase of β_a would be very small ($\frac{\partial \hat{\beta}_a}{\partial \beta_a} = \alpha(1 - \beta_m)$). The net effect of a decrease in V_m would be an increase in the market value of the human capital signal, but a decrease in the value of the parental income one. There would be a trade-off between meritocracy and advantages as the crowding-out effect dominates.

Now imagine the polar opposite case where V_m is very large. In such a case β_m would be close to 0 and background information would play the dominant role in the determination of human capital. Any increase in the quality of information on ability would increase the market value of both signals because, in this case, the effect of an increase in β_a on $\hat{\beta}_a$ is relatively large.

In any case, notice that the degree of intergenerational mobility *always* decreases whenever the society becomes more meritocratic as a consequence of a decrease of V_m . This occurs both when there is a trade-off (and advantages become less important) or when there is not. This is a consequence of inheritance. The talented become richer, and thus incomes are bound to get more dispersed. This increased dispersion of incomes is going to be translated into a further dispersion of abilities as the rich invest more in the children. Abilities then, being better evaluated, translate

into more income for the children of the rich *even if it is perfectly possible that firms care less about the background of agents*.

This is one of the main insights of our paper. Meritocracy in and of itself is not going to increase intergenerational mobility or decrease the prevalence of inheritance. And this is bound to happen even if an increase in meritocracy does produce a decrease in the advantages associated with being from a good background, which is by no means a foregone conclusion.

This is not to say that meritocracy is a bad thing. In the next section we show that it increases the share of income invested in human capital. The significance of our result is to notice that there are several roads that lead to countries having low intergenerational mobility and high inequality: one is the “aristocratic” route, where the children of the rich benefit from positive statistical discrimination as the rich are known to invest more heavily in their children’s education than the poor; but a very different road leading to the same mobility and inequality is the meritocratic one. If the aspects of reality that one focuses on are limited to the degree of mobility and inequality, meritocracy and aristocracy are almost equivalent. Both of them produce a negative correlation between inequality and mobility, reproducing the observed data. Thus, just looking at the data we cannot say if a society is more or less meritocratic. For doing so we need to find a variable that reacts differently to advantages and to meritocracy. We do so in the next section.

6 Equilibrium

The log-normal structure of the model has allowed us to do the rather unusual exercise of solving for the comparative statics in steady state of a set of the endogenous variables (V_y , β_a and β_m) without solving the complete model. This facilitates our analysis enormously by making the pricing decisions independent of the share of investment in education (λ), insofar as all parents invest the same share of their income, which we saw in section 4 was the case. Now we solve for the complete equilibrium, including investment in human capital.

An equilibrium consists of (i) a rule for parent’s investment behavior and (ii) an income determination process such that:

1. The investment behavior of parents is optimal given the income determination process.
2. The income determination process is such that the wage of each worker equals her expected productivity given the information available on the worker (Ω_t^i) and the investment

behavior of parents.

In section 4 we saw that if parents believe that the income of their children is going to be determined by $Y_{t+1}^i = e^{\gamma_0} (Y_t^i)^{\gamma_1} (X_t^i)^{\gamma_2} e^{\varepsilon_{t+1}^i}$ (for some values of γ_1, γ_2 and stochastic process ε_{t+1}^i), then they choose optimally to invest a fixed proportion of their income $\lambda = \frac{\gamma_2}{1+\delta-\gamma_1}$.

In section 5 we have seen what the stochastic process of income is if parents invest a fixed proportion λ of their income in their children.

Thus, in order to prove that we have located an equilibrium it remains to be shown that there exists values of $\gamma_0, \gamma_1, \gamma_2$ and a well defined stochastic process ε_{t+1}^i such that equation (5) is a representation of the stochastic process of income defined by (18) in result 4 for the values of V_y, β_a and β_m obtained in result 5. We do so in the following result.¹⁴

Result 8. Equilibrium. *The equilibrium stochastic process of income as a function of parental income and investment is $Y_{t+1}^i = e^{\gamma_0} (Y_t^i)^{\gamma_1} (X_t^i)^{\gamma_2} e^{\varepsilon_{t+1}^i}$ with:*

$$\gamma_0 = \ln Z + \alpha(1 - \beta_m) [(1 - \beta_a) \mu_y + \ln \lambda] + \frac{1}{2} [\beta_m V_m - V_\omega] \quad (31)$$

$$\gamma_1 = \alpha \beta_a (1 - \beta_m) \quad (32)$$

$$\gamma_2 = \alpha \beta_m \quad (33)$$

$$\varepsilon_{t+1}^i = \alpha \beta_a (1 - \beta_m) \varepsilon_{t+1}^{ai} + \beta_m (\omega_{t+1}^i + \varepsilon_{t+1}^{mi}) \quad (34)$$

and, consequently, the equilibrium share of income invested in children's education is:

$$\lambda = \frac{\alpha \beta_m}{1 + \delta - \alpha \beta_a (1 - \beta_m)} \quad (35)$$

The elasticities of income with respect to parental income and investment are $\alpha \beta_a (1 - \beta_m)$ and $\alpha \beta_m$ respectively. The portion of talent which is not derived from parental income (ω_t^i) plays a larger role in the determination of income when β_m is higher, implying it has a more substantial impact in more meritocratic societies.

From results 5 and 8 we derive the following corollary, which describes the equilibrium.

Result 9. *If $\alpha < 1$ there exists an unique steady state. In steady state the values of V_y, β_a and β_m are the unique solution to:*

$$V_y = \frac{\beta_m V_\omega}{1 - \alpha^2 [\beta_a (1 - \beta_m) + \beta_m]}; \quad \beta_a = \frac{V_y}{V_y + V_a}; \quad \beta_m = \frac{\alpha^2 \beta_a V_a + V_\omega}{\alpha^2 \beta_a V_a + V_\omega + V_m}$$

¹⁴Proof in appendix I.

$$\rho = \alpha [\beta_a (1 - \beta_m) + \beta_m]$$

and the share of income invested in human capital is in all families identical, and equals

$$\lambda = \frac{\alpha \beta_m}{1 + \delta - \alpha \beta_a (1 - \beta_m)}$$

Given that we have already seen the comparative statics of V_y , ρ , β_m and $\hat{\beta}_a$ with respect to V_m and V_a , all what it rests to show is how the investment rate λ react with respect to those variables.

We start with V_m . It is quite intuitive that if firms are very well informed about the productivity of workers, parents will want to invest a large share of their income in their children education. This is because, in that case children's income necessarily depends on their productivity, and not in other considerations that could be inferred from their background. Consequently, the return of investment in education becomes large: it translates into future income readily.

Thus, it is not surprising the following result:¹⁵

Result 10. *An increase in the accuracy of the human capital signal (a decrease of V_m) results, in steady state, in an increase in the proportion of income invested in education.*

Notice that in our context parents want to educate their children only in the measure that education makes them *look* as having more human capital. Wages are a result of how productive agents *look*, not of how productive they *are*. Thus, educational investment is very sensitive to β_m . If its value is very low, you will not invest in your children, as the investment would not change the perception that the market has on them. You would rather eat the resources.

This is a strong force by which parents respond to the greater rewards to education arising from a more accurate human capital signal by investing more heavily in the human capital of their child.

This force works in exactly the opposite in response to an increase in the accuracy of the information available on background. As we saw in result 7, a decrease of V_a produces a steady state decrease of β_m , decreasing the return of investing in education.

There is, though, an additional force in play that complicates the analysis: A larger effect of the perception of background on wages is similar to a reduction of the discount rate. The

¹⁵Proof of results 10 and 11 is in appendix J.

reason is as follows. A larger effect of background on wages (larger $\hat{\beta}_a$) means that any increase in the wage of your child that comes as a consequence of you increasing her education will have a larger effect in your grandchild. The background of your grandchild will be better, and she will be better treated by the market. Thus, the effects of the investment last longer across the generations, something that you appreciate.

Result 10 demonstrates that is not strong enough to overturn the first one in response of investment to an increase in the accuracy of the merit signal. Nevertheless, it complicates the analysis of the response of investment to an increase in accuracy of the signal on background, which we summarise this in the following result ¹⁶

Result 11. *An increase in the accuracy of the signal on background (a decrease of V_a) may increase or decrease investment. A sufficient condition for $\frac{d\lambda}{dV_a} < 0$ is:*

$$\frac{V_m}{V_\omega} \left[\frac{1 - \alpha^2 [\beta_a (1 - \beta_m) + \beta_m]}{[(1 + \delta) - \alpha \beta_a]} \right] > 1 \quad (36)$$

Thus, V_m and V_a have different effects on λ . In the next section we will use these differences to identify the degree of *meritocracy* and the prevalence of advantages in different societies.

7 Calibration

The existing data on on inequality and intergenerational mobility is often interpreted in public discussion in terms of meritocracy and the prevalence of advantages. Thus, when a society is shown to have relatively low intergenerational mobility it is often understood to be *not meritocratic*, at least vis a vis another society with more intergenerational mobility. Our model contradicts such a view, our main insight being that given the complexities of the transmission of inheritance one should be cautious when extrapolating from the existing data on intergenerational mobility and inequality into which are their causes.

Thus, albeit the motivation and the main contribution of this the paper is eminently theoretical, it is interesting to see how the model interprets such data. We have seen that two societies showing similar levels of inequality and mobility could have arrived to such a point via different routes. For instance, Italy and the US have similar levels of mobility and inequality, but it is perfectly possible that one may have a high degree of *meritocracy* and low prevalence of advantages, while the other may be just the reverse.

¹⁶Proof in appendix J.

Moreover, at the light of the model it is perfectly possible that a society has less intergenerational mobility than another while paying individuals much more according to their objective measures of productivity and less in function of their background. Thus, in principle the US could have a higher correlation of income parent-son than Sweden, while being more meritocratic.

The objective of this section is to look at how our model maps the existing data into the degree of meritocracy and the prevalence of advantages of different societies. It is by no mean an empirical “test” of the model, and we do not want to read the results in a literal sense. All models are “wrong”, in the sense of being a simplification, and when dealing with complex issues (such as inheritance and advantages) in order to be comprehensible they need to be abstract from many issues that are undoubtedly relevant.

Nevertheless, a model is better than no model. The literal and naive reading of the data that is normally made (less mobility equals less meritocracy and more prevalence of advantages) implies a much bigger oversimplification, if not a logical fault. Consequently, the natural next step is to entail the degree of meritocracy and the prevalence of advantages that our model suggests that exist in different societies.

7.1 Procedure

Obviously, there exists no direct data on β_m or $\hat{\beta}_a$, much less on V_a and V_m . All that is available is data on ρ , V_y and λ . Our goal is to use these data in order to reveal (for any society j) the values of V_a^j and V_m^j (and thus β_m^j and $\hat{\beta}_a^j$) that would generate within the model values of ρ^j and V_y^j , and λ^j that are the closest possible to the data. To clarify notation, from now on we will include an index $j \in \{1, \dots, J\}$, indicating the society (country) to which we are referring.

Our starting assumption is that societies differ in the amount of information available to firms to evaluate the workers productivity (V_m^j), and in the amount of information available to firms to evaluate the workers background (V_a^j). The other three parameters in the model (α , V_ω and δ) are assumed to take the same value in all societies.

The discount rate δ is assumed to be 1% annually. It is much more difficult to give numerical values to the elasticity of human capital accumulation to educational investment (α) or the variance in the process of exogenous shocks in the acquisition of human capital (V_ω). Thus, we will calibrate then within a very comprehensive and internally coherent data set.

In addition to V_a^j and V_m^j societies could differ in the process of human capital accumulation. In our model these differences would be captured in the intercept Z in equation 6. In principle

we could also calibrate for those variables Z^j in each country j , and include average income as another target, but notice that such extension would be irrelevant. The reason is that differences of productivity in the human capital accumulation function would map one to one into differences in average income, while they do not affect at all the rest of the variables in the model. We would match average income exactly, but it is just because for each country we have an extra variable dedicated to that task. Given that this variable does not affect at all our variables of interest it would be a futile exercise that we ignore.

Thus, summarizing. We will assume that all societies have a common discount rate of $\delta = 1\%$ annually, and they also have common values for α and V_w , but the values of these variables will be discerned by calibrating the model. Societies are different because they differ in the values of V_m^j and V_a^j , which are the exogenous force of variation in the degree meritocracy and the prevalence of advantages across societies. Societies may also differ on Z_j , but that is irrelevant and we can ignore it.

These different values of V_m^j and V_a^j , together with the common values of α , δ and V_w , generate within the model values for the variables for which we have counterparts in the data. Specifically, within the model we get the degree of intergenerational mobility ρ^j , the degree of inequality V_y^j , and the investment in education λ^j . Our goal is to find for all countries for all countries $j \in \{1, \dots, J\}$ the values of V_a^j and V_m^j , and the global variables α and V_w such that the model values of these moments (V_y^j , ρ^j and $\lambda^j \forall j$) within the model are the closest possible to their data counterparts.

We do this exercise in two different contexts. First we look at the configuration of meritocracy and advantages across the geography of the US. For this we make use of the extraordinary data made available by Chetty et al. (2014). They use administrative data in order to estimate intergenerational mobility in 741 US Commuting Zones. Moreover they report the degree of inequality and education investment in each of them (along with many other variables). This is, assuming that all US Commuting zones share the same value of α and V_w , we calculate the values of V_a^j and V_m^j (plus α and V_w) that make the model closest to the data. We determine their geographical distribution, and we correlate the results with other macroeconomic variables across commuting zones. Thus, we can look not only at the correlation of meritocracy with advantages, but also at which variables are correlated with each of them.

In contrast with this set of US data (that is extremely rich and with very accurate measurements of intergenerational mobility), international data on intergenerational mobility is scarce

and not always directly comparable. Nevertheless, in the last 10 years there has been a major, and overall successful, effort to generate comparable measures. Thus, the values reported by Corak (2013) have received much attention as they show that mobility and inequality are negatively correlated across countries (the so called “Great Gatsby Curve”), and that the US is in the dismal extreme of this relationship (i.e., the US has relatively low mobility and high inequality). Our second exercise is to look at how our model filters these data. This is, we look at the degree of meritocracy and the prevalence of advantages that our model suggests that could have generated the observed data across the different countries. Countries that look similar in the “Great Gatsby Curve” are very much apart in the meritocracy versus advantages dimension. Some countries with very low mobility appear to be highly meritocratic.

Still, for doing this second exercise we will make use of knowledge gained with the much more reliable US data. So we turn to that problem first.

7.2 A Geography of Merit and Inherited Advantages in the US

As we mentioned, we obtain all data from Chetty et al. (2014).¹⁷ From their extensive data set we use principally (albeit not exclusively, see below) data determining the degree of inequality, the degree of intergenerational mobility and the investment in education in 741 US commuting Zones.

Summarizing our approach: we assume a discount rate of 1% annually (this is, $\delta = 0.348$ counting that a period is one generation which we take to be 30 years), and find the values of V_a^j and V_m^j (in each of the commuting zones) and the values of α and V_ω that maximize the correlation between model generated moments of ρ^j , V_y^j and λ^j with their empirical counterparts. We explain the procedure in detail below, but first we describe the data we use, and discuss their caveats.

7.2.1 Description of US Data

Data on Inequality

Chetty et al. (2014) document the GINI index of the distribution of income at Commuting Zone level. This is a different measure of inequality than the variance of log income, which is the one we use in the model. In order to run the calibration procedure we translate the GINI

¹⁷Data available from “Chetty, Hendren, Kline and Saez (2014): Descriptive Statistics by County and Commuting Zone”. Section of <http://equality-of-opportunity.org/index.php/data>

index into the variance of log income assuming that income is distributed with a log normal.¹⁸

Data on Intergenerational Mobility

The preferred measure of intergenerational mobility in Chetty et al. (2014) is the rank-rank correlation, which differs from the parent child correlation in that it looks only at persistence in the relative position, not weighting at the differences in this position vis-a-vis the other members of society. There are good reasons for centering in the rank-rank correlation, as it controls better for the relatively scarce information on lifetime income of the children. Nevertheless, this measure does not map directly into our model (or into any formal model that we know of), and it seems reasonable to translate it into parent-child income correlation before calibrating the model.

We do so by determining which rank-rank correlation would be generated by an income generating process following a AR(1) with normal shocks, with exogenous mean, variance and autocorrelation, and then inverting this relationship. In appendix N we explain in detail the procedure followed. The relationship does seem to be independent of $\bar{y}$ and V_y , and ρ and the rank-rank correlation are closely associated (as it can be seen in figure 8). So much so, that there is almost no difference between using the rank-rank correlation and the intergenerational correlation implied.¹⁹

Educational Investment

As an empirical counterpart of λ we use an transformation of the ratio of spending per student to the mean household income of the each Commuting Zone. The transformation, and its rational, is as follows.

Spending in education of the children is larger, perhaps substantially larger, than what is spent in formal schooling. It includes, for instance, the spending in real state incurred in order to enjoy the externalities generated from the presence of high income students in the neighborhood, or the spending in extra-curricular activities (from violin lessons to trips to museums). We have no way of including data on all this activities, and Chetty et al. (2014) only report on spending per pupil.

It is reasonable to expect, nevertheless, that total spending will be related to formal spending per pupil. Thus, we define a variable k^j per each country j capturing the standardized (i.e. zero mean and variance equal to one) ratio of spending per student divided by mean household income

¹⁸Specifically, we invert the relationship: $GINI = 2\Phi\left(\frac{\sigma_y}{\sqrt{2}}\right) - 1$. See Aitchison & Brown (1963)

¹⁹ We performed the same exercise using directly the rank-rank correlations instead of our transformation. The calibrated parameters are almost identical, and the correlation of the values of β_m^j and $\hat{\beta}_a^j$ with our benchmark is larger than 97%

of each CZ. We then define $\hat{\lambda}_i = \mu + \sigma \hat{\kappa}_i$ as our proxy to the values of the data moment, and our empirical counterpart to λ^j . μ and σ are chosen by the calibration routine (see below) in order to equalize the mean and variance of the model value of λ^j to its empirical counterpart. As our goal is to maximize the correlation between empirical and model variables across countries (and not to equate their mean and variance), the inclusion of μ and σ is innocuous except for the mostly aesthetic effect of equalizing the distribution of empirical and model moment: this does not affect the correlation across commuting zones.

Not all cells in the data are populated, thus we eliminate all commuting zones on the basis of any missing values in the variables needed to construct $\left\{ \hat{V}_y(i), \hat{\rho}_i, \hat{\kappa}_i : \forall i \right\}$. This leaves a data set of 701 commuting zones.

7.2.2 Calibration Procedure

We look for the values of $\{V_m^j, V_a^j\} \forall j$ and $\{\alpha, V_\omega\}$ that minimize the following objective function:

$$\min_{\{V_m^j, V_a^j\}_{\forall j}, \alpha, V_\omega} L = \left\{ \begin{aligned} & (1 - \text{Corr}[\hat{\rho}, \rho])^2 + \left(1 - \text{Corr}[\hat{V}_y, V_y]\right)^2 + \left(1 - \text{Corr}[\hat{\lambda}, \lambda]\right)^2 \\ & + (\ln \bar{\rho}^D - \ln \bar{\rho}^M)^2 + (\ln \bar{V}_y^D - \ln \bar{V}_y^M)^2 \end{aligned} \right\}$$

where the elements of the first line of the loss function are (one minus) the correlation across commuting zones between the model moments and their data counterparts. Thus, for instance, $\text{Corr}[\hat{\rho}, \rho]$ is the correlation across all commuting zones between the model generated values of ρ (as a consequence of $\{V_m^j, V_a^j\}_{\forall j}, \alpha$ and V_ω) and their data counterparts $\hat{\rho}$. The second line of the loss function is the square of the difference between the model and data values of the means of ρ and V_y in all commuting zones.

We maximize the correlations, instead of minimizing the square errors, of the moments because in this way we normalize the magnitudes of the three variables, giving them effective equal weight. The loss is convex in the square of the correlation mistake, indicating that we prefer to equate the loss across the three moments. Notice that this does not fix either the mean or the variance of the moments. This is, the mean and variance of the moments of the data could be vastly different than the ones in the model, while the correlation could be high. In order to avoid this we also account (in the second line) for the deviations of the mean of ρ and V_y from the ones observed in the data.

Note two things. First, we could as well account for deviations of the variances. We choose not to do it in order to have some untargeted moments where to check how well we are doing in

the calibration. Second, we do not include the mean or the variance of λ , because these already perfectly match by construction via the inclusion of σ and μ (as explained above).

The optimization algorithm²⁰ runs an outside loop to calibrate the global parameters α and V_w . It then runs an inner loop to calibrate the local parameters V_m^j , V_a^j for each location j conditional the current values of α and V_w .

7.2.3 Calibration Results

In table 1 and figure 1 we report the calibrated parameters and the fitness of the targeted moments. Each scatter plot draws the value of the data and the moment value for one of the targeted moments in each of the CZ's. In spite of having many degrees of freedom, the correlations are substantially high and, less surprisingly, the mean of ρ and V_y across CZ's is pinpointed very accurately. Nevertheless, being targeted moments, it is difficult to interpret this as goodness of fit.

The variances of ρ and V_y across CZ's are the only untargeted moments that we have left. On this we do quite well in the first, but less so in the second. The standard deviation of ρ in data is 0.060, in the calibration it's remarkably close: 0.063. On the other hand, we do very poorly in the variance of inequality across CZ's, its standard deviation in the data is 0.304, in calibration it's 0.095.

The calibrated values of α and V_w are 0.409 and 1.038 respectively.²¹ Instead of looking at the of V_m^j and V_a^j it is more convenient to look at the precisions. We define $P_m^j = \log\left(\frac{1}{V_m^j}\right)$ and $P_a^j = \log\left(\frac{1}{V_a^j}\right)$. Then, in figure 2 we plot the distribution of the logs of the implied precision of the merit and background signals across all CZ's, which visually suggests that they are negatively correlated across CZ's. This is indeed the case. In table 1 we show that more meritocratic places have in average less advantages: the correlation of P_a with P_m is negative, which translates into a negative correlation of $-.54$ between the rewards on merit and background. This covariance between the two sources of exogenous variation across CZ's explains why when comparing across CZ's:

- CZ's with more information on background have less mobility, but also less inequality (positive correlation of P_a^j with data and model ρ^j)

²⁰It is explained in detail in appendix O

²¹We have performed the same exercise assuming instead values of the discount of 0.5% and 1.5% annually (instead of the 1% of our benchmark). The change in the calibrated parameters is small, and the correlation of the calibrated values of β_m^j and β_a^j with the ones in our benchmark calibration is always larger than 95%

Calibrated parameters

α = 0.409
 V_ω = 1.038
 μ = 0.1694
 σ = 0.0219

 $Loss\ F_n$ = 0.285

Targeted Moments

Correlation at CZ level
data and model

$Corr\left(\hat{\rho}, \rho\right) = 0.84$
 $Corr\left(\hat{V}_y, V_y\right) = 0.64$
 $Corr\left(\hat{\lambda}, \lambda\right) = 0.64$

	Model	Data
$E(\rho)$	0.309	0.308
$E(V_y)$	0.613	0.616

Modelled

	P_a	P_m	$1-\rho$	$Gini$
P_a	100%			
P_m	-22%	100%		
$1-\rho$	-81%	-3%	100%	
$Gini$	-26%	80%	-6%	100%

Modelled

	$\hat{\beta}_a$	β_m	$1-\rho$	$Gini$
$\hat{\beta}_a$	100%			
β_m	-54%	100%		
$1-\rho$	-90%	12%	100%	
$Gini$	-38%	98%	-6%	100%

Data

	P_a	P_m	$\hat{\beta}_a$	β_m	$1-\rho$	$Gini$
P_a	100%					
P_m	-22%	100%				
$\hat{\beta}_a$	87%	-34%	100%			
β_m	-42%	82%	-54%	100%		
$1-\rho$	-68%	5%	-77%	12%	100%	
$Gini$	-18%	53%	-15%	59%	-29%	100%

Table 1: Calibrated Parameters and Fitness of Targeted Moments.

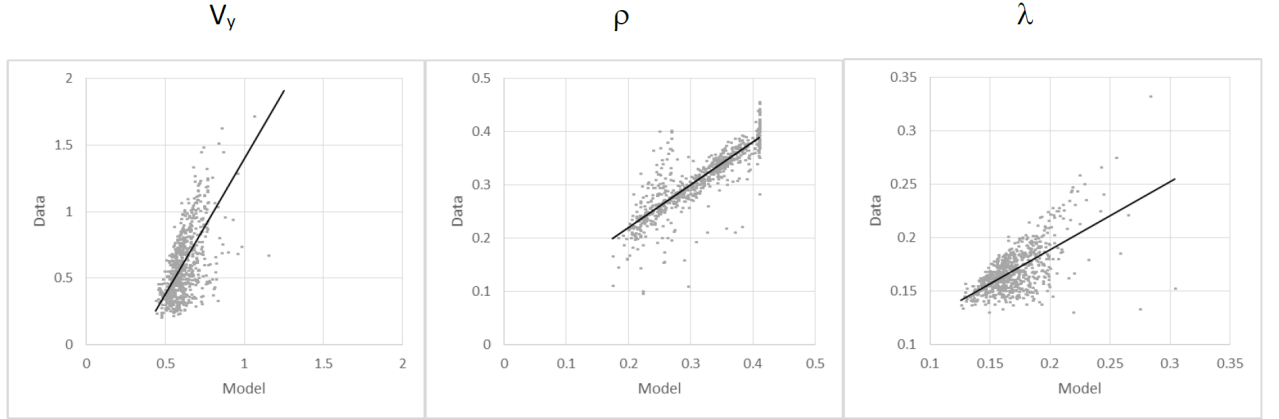


Figure 1: Calibrated Common parameters & Model and Data Moments for all CZ.

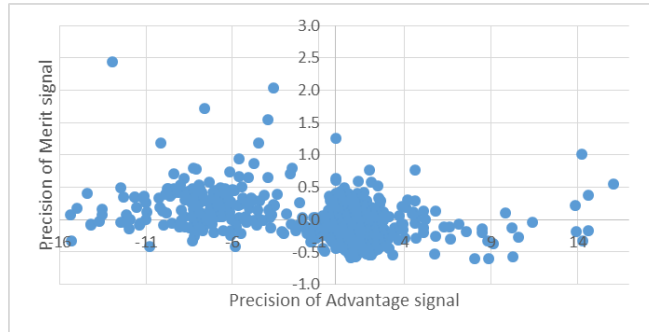


Figure 2: Distribution of calibrated values of the precision of the signals: P_m^j and P_a^j for all CZ's.

- CZ's with more information on merit are suffer more inequality but have almost no different mobility (P_m^j is very positively correlated with data and model V_y^j , but only very marginally positively correlated with the model ρ^j , and very marginally correlated with the data $\hat{\rho}^j$).

In figure 3 we plot the geographical distribution of our proxits for merit and reward, along with the data values of mobility and inequality for comparison. Data is in green in quintiles: darkest (lightest) green is highest (lowest) quintile of variable being mapped. We plot in red the CZ with insufficient data. Here the same pattern is apparent. Due to the negative correlation between the two precisions, the precision of merit seem to explain inequality better, while the precision of background seems to explain mobility better.

It may also be of interest to correlate the implied precision of both signals and the rewards to merit and background to a relatively large set of Commuting Zone characteristics available from Chetty et al. (2014). In table 2 we do so for the 418 Commuting Zones that have information on all the properties that we look. We run a regression²² of these characteristics against both mobility and inequality in each CZ, but also on P_a , P_m , $\hat{\beta}_a$ and β_m . We only report the values for theose characteristics that are significative in at least one regression.

Income level is not very significantly related to either precision, while income *growth* is negatively related to the prevalence advantages and positively related to the prevalence of merit.

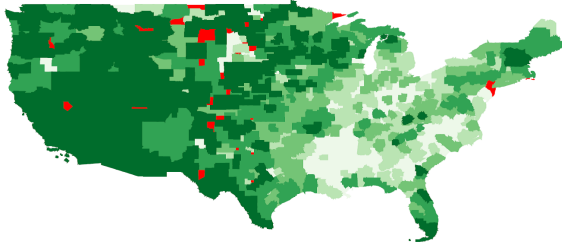
Racial segregation and the share of manufacturing, particularly the second, seem to be the variables that correlate better with the precision of the background signal (and with its reward $\hat{\beta}_a^j$). It also correlates negatively with the size of the foreign community in the CZ, and with education achievement indicators (teacher student ratio and college graduation rate).

The share of manufacturing is also the variable that better correlates with the precision of the merit signal, but negatively. The fraction of foreign born correlates very positively. This suggest that areas with relatively large service sectors and more migrants tend to be more meritocratic. The reward to merit also decreases with the proportion of African Americans. Educational achievement (share of college graduates) correlates positively with it, while tax progressivity does it negatively.

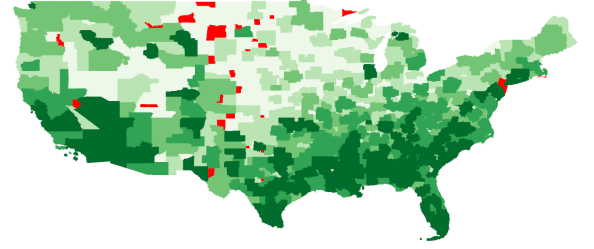
Finally, Social Capital seems to correlate positively with the incidence of background and negatively with the incidence of merit.

In table 3 we perform the same exercise in a sample including all the commuting zones but less characteristics. It shows the same general patterns.

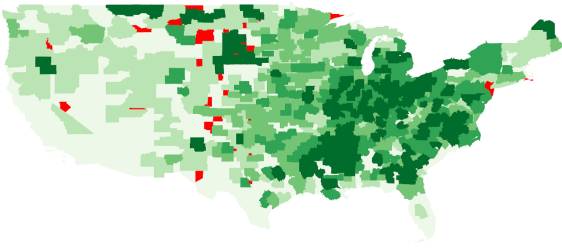
²²All variables in all exercises have been standardized by subtracting mean and dividing by standard deviation.



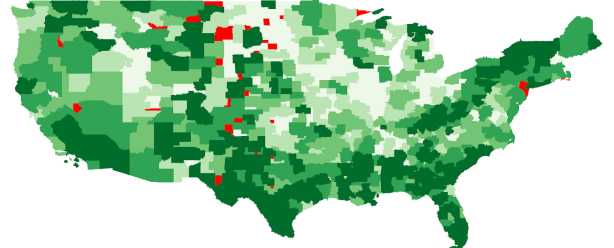
(a) Mobility, ρ



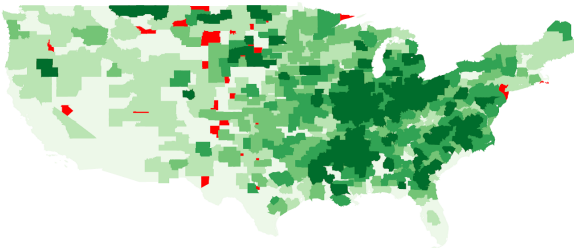
(b) inequality, Gini Index.



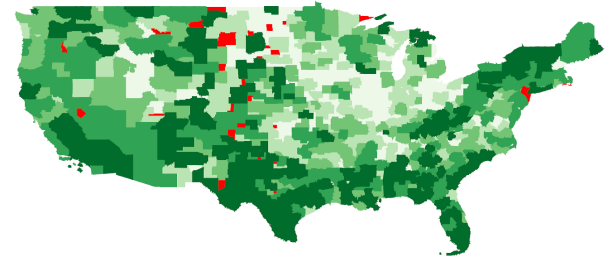
(c) Precision of Advantage Signal



(d) Precision of Merit Signal



(e) Reward to Advantage Signal, $\hat{\beta}_a$



(f) Reward to Merit Signal, β_m

Figure 3: Geographical patterns in the US

Sample Correlations.						
	Mobility	Gini	1/ V_a	1/ V_m	$\hat{\beta}_a$	β_m
fraction black	-62%	58%	24%	9%	25%	18%
racial segregation	-43%	30%	12%	6%	19%	15%
fraction with commute 15 mins	47%	-58%	-12%	-11%	-21%	-14%
household income per capita	13%	-8%	-10%	-20%	-2%	-35%
state income tax progressivity	4%	-19%	6%	-10%	6%	-18%
state eitc exposure	14%	-31%	3%	-7%	0%	-9%
teacher student ratio	2%	16%	-5%	-12%	-3%	-18%
high school dropout rate income adjus	-35%	47%	10%	17%	9%	23%
college graduation rate income adjus	4%	-17%	-2%	-4%	-2%	-2%
labor force participation rate	32%	-35%	-11%	-32%	-7%	-44%
manufacturing employment share	-40%	-3%	40%	-16%	52%	-32%
teenage labor force participation rate	52%	-64%	-13%	-23%	-16%	-38%
migration inflow rate	24%	8%	-24%	-7%	-25%	-5%
migration outflow rate	27%	2%	-22%	-8%	-27%	-1%
fraction foreign born	29%	22%	-31%	13%	-30%	25%
social capital index	32%	-66%	4%	-29%	4%	-46%
fraction of children with single mother	-65%	62%	25%	17%	23%	31%
fraction of adults divorced	-13%	20%	-2%	-5%	4%	-5%
income growth	51%	-35%	-32%	2%	-36%	2%

Partial Correlations.						
	Mobility	Gini	1/ V_a	1/ V_m	$\hat{\beta}_a$	β_m
fraction of black	-13%	11%	-4%	-22%	3%	-22%
	-2.26	1.69	-0.42	-1.85	0.36	-2.93
racial segregation	-30%	0%	9%	-7%	19%	-1%
	-9.68	-0.10	1.75	-1.14	4.08	-0.23
fraction with commute 15 mins	17%	-21%	-19%	-18%	-21%	-12%
	3.37	-3.39	-2.26	-1.65	-2.75	-1.81
household income per capita	5%	17%	-21%	8%	-15%	-8%
	0.98	2.63	-2.38	0.70	-1.91	-1.16
state income tax progressivity	-1%	-6%	5%	-12%	5%	-18%
	-0.30	-1.47	0.86	-1.74	1.08	-4.14
state eitc exposure	6%	-7%	-4%	7%	-10%	17%
	2.19	-2.01	-0.82	1.15	-2.43	4.42
teacher student ratio	23%	-7%	-8%	-14%	-12%	-20%
	6.33	-1.51	-1.28	-1.90	-2.24	-4.13
high school dropout rate income adjusted	-5%	13%	3%	18%	0%	13%
	-1.55	3.40	0.63	2.68	-0.04	3.11
college graduation rate income adjusted	-1%	2%	-12%	2%	-9%	11%
	-0.24	0.71	-2.62	0.41	-2.24	3.00
labor force participation rate	6%	-8%	0%	-30%	-1%	-18%
	1.29	-1.48	0.01	-3.02	-0.14	-2.82
manufacturing employment share	-24%	-21%	32%	-21%	48%	-36%
	-6.49	-4.64	5.12	-2.63	8.66	-7.31
teenage labor force participation rate	15%	-17%	9%	10%	4%	-10%
	2.53	-2.49	0.93	0.81	0.42	-1.27
migration inflow rate	-8%	14%	-11%	7%	-9%	9%
	-1.51	2.18	-1.17	0.57	-1.03	1.28
migration outflow rate	12%	-24%	10%	-21%	7%	-21%
	2.55	-4.12	1.22	-2.05	0.91	-3.24
fraction foreign born	18%	19%	-23%	8%	-16%	22%
	4.65	4.02	-3.51	0.99	-2.73	4.24
social capital index	-25%	-21%	31%	-37%	38%	-32%
	-4.50	-3.18	3.37	-3.20	4.67	-4.36
fraction of children with single mother	-24%	20%	21%	22%	12%	30%
	-3.93	2.81	2.10	1.69	1.28	3.72
fraction of adults divorced	-1%	-9%	-8%	-34%	4%	-27%
	-0.26	-1.86	-1.16	-3.96	0.71	-4.98
income growth	13%	2%	-19%	15%	-12%	14%
	3.35	0.43	-2.99	1.85	-2.14	2.70
R^2	75%	68%	34%	27%	46%	59%

Table 2: 418 CZ without blank data. In the Partial Correlations panel each column is a different LHS variable. In the first line of each row we place the coefficient of the variable (in %), in the second line the t-statistic.

Sample Correlations.						
	Mobility	Gini	$1/V_a$	$1/V_m$	$\hat{\beta}_a$	β_m
fraction black	-57%	54%	25%	10%	25%	14%
racial segregation	-41%	28%	14%	3%	22%	6%
fraction with commute 15 mins	41%	-56%	-8%	-6%	-18%	-5%
household income per capita	17%	7%	-17%	-18%	-9%	-25%
local government expenditures per capita	31%	-10%	-21%	6%	-26%	12%
state income tax progressivity	14%	-8%	-11%	-7%	-8%	-9%
state eitc exposure	13%	-25%	1%	-4%	-4%	-4%
labor force participation rate	22%	-26%	-6%	-29%	-2%	-35%
manufacturing employment share	-44%	-2%	39%	-20%	51%	-34%
teenage labor force participation rate	44%	-58%	-10%	-22%	-13%	-30%
fraction foreign born	32%	30%	-40%	18%	-37%	30%
fraction religious	6%	-29%	7%	-11%	5%	-17%
fraction of children with single mother	-59%	58%	24%	19%	23%	27%

Partial Correlations.						
	Mobility	Gini	$1/V_a$	$1/V_m$	$\hat{\beta}_a$	β_m
fraction black	-12%	17%	4%	1%	1%	0%
	2.80	3.88	0.77	0.09	0.26	0.03
racial segregation	-27%	-3%	15%	-1%	23%	-1%
	-9.89	-1.15	4.00	-0.36	6.49	-0.30
fraction with commute 15 mins	9%	-25%	3%	-7%	0%	-1%
	2.11	-5.33	0.58	-1.09	0.04	-0.19
household income per capita	3%	18%	-10%	-13%	-2%	-20%
	0.75	4.75	-1.98	-2.41	-0.48	-4.33
local government expenditures per capita	6%	-3%	-4%	9%	-8%	12%
	2.11	-0.91	-1.06	2.09	-2.24	3.41
state income tax progressivity	5%	-6%	-4%	-10%	0%	-14%
	1.92	-2.38	-1.27	-2.80	0.03	-4.70
state eitc exposure	6%	-5%	-3%	12%	-13%	21%
	2.29	-1.78	-0.75	2.98	-3.63	6.11
labor force participation rate	4%	-12%	1%	-22%	1%	-19%
	1.18	-3.30	0.17	-4.40	0.22	-4.48
manufacturing employment share	-29%	-15%	35%	-19%	47%	-33%
	-10.50	-5.11	9.19	-4.60	13.66	-9.29
teenage labor force participation rate	2%	-16%	12%	-2%	12%	-15%
	0.35	-2.87	1.72	-0.24	1.86	-2.23
fraction foreign born	27%	21%	-27%	17%	-26%	27%
	9.68	7.00	-7.18	4.20	-7.43	7.68
fraction religious	-8%	3%	5%	-6%	6%	-10%
	-2.92	0.85	1.25	-1.46	1.64	-2.85
fraction of children with single mother	-29%	25%	18%	7%	14%	12%
	-6.44	5.11	2.93	1.08	2.44	2.15
R^2	65%	59%	34%	19%	44%	42%

Table 3: 701 CZ. In the Partial Correlations panel each column is a LHS variable. In the first line of each row we place the coefficient of the variable (in %), in the second line the t-statistic.

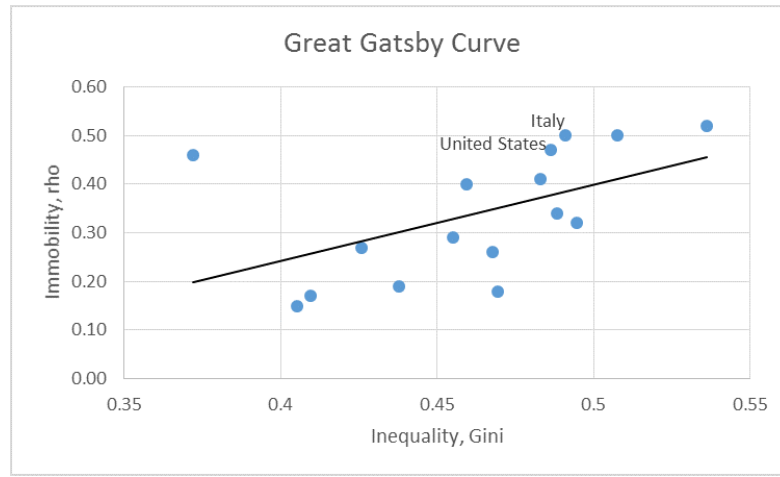


Figure 4: The Great Gatsby Curve

Summarizing. When we calibrate our model to US data, the implied values of the precision of the merit signal are negatively correlated with the implied precision of the background signal. The prevalence of advantages due to background is very negatively correlated with the degree of mobility, not so much to inequality. High rewards to merit, on the contrary, seem correlated to inequality and only little with mobility. Racial characteristics of a CZ have the expected effects: segregation and racial diversity increase the prevalence of advantage and decrease meritocracy. State tax progressivity is negatively related to the precision of the merit signal. The productive structure of the commuting zone is what most clearly correlates with both merit and advantages: manufacturing associates positively with advantages and negatively with merit. Finally, foreign migrants tend to live in CZ's where the implied precision of merit signals is high and where the one of background is low.

7.3 Calibration to International Data

The discussion on the preeminence of inheritance and the possible end of the American dream has focused a substantial amount of attention in the so-called “Great Gatsby Curve”.²³ This is a relationship between the degree of inequality and a set of comparable measures of the intergenerational persistence of income across countries, which we depict in figure 4.²⁴ It indicates that societies showing less inequality are more likely to show more mobility. In itself it is not surprising, and it has been extensively discussed in the literature²⁵ What did get the public attention was the fact that the US is at the dismal extreme of the curve: the US has a

²³The origin of the name is obscure and difficult to ascertain. It is also quite confusing. The relationship became prominent following Krueger (2012), albeit it was based in previous work developed by Corak (2013).

²⁴This positive correlation has been documented across countries (Corak (2013)), across US commuting zones (Chetty et al. (2014)) and across Italian provinces (Güell et al. (2015)).

²⁵As mentioned above Hassler et al. (2007) and Solon (2004), for instance, discussed the theoretical relationship between inequality and mobility before this was documented by Corak

very low degree of mobility and a large degree of inequality. It is very much in the same place than Italy and much worse than, for instance, the Scandinavian countries. This has been read as indicating that the US is no more the land of opportunity, but one where merit is secondary to privilege.

It is naturally appealing to check at whether our model reads the data as indicating that the prevalence of merit within the US is relatively large or small because, as we have seen, from the point of view of theory our could as well imply the same reading than the naive equalization of meritocracy with mobility (i.e., relatively low values of β_m in the US), or the opposite. Moreover, the calibration of the model to international data allows us to do a sensible simulation of the model, indicating the magnitudes that we are talking about.

Nevertheless, to do so we need to address even more issues than in our previous section. First due to the difficulty of obtaining comparable data across countries, and second because our model obviates many institutional issues that should be taken into account when doing this exercise.

7.3.1 Data and Procedure

Our approach, like in the previous section, is to presume that the exogenous source of variation between countries lies in the precision of their signals on merit and background. Consequently δ , α and V_ϵ are kept constant across all countries. As before, we assume $\delta = 0.348$, which amount to a 1% annual discount. Instead of re-calibrating the values of α and V_ω we use the values that we obtained in the previous exercise with US Commuting Zones. This is: $\alpha = 0.409$ and $V_\omega = 1.038$.

We use two sources of international data. We obtain from Corak (2013) values of the inter-generational correlation of income between parents and sons for a selection of countries. The OECD provides data on the distribution of income and the degree of educational investment. In order to abstract from issues that might depend on the stage of development we opt to do our exercise only with developed countries.²⁶ The intersection of two sets consists on 15 countries.²⁷

When looking at international data we need to take into account that in the model we have abstracted from institutional differences across countries. In particular we have abstracted from (i) the possibility of taxation and redistribution, and (ii) from the possibility of the procurement

²⁶The only country at a different stage of development for which we have data is Chile. We also performed our exercise including Chile and obtain the same qualitative results.

²⁷The countries are: Australia, Canada, Denmark, Finland, France, Germany, Italy, Japan, New Zealand, Norway, Spain, Sweden, Switzerland, United Kingdom, United States.

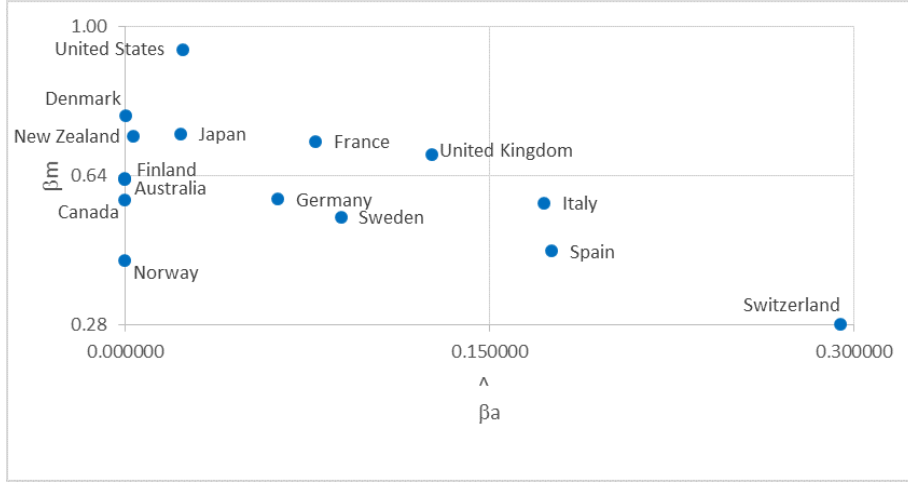
Country	ρ	V_y		Education over GDP	
		PreTax	PostTax	TotalEd	PrivateEd
Australia	0.26	0.777494	0.375999	5.51%	1.34%
Canada	0.19	0.669317	0.340614	5.77%	0.70%
Denmark	0.15	0.563195	0.190069	8.05%	1.49%
Finland	0.18	0.783977	0.230952	5.60%	0.40%
France	0.41	0.837894	0.282429	6.07%	0.63%
Germany	0.32	0.885219	0.2708	4.87%	0.22%
Italy	0.50	0.87078	0.332033	5.12%	0.82%
Japan	0.34	0.859455	0.375552	5.91%	1.93%
New Zealand	0.29	0.730179	0.362101	6.51%	1.18%
Norway	0.17	0.576651	0.203958	5.16%	0.09%
Spain	0.40	0.746395	0.328532	4.98%	0.59%
Sweden	0.27	0.628584	0.21964	5.96%	0.00%
Switzerland	0.46	0.467887	0.292865	5.23%	0.08%
United Kingdom	0.50	0.941469	0.391376	5.38%	1.08%
United States	0.47	0.851498	0.484567	7.00%	1.47%

Table 4: International Data

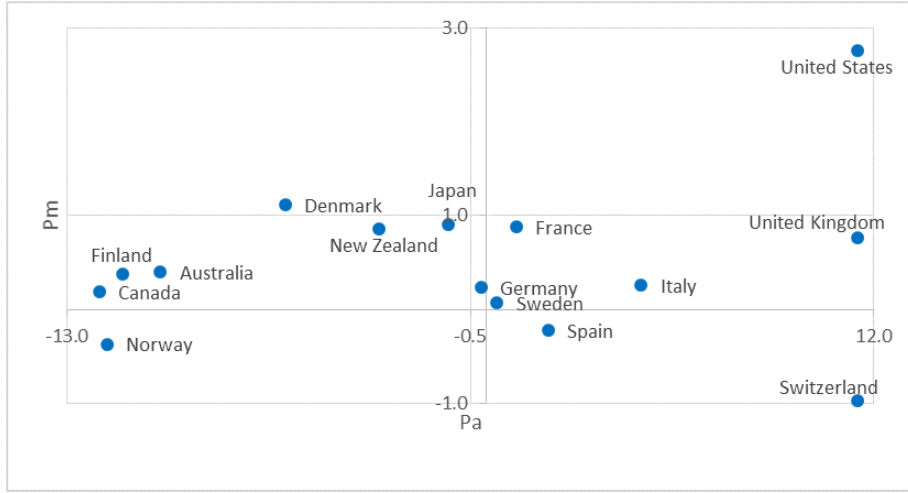
of public education.

Both aspects are of obvious importance, and both affect the quantitative implications of the model. The only reason for not incorporating their analysis in the model is that it would have impeded us to obtain analytic results. Including these aspects while preserving the log-linear nature of the model, which allows for it to be analytically tractable, is a complicated task that is part of our future research on this topic. For the present paper instead of developing such an extended theory we opt to adapt the data to our model. This is, we calibrate the model to (1) both the distribution of income before and after taxes, and (2) both the investment in private education and in total education (i.e., both private plus public). The qualitative results are robust to any of the four possible data configurations. Given that the intention of our exercise is not to provide with a numerical value to the degree of “meritocracy”, but to show that our way to looking at the problem can be potentially important, we believe that this should suffice to convince the reader.

In table 4 we present our data. The data on mobility (ρ) is directly defined as in the model and we use it as it is. We map the GINI coefficient of the distribution of income (provided by OECD) into the variance of log income in the same manner as in the previous section. OECD data allows us to do this for the distributions of income before and after taxes. In order to proxy for human capital investment we use either the total educational or private educational spending divided by GDP, both figures from OECD. As in the previous section we consider true human capital investment effort λ to be a linear transformation of the data, with the values of



(a) Rewards to merit and background $\beta_m^j, \hat{\beta}_a^j$



(b) Precisions of signals on merit and background P_m^j, P_a^j

Figure 5: Implied Rewards and Precisions of Merit and Background signals for all countries for the income distribution before taxes and total educational investment.

the transformation (μ and σ) being obtain from the calibration.²⁸

Thus, we have four data configurations: (pre-tax income, total educational spending), (post-tax income, total educational spending), (pre-tax income, private educational spending) and (post-tax income, private educational spending). In each of them we apply the same calibration procedure than in the previous section for selecting $(V_m^j, V_a^j) \quad \forall j$ and for μ and σ , but we do not adjust α or V_ω .

7.3.2 Results of Calibration to International Data

In figure 5 we plot the implied rewards to merit and advantage (in panel (a)), and the implied precision of both signals (in panel (b)) using the income distribution before taxes and the total spending in education. We can extract three lessons from these graphs.

- It may be the case that in a country the *unconditional* effect of background on income be high while the effect *conditional* on the individual observed abilities is very low.

We know that parental background matters a lot in the US relative to other countries because that is precisely what it means that its mobility is low. Moreover, in the calibration the imputed precision of the signal on background is among the highest. Nevertheless, the implied reward of having a good background $\hat{\beta}_a$ is among the lowest. The reason is that the implied precision of the direct objective signal on the productive ability of individuals is so much higher in the US than anywhere else. This results in that the rewards of merit $\hat{\beta}_m$ are very large, and β_a very low. The unconditional effect of background is large, but the effect *conditional* to ability is very small.

Perhaps this is our most important insight: Rewards to merit may be very large while having a low degree of mobility. Actually, large rewards to merit will cause low mobility and high inequality. Moreover, that is how the model reads the situation of the US.

- Italy and the US look very much the same in the Great Gatsby curve, but they are very far apart in the maps of imputed rewards and precisions.

The reason is that they are very different in the educational effort that they make. The US invests in education much more than the average country, while Italy is very much an usual on this respect. The model interprets this as indicating that the return to educational investment must be perceived as larger in the US, which is what happens if the precision of merit signals is much larger in the US. Actually, that is the differentiating aspect of the US: it invests a lot in education (see in table 4).

Thus, two countries which are very much close in the Great Gatsby curve may have arrived to this point via very different routes. In particular, the model suggests that if they differ in the educational investment is a sign that the one that invests the most is likely to have

²⁸Notice again that μ or σ only adjust the mean and variance of λ across countries between model and data, but their value does not affect the correlations, and in this respect they are there for aesthetic value. In any case, we perform the exercise also using spending in education directly (without any transformation) and obtain the same qualitative results.

more accurate objective information on the ability of workers, larger reward to merit and less incidence of background *conditional* on merit. This generates the larger demand of educational investment.

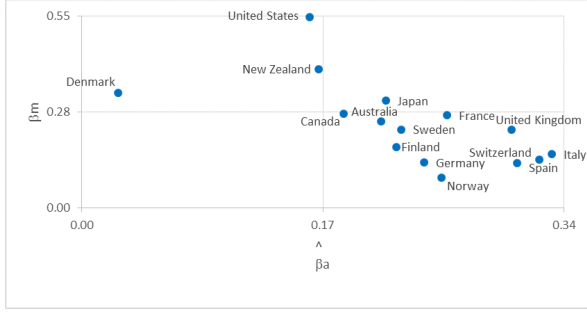
- A country may have much less mobility than another, and in spite of it the effect of background be much larger. This is very clear when comparing the US and Denmark in the data. The interesting bit is that Denmark is the only country with a larger educational effort than the US. In spite of it the model imputes Denmark less precision in both accounts signals, as it is the only way of making it compatible with the low inequality and mobility observed.

In figure 6 we present the result of the same exercise performed under alternative data configurations. As it is clear, qualitatively the same results arise. Albeit the behavior of the US it is somewhere less radical, it is always among the countries with the highest precision on the merit signal, and its largest reward. With respect to Italy, the US puts always considerably more value on merit, and there is less prevalence of background conditional on merit (albeit with after taxes and private education data only the differences are not huge, they always go in the right direction). With respect to Denmark the US has always larger values of both precisions, and larger rewards on both signals.

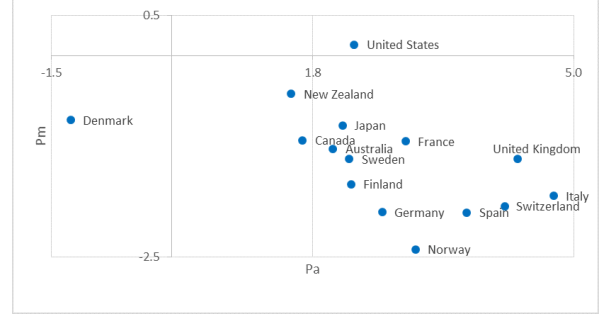
The main point of this section is that when confronted with data, the suggestions of the paper seem at least plausible. A naive interpretation of intergenerational mobility as degree of “meritocracy” may be a dangerous oversimplification. The US might be still a land of opportunities *conditional* to have the right ability. Of course, those abilities are to a large extent inherited, and that is the reason why the intergenerational persistence of income is so large. But notice that the reason why they are so inheritable might precisely be that merit is highly rewarded; that in a deep sense, the US is a very meritocratic society.

8 Summary, Conclusion and Further Research

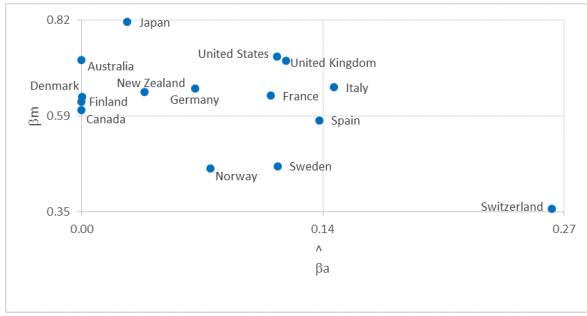
In this paper we have argued that it is reasonable to understand merit as rewards to objective measures of productive ability, as opposed to rewarding inferences made on an agents productivity based on her background. Perhaps our main contribution is to articulate that a low degree of intergenerational mobility is not at all incompatible with a large prevalence of merit in a society. Very much the opposite.



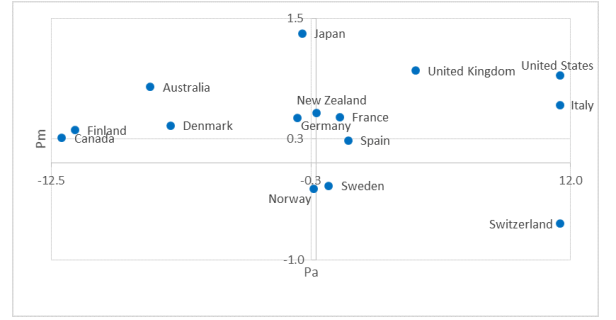
(a) $\beta_m, \hat{\beta}_a$. After taxes. Total education.



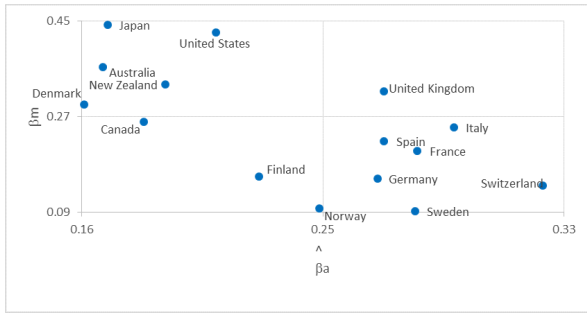
(b) P_m, P_a . After taxes. Total education.



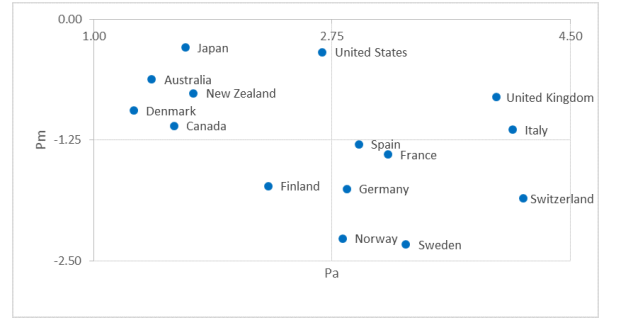
(c) $\beta_m, \hat{\beta}_a$. Before taxes. Private Education.



(d) P_m, P_a . Before Taxes. Private Education.



(e) $\beta_m, \hat{\beta}_a$. After taxes. Private education.



(f) P_m, P_a . After taxes. Private education.

Figure 6: Implied Rewards and Precisions across countries for other data configurations.

We have presented a model of statistical discrimination. On one hand families take a rational decision on consumption and educational investment in the absence of capital markets. As a consequence richer families invest more resources in the education of their children, who in average become more productive than the children of poor families.

Then firms need to determine the wage of workers whose productivity is observed with noise. Consequently, firms use any information on the background of agents in order to help infer their productivity. Two individuals that provide the market with the same observable productivity will not receive the same wage. The one with richer parents will be rewarded more generously.

We show that, not surprisingly, the more accurate is the information on the background of individuals, the easier it is to discriminate, the larger the reward of having a good background, and the smaller the reward of merit. All of it leading to more inequality and less intergenerational mobility.

Perhaps more surprisingly, we also show the accuracy in the direct signal on productivity has qualitatively the same effects. More precise information on productivity leads to having a larger reward to merit, but also to a society that not only is more unequal society, but where the income of a child is more associated with the one of her parent.

Providing with information (be it on agents productivity or on their background), allows society to differentiate between agents, paying more to those deemed to be more productive. This increase of inequality feeds back to itself, because heterogeneity among the parents increases the ex-ante uncertainty that firms have on the human capital of their children, which itself leads to more use of the information that differentiates people (the signals), which leads to further inequality, and so on...

Notice that inheritance is ingrained in this process. Things hinge in the fact that more inequality in the income of the parents translates into more inequality in the human capital of their children. Because it does, not only more access of information feeds up inequality, but also the linkages of income across generations. Both things (inequality and lack of mobility) positively correlate with each other, implying a Great Gatsby Curve, but they could be caused by either a prevalence of merit or a prevalence of inherited advantages.

Nevertheless, there is an important aspect in which the effects of more information on merit differ qualitatively from the effects of more information on parental background. The first encourages human capital accumulation much more. Rewards of merit directly determines the return of human capital accumulation, because in the model you do not want to educate your

children because it is rewarding in itself (there is no warm glove motive), You want to educate your children in the measure that it makes them look productive, and thus makes them rich. For that to be the case, it is necessary that firms value their talent itself.

More accurate information on background has a much milder effect on the incentives to invest. On one hand has a second order incentive effect on investment because your grandchildren (not your children) will enjoy a larger income as a consequence of your kids will getting richer if you invest more on them. But on the other hand, the rewards to merit will undoubtedly decrease as consequence of it, decreasing the link between your investment and your children income.

This difference allow us to try to do the exercise of calibrating the model to the available data, first across countries US commuting zones, and then across countries. We do not perform this exercise with the aim of testing the model, and we do not take our results as literal measurement of the rewards to merit and background across societies. With this exercise we make the point that:

- The measures that we obtain from the model seem to have reasonable correlations with characteristics of the society. Thus, across US commuting Zones the prevalence of merit seems to correlate well with having a large service sector and large foreign communities, while the fraction of African-Americans and racial segregation correlate well with the prevalence of inherited advantages.
- It may very well be the case that in societies where parental income has a very large unconditional effect on the child income, this same parental income may have a very small effect when controlling for the abilities of the child. Thus, the US shows very little mobility when compared with most other developed countries, while the inferred the prevalence of merit is among the largest, if not the largest.
- Two societies that look similar in the Great Gastby Curve, may look very different in the space defined by the rewards to merit and background, and this may be reflected in the effort that they make on educational investment. For instance, the US and Italy look very much the same in what respects to their levels of inequality and mobility, but the US invests much more in education. Consequently, the model suggests that merit is much more prevalent in the US than in Italy.
- A society may have much more prevalence of merit, and less rewards to advantages, than another that is much more equal and that has much more intergenerational mobility. For

instance, the implied prevalence of merit in the US is larger than in Denmark, albeit it is much more unequal and less mobile. This is interesting because Denmark is the country that has more investment in education (both public and private) in the sample. The reason for getting this ranking is that in order to be provide with low inequality and mobility to Denmark, the model can not allow her to be very meritocratic. In the model, and when we simulate it with reasonable numbers, a consequence of having accurate information on the productivity of people is to be in the sad end of the Great Gatsby Curve.

8.1 Caveats and further research

It is important to remark things that we don't do. Our exercise is not to explain the great Gatsby curve. It is almost evident that institutional differences in redistribution and public education spending must be an important part of it. Our measures completely abstract from it, as we have not modeled them. Our point, thus, is just to show that the model suggests that the rankings of merit *could* differ, even very substantially, from the rankings of intergenerational mobility. That is precisely what Michael Young had in mind when he invented the word "*meritocracy*". Moreover, doing a simple (but informed) simulation of the model suggests that in vary radical way, by placing the US in the top. Our point is to indicate that these effects have the potentiality of being large, and to have them in mind might be a useful tool for understanding current events.

There are several issues that should address in the future. (1) First of all, and most obviously, we should model also redistribution and government education. (2) We should also try to find direct empirical measures of merit, and correlate it with our model-based inference. (3) We should model capital markets in more inclusive way, and allow for different information structures at the time of investment.

Public education could have the effect of disconnecting one's productivity from their parental income under two conditions: If it was really equal for all and if it substituted away completely from private educational decisions. If both were literally true, there would be no connection between parental income and a person human capital, which is one of the fundamental building blocks of our model. There would be effectively no inheritance due to the economic position of the parents.

Nevertheless, it is most unlikely that both held true in such a literal sense. First of all, as we have argued, investment in education far exceeds the income spend in formal schooling. Externalities in human capital accumulation suggest that people will segregate according to

income.²⁹ Richer people is likely to pay more in order to live near rich people, and to make sure that their children interact with other a superior station of life than their own, resulting in a segregated distribution of population and neighborhood gentrification. So, even in the presence of generous public education the same effects that we have model may take place. Moreover, the provision of public education does not preclude the existence of a large degree of private educational investment. Some of it in a very formal sense, like paying private schools that provide better education, or paying private tutoring. Some of it in the form of more time, and more time of highly qualified individuals, devoted to improving the future productive ability of children. This could be teaching them to play violin, but also reading them books at night, or doing homework, or having motivating conversations with them, and so on. Thus, even in societies that could prohibit formal private education, it is unlikely to imagine that the better of will not invest more in their children.³⁰ We are quite confident that the kind of effects we have talked about would survive in a model with formal education. The issue is quantitative, not qualitative. The inclusion of formal public education would allow us to put more trust in the calibration exercise.

Somehow the same happens with redistribution. Unless we were talking of a Maoist regime that would cancel any ex-post heterogeneity, the same effects that we are talking about would survive, albeit their magnitude could be very different.

In this paper we have abstracted from these most obvious issues because including them would preclude us from providing analytical solutions. We believe that that is important, as this is an eminently theoretical paper that aims to stress a channel that has been sidelined by the literature.

In any case, to address these issues further research should provide with a model allowing for these effects, in order to provide with a more sensible (while still rough) calibration. Likewise another issue that should be addressed in further research is the inclusion of more realistic capital markets, as it is clear that student credit is a first order related issue and parents do leave inheritance to their children in other forms than human capital investment.

In principle it would be possible to study all those effects in a structural model, and perform hypothetical exercises, measuring for instance the effects of different policies, and the relative relevance of the mechanism that we exposed. While it is not impossible to include all these

²⁹See for instance the seminal paper Benabou (1993), and the large literature thereafter.

³⁰And all this assuming that public education benefits the poor more than the rich, which could not be the case in some circumstances. Checchi et al. (1999) for instance argue that subsidies in the tertiary educational system decrease mobility in Italy when compared to the US precisely because they are in practice a subsidy to the rich

effects in an encompassing model, it is an extremely difficult task. We will certainly explore the possibility of doing so, but there is another possible way of dealing at least with the relevance of our mechanism.

We stress the difference between the unconditional effect of parental income, and the effect of parental income once conditioning for individual characteristics. That is something that in principle it should be possible to do. It is difficult to compare across countries with identical measures, but perhaps there are ways to deal with the problems. We plan to pursue that line of research.

In the model we do allow for the inheritance of talent itself.

Two theoretical caveats... inheritance of talent and information of children ability at the time of investment

There are, of course, many caveats to this story.
and possible extensions.

- Inheritance of talent, genetic or otherwise
 - Assortative Mating
- Grammar schools. parents invest more to get their children in the right track.
- Not a caveat, but a possible add on, conspicuous consumption.
- Also that the signal on parental income could be educational investment... conspicuous education.
- Look at measures across countries of conditional and unconditional effect of parents... perhaps using educational attainment. Empirical task.

References

- Abbott, B., & Gallipoli, G. (forthcoming). Human capital spill-overs and the geography of intergenerational mobility. *Review of Economic Dynamics*.
- Aitchison, J., & Brown, J. A. C. (1963). *The Lognormal Distribution*. Cambridge University Press.
- Arrow, K. J. (1973). The theory of discrimination. In O. Ashenfelter, & A. Rees (Eds.) *Discrimination in Labor Markets*, chap. 1. Princeton University Press.
- Becker, G. S. (1957). *The Economics of Discrimination*. University of Chicago Press, 1st ed.
- Becker, G. S., & Tomes, N. (1979). An Equilibrium Theory of the Distribution of Income and Intergenerational Mobility. *Journal of Political Economy*, 87(6), 1153–1189.
- Becker, G. S., & Tomes, N. (1986). Human Capital and the Rise and Fall of Families. *Journal of Labor Economics*, 4(3), S1–S39.
- Benabou, R. (1993). Workings of a City: Location, Education, and Production. *The Quarterly Journal of Economics*, 108(3).
- Checchi, D., Ichino, A., & Rustichini, A. (1999). More equal but less mobile? Education financing and intergenerational mobility in Italy and in the US. *Journal of Public Economics*, 74, 351–393.
- Chetty, R., Hendren, N., Kline, P., & Saez, E. (2014). Where is the Land of Opportunity? The Geography of Intergenerational Mobility in the United States. *The Quarterly Journal of Economics*, 129(4), 1553–1623.
- Coate, S., & Loury, G. (1993). Will Affirmative Action Policies Eliminate Negative Stereotypes? *American Economic Review*, 83(5), 1220–1240.
- Corak, M. (2013). Inequality from Generation to Generation: The United States in Comparison. In R. Rycroft (Ed.) *The Economics of Inequality, Poverty, and Discrimination in the 21st Century*. ABC-CLIO.
- Galor, O., & Zeira, J. (1993). Income distribution and macroeconomics. *The Review of Economic Studies*, 60, 355–392.

- Güell, M., Pellizzari, M., Pica, G., & Rodríguez Mora, J. V. (2015). Correlating Social Mobility and Economic Outcomes. *CEPR Discussion Papers*, 10496.
- Hassler, J., Rodríguez Mora, J. V., & Zeira, J. (2007). Inequality and Mobility. *Journal of Economic Growth*, 12, 235–259.
- Krueger, A. B. (2012). The Rise and Consequences of Inequality in the United States. Speech to the Center for American Progress.
- Moro, A., & Norman, P. (2004). A General Equilibrium Model of Statistical Discrimination. *Journal of Economic Theory*, 114, 1–30.
- Norman, P. (2003). Statistical discrimination and efficiency. *Review of Economic Studies*, 70, 615–627.
- Solon, G. (2004). A Model of Intergenerational Mobility Variation over Time and Place. In M. Corak (Ed.) *Generational Income Mobility in North America and Europe*, chap. 2, (pp. 38–47). Cambridge University Press.
- Young, M. (1958). *The Rise of Meritocracy, 1870 - 2033*. London: Thames and Hudson.

A Proof of Result 1

Proof. Solving the program:

$$W(Y_t^i) = \max_{X_t^i} \left\{ \ln[Y_t^i - X_t^i] + \frac{1}{1+\delta} EW \left(e^{\gamma_0} (Y_t^i)^{\gamma_1} (X_t^i)^{\gamma_2} e^{\varepsilon_t^i} \right) \right\} \quad (37)$$

First we prove that parents invest a fixed percentage of their income in their children:

$$X_t^i = \lambda Y_t^i \quad (38)$$

To do we guess

$$W(Y_t^i) = A + B \ln Y_t^i \quad (39)$$

which we will later verify. The Euler equation is:

$$\frac{1}{Y_t^i - X_t^i} = \frac{1}{1+\delta} \frac{\partial EW \left(e^{\gamma_0} (Y_t^i)^{\gamma_1} (X_t^i)^{\gamma_2} e^{\varepsilon_t^i} \right)}{\partial X_t^i} \quad (40)$$

$$EW \left(e^{\gamma_0} (Y_t^i)^{\gamma_1} (X_t^i)^{\gamma_2} e^{\varepsilon_t^i} \right) = A + B [\bar{\varepsilon} + \gamma_0 + \gamma_1 \ln Y_t^i + \gamma_2 \ln X_t^i] \quad (41)$$

So, the Euler becomes:

$$\frac{1}{Y_t^i - X_t^i} = \frac{1}{1+\delta} B \gamma_2 \frac{1}{X_t^i} \quad (42)$$

Implied:

$$X_t^i = \frac{B \gamma_2}{(1+\delta) + B \gamma_2} Y_t^i \quad (43)$$

and

$$C_t^i = \frac{(1+\delta)}{(1+\delta) + B \gamma_2} Y_t^i \quad (44)$$

Now, substituting X_t^i into the expectation we get:

$$EW(Y_{t+1}^i | Y_t^i, X_t^i) = EW \left(e^{\gamma_0} (Y_t^i)^{\gamma_1} (X_t^i)^{\gamma_2} e^{\varepsilon_t^i} \right) \quad (45)$$

$$= A + B \left[\bar{\varepsilon} + \gamma_0 + (\gamma_1 + \gamma_2) \ln Y_t^i + \gamma_2 \ln \frac{B \gamma_2}{(1+\delta) + B \gamma_2} \right] \quad (46)$$

And the value function will be:

$$W(Y_t^i) = \ln \frac{(1+\delta)}{(1+\delta) + B \gamma_2} + \ln Y_t^i + \frac{A}{(1+\delta)} \quad (47)$$

$$+ \frac{B}{(1+\delta)} \left[\bar{\varepsilon} + \gamma_0 + (\gamma_1 + \gamma_2) \ln Y_t^i + \gamma_2 \ln \frac{B \gamma_2}{(1+\delta) + B \gamma_2} \right] \quad (48)$$

So, if the guess is right:

$$A = \ln \frac{(1+\delta)}{(1+\delta) + B \gamma_2} + \frac{A}{(1+\delta)} + \frac{B}{(1+\delta)} \left[\bar{\varepsilon} + \gamma_0 + \gamma_2 \ln \frac{B \gamma_2}{(1+\delta) + B \gamma_2} \right] \quad (49)$$

and

$$B = \frac{B}{(1+\delta)} (\gamma_1 + \gamma_2) + 1 \quad (50)$$

Solving for B

$$B = \frac{(1+\delta)}{(1+\delta) - (\gamma_1 + \gamma_2)} \quad (51)$$

which should be positive. We will show that the equilibrium values of γ_1 and γ_2 are always such that this happens. Finally, solving for A

$$\delta A = (1+\delta) \ln \frac{(1+\delta)}{(1+\delta) + B \gamma_2} + \frac{(1+\delta)}{(1+\delta) - (\gamma_1 + \gamma_2)} \left[\bar{\varepsilon} + \gamma_0 + \gamma_2 \ln \frac{B \gamma_2}{(1+\delta) + B \gamma_2} \right] \quad (52)$$

It is useful to notice that

$$1 - \lambda = \frac{(1 + \delta)}{(1 + \delta) + B\gamma_2} = \frac{(1 + \delta) - (\gamma_1 + \gamma_2)}{(1 + \delta) - \gamma_1} = 1 - \frac{\gamma_2}{[(1 + \delta) - \gamma_1]} \quad (53)$$

so:

$$\lambda = \frac{\gamma_2}{[(1 + \delta) - \gamma_1]} \quad (54)$$

$$\delta A = (1 + \delta) \ln \left[1 - \frac{\gamma_2}{[(1 + \delta) - \gamma_1]} \right] + \frac{(1 + \delta)}{(1 + \delta) - (\gamma_1 + \gamma_2)} \left[\bar{\varepsilon} + \gamma_0 + \gamma_2 \ln \frac{\gamma_2}{[(1 + \delta) - \gamma_1]} \right] \quad (55)$$

$$A = \frac{\bar{\varepsilon} + \gamma_0 + \ln [(1 + \delta) - (\gamma_1 + \gamma_2)]^{[(1 + \delta) - (\gamma_1 + \gamma_2)]} + \ln \frac{\gamma_2^{\gamma_2}}{[(1 + \delta) - \gamma_1]^{[(1 + \delta) - \gamma_1]}}}{[(1 + \delta) - (\gamma_1 + \gamma_2)] \frac{\delta}{1 + \delta}} \quad (56)$$

□

B Proof of Result 2

Proof. This follows directly from the human capital accumulation equation 6 and the investment rule, $X_t^i = \lambda Y_t^i$. □

C Proof of Result 3

Proof. Since $(h_{t+1}^i, a_{t+1}^i, m_{t+1}^i)$ has a multivariate normal distribution, we can appeal to the conditional normal p.d.f. to find the distribution of $h_{t+1}^i | \Omega_{t+1}^i$.

Let $\mathbf{X}$ be a partitioned multivariate normal random vector with $\mathbf{X}^T = [\mathbf{X}_1 \quad \mathbf{X}_2]$, where $\mathbf{X}_1 = [h_{t+1}^i]$ and $\mathbf{X}_2^T = [a_{t+1}^i \quad m_{t+1}^i]$. The mean of $\mathbf{X}$ is given by:

$$\mu = \begin{bmatrix} \mu_1 \\ \mu_2 \end{bmatrix} \quad (57)$$

where $\mu_1 = [\mu_h]$ and $\mu_2^T = [\mu_y \quad \mu_m]$.

The variance-covariance matrix of $\mathbf{X}$ is given by:

$$\Sigma = \begin{bmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{21} & \Sigma_{22} \end{bmatrix} \quad (58)$$

where:

$$\Sigma_{11} = [V_h] \quad (59)$$

$$\Sigma_{12} = [\alpha V_y \quad V_h] = \Sigma_{21}^T \quad (60)$$

$$\Sigma_{22} = \begin{bmatrix} V_y + V_a & \alpha V_y \\ \alpha V_y & V_h + V_m \end{bmatrix} \quad (61)$$

Then the distribution of h_{t+1}^i conditional on Ω_{t+1}^i is univariate normal with the following mean and variance:

$$E(h_{t+1}^i | \Omega_{t+1}^i) \quad (62)$$

$$= \mu_1 + \Sigma_{12} \Sigma_{22}^{-1} (\mathbf{x}_2 - \mu_2) \quad (63)$$

$$= \mu_h + \frac{1}{(V_h + V_m)(V_y + V_a) - \alpha^2 V_y^2} \begin{bmatrix} \alpha V_y & V_h \end{bmatrix} \begin{bmatrix} V_h + V_m & -\alpha V_y \\ -\alpha V_y & V_y + V_a \end{bmatrix} \begin{bmatrix} a_{t+1}^i - \mu_y \\ m_{t+1}^i - \mu_h \end{bmatrix} \quad (64)$$

$$= \mu_h + \alpha \beta_a (1 - \beta_m) [a_{t+1}^i - \mu_y] + \beta_m [m_{t+1}^i - \mu_h] \quad (65)$$

$$= \beta_m m_{t+1}^i + (1 - \beta_m) \left[\ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} + \alpha \beta_a a_{t+1}^i + \alpha (1 - \beta_a) \mu_y \right] \quad (66)$$

and,

$$Var(h_{t+1}^i | \Omega_{t+1}^i) \quad (67)$$

$$= \Sigma_{11} - \Sigma_{12} \Sigma_{22}^{-1} \Sigma_{21} \quad (68)$$

$$= V_h - \frac{1}{(V_h + V_m)(V_y + V_a) - \alpha^2 V_y^2} \begin{bmatrix} \alpha V_y & V_h \end{bmatrix} \begin{bmatrix} V_h + V_m & -\alpha V_y \\ -\alpha V_y & V_y + V_a \end{bmatrix} \begin{bmatrix} \alpha V_y \\ V_h \end{bmatrix} \quad (69)$$

$$= V_h - \alpha^2 V_y \beta_a (1 - \beta_m) - \beta_m V_h \quad (70)$$

$$= \alpha^2 V_y (1 - \beta_m) (1 - \beta_a) + (1 - \beta_m) V_\omega \quad (71)$$

$$= \beta_m V_m \quad (72)$$

□

D Proof of Result 4

Proof. Given the conditional distribution of the log of human capital is normal, the conditional distribution of human capital is log normal with mean:

$$Y_{t+1}^i = E(\exp\{h_{t+1}^i\} | \Omega_{t+1}^i) = \exp\left\{\mu_{h_{t+1}^i | \Omega_{t+1}^i}\right\} \exp\left\{\frac{V_{h_{t+1}^i | \Omega_{t+1}^i}}{2}\right\} \quad (73)$$

Substituting and taken logarithms produces expression 18.

□

E The Multiplier Effect

This appendix proves the existence of the multiplier effect described in section 5.2 and calculates its magnitude.

Proof. Imposing the steady state condition, $V_{y_t} = V_{y_{t-1}} = V_y$ on the law of motion of the variance of log income given in equation 20, gives:

$$V_y = \alpha^2 [\beta_a (1 - \beta_m) + \beta_m] V_y + \beta_m V_\omega \quad (74)$$

Solving this gives the equation for steady state variance of log income given in result 5. We will call the right-hand side of equation 74: $\Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)$.

We can then examine two things: the shift in Ψ from a change in one of the exogenous parameters, keeping V_y fixed; and the total derivative of V_y with respect to the same parameter. The ratio of the latter to the former is the multiplier.

Take V_a as an example. For a given V_y , the effect of a change in V_a on the Ψ function is given by:

$$\frac{\partial \Psi}{\partial V_a} = \frac{\partial \Psi}{\partial \beta_a} \frac{\partial \beta_a}{\partial V_a} + \frac{\partial \Psi}{\partial \beta_m} \left[\frac{\partial \beta_m}{\partial V_a} + \frac{\partial \beta_m}{\partial \beta_a} \frac{\partial \beta_a}{\partial V_a} \right] \quad (75)$$

This is the *partial* effect for a fixed V_y . Allowing V_y to fully adjust however, we find the effect of a change in V_a to be:

$$\begin{aligned}
\frac{dV_y}{dV_a} &= \frac{\partial \Psi}{\partial V_y} \frac{dV_y}{dV_a} + \frac{\partial \Psi}{\partial \beta_a} \frac{d\beta_a}{dV_a} + \frac{\partial \Psi}{\partial \beta_m} \frac{d\beta_m}{dV_a} \\
&= \left[\frac{\partial \Psi}{\partial V_y} + \left[\frac{\partial \Psi}{\partial \beta_a} + \frac{\partial \Psi}{\partial \beta_m} \frac{\partial \beta_m}{\partial \beta_a} \right] \frac{\partial \beta_a}{\partial V_y} \right] \frac{dV_y}{dV_a} + \frac{\partial \Psi}{\partial \beta_a} \frac{\partial \beta_a}{\partial V_a} + \frac{\partial \Psi}{\partial \beta_m} \left[\frac{\partial \beta_m}{\partial V_a} + \frac{\partial \beta_m}{\partial \beta_a} \frac{\partial \beta_a}{\partial V_a} \right] \\
&= \frac{1}{1 - \left[\frac{\partial \Psi}{\partial V_y} + \left[\frac{\partial \Psi}{\partial \beta_a} + \frac{\partial \Psi}{\partial \beta_m} \frac{\partial \beta_m}{\partial \beta_a} \right] \frac{\partial \beta_a}{\partial V_y} \right]} \left\{ \left[\frac{\partial \Psi}{\partial \beta_a} + \frac{\partial \Psi}{\partial \beta_m} \frac{\partial \beta_m}{\partial \beta_a} \right] \frac{\partial \beta_a}{\partial V_a} + \frac{\partial \Psi}{\partial \beta_m} \frac{\partial \beta_m}{\partial V_a} \right\} \quad (76)
\end{aligned}$$

The first term on the right-hand side (i.e the fraction) is the multiplier term. The second term is the shift in the Ψ curve calculated in 75. If the multiplier is greater than 1, it tells us that the total (or long-run) effect on V_y of a change in V_a is larger than the partial effect when V_y is fixed (or the short-run effect, since $V_{y,t-1}$ is fixed).

By substitution, the multiplier term is:

$$\frac{1}{1 - \alpha^2 \left[1 - (1 - \beta_a)^2 (1 - \beta_m)^2 \right]} > 1 \quad (77)$$

Thus any change in V_a is amplified since the additional income variance which it creates feeds into next period's income variance and the levels of discrimination which individuals in that generation are subjected to.

We can carry out exactly the same exercise for V_m , V_ω and α . Although the partial effect of a change in each of the parameters differs, the multiplier is the same in each case. In all cases it is given by equation 77. $\square$

F Proof of Result 5

Proof. $\Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)$ is defined as in appendix E. It follows that:

$$\frac{d\Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)}{dV_y} \quad (78)$$

$$= \frac{\partial \Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)}{\partial V_y} + \frac{\partial \Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)}{\partial \beta_a} \frac{d\beta_a}{dV_y} + \frac{\partial \Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)}{\partial \beta_m} \frac{d\beta_m}{dV_y} \quad (79)$$

$$= \alpha^2 [\beta_a(1 - \beta_m) + \beta_m] + \alpha^2(1 - \beta_m)V_y \frac{d\beta_a}{dV_y} + [\alpha^2 \beta_a V_a + V_\omega] \frac{d\beta_m}{dV_y} \quad (80)$$

Since,

$$\frac{d\beta_a}{dV_y} = \frac{1 - \beta_a}{V_y + V_a} > 0 \quad (81)$$

and,

$$\frac{d\beta_m}{dV_y} = \alpha^2 \frac{(1 - \beta_m)(1 - \beta_a)^2}{\alpha^2 \beta_a V_a + V_\omega + V_m} > 0 \quad (82)$$

it follows that,

$$\frac{d\Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)}{dV_y} = \alpha^2 [\beta_a(1 - \beta_m) + \beta_m] \{1 + (1 - \beta_m)(1 - \beta_a)\} \quad (83)$$

$$= \alpha \rho [1 + (1 - \beta_m)(1 - \beta_a)] > 0 \quad (84)$$

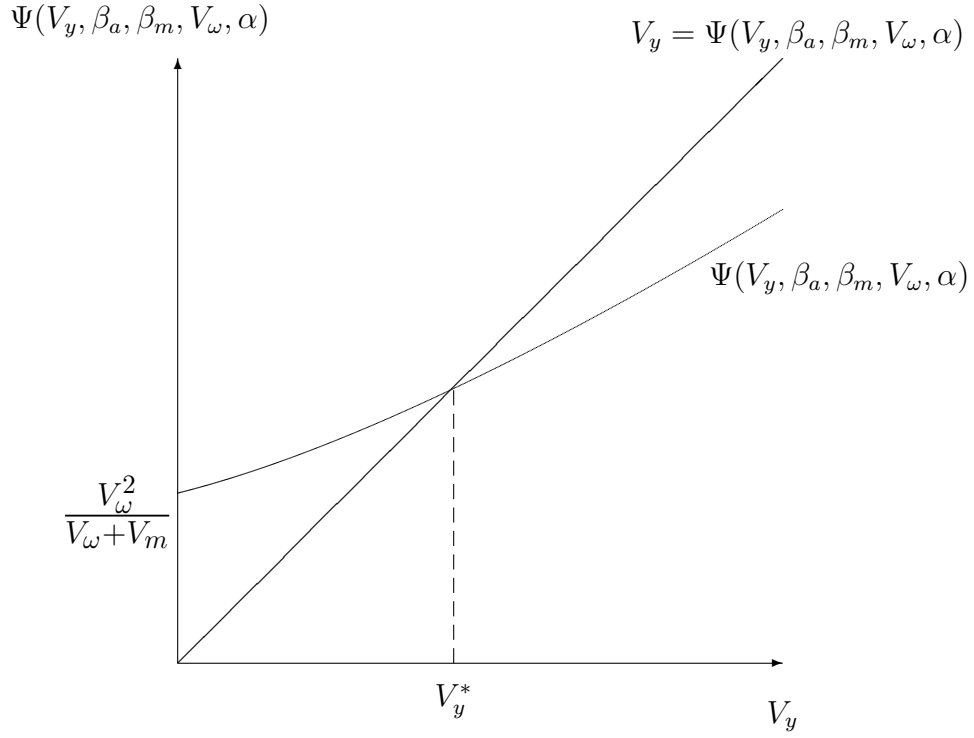


Figure 7: Law of motion of V_y

while the second derivative is:

$$\frac{d^2 \Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)}{dV_y^2} = \frac{\partial \left[\frac{d\Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)}{dV_y} \right]}{\partial \beta_a} \frac{d\beta_a}{dV_y} + \frac{\partial \left[\frac{d\Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)}{dV_y} \right]}{\partial \beta_m} \frac{d\beta_m}{dV_y} \quad (85)$$

$$= 2\alpha^2 (1 - \beta_a) (1 - \beta_m) \left[(1 - \beta_m) \frac{d\beta_a}{dV_y} + (1 - \beta_a) \frac{d\beta_m}{dV_y} \right] > 0 \quad (86)$$

We also know that:

$$\Psi(0, \beta_a, \beta_m, V_\omega, \alpha) = \frac{V_\omega^2}{V_\omega + V_m} \quad (87)$$

$$\Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha) \rightarrow \infty \text{ as } V_y \rightarrow \infty \quad (88)$$

$$\frac{d\Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)}{dV_y} \rightarrow \alpha^2 \text{ as } V_y \rightarrow \infty \quad (89)$$

Since $\Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)$ starts above the 45 degree line, is upward sloping and convex in V_y , and has a maximum slope of $\alpha^2 < 1$, it must cut the 45 degree line once from above. This gives the unique, stable, steady state value of V_y .

Note that the shape of the curve implies there will be a multiplier effect from changes in the parameters. Any parameter change which causes a shift in $\Psi(V_y, \beta_a, \beta_m, V_\omega, \alpha)$ will lead to a larger change in V_y . This multiplier effect will be discussed further when we consider the comparative statics of the model.

To find the steady state value of the mean of log income, we take the mean of equation 19 and apply the steady state condition $\mu_{y_t} = \mu_{y_{t+1}} = \mu_y$ to get:

$$\mu_y = \ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} + \frac{\beta_m V_m}{2} + \alpha \mu_y \quad (90)$$

This solves to give the equation for mean log income.

□

G Proof of Result 6

Proof. Notice:

$$\frac{d\beta_a}{dV_a} = \frac{\partial\beta_a}{\partial V_y} \frac{dV_y}{dV_a} + \frac{\partial\beta_a}{\partial V_a} = \frac{1}{V_y + V_a} \left[(1 - \beta_a) \frac{dV_y}{dV_a} - \beta_a \right] \quad (91)$$

and

$$\frac{d\beta_m}{dV_a} = \frac{\partial\beta_m}{\partial\beta_a} \frac{d\beta_a}{dV_a} + \frac{\partial\beta_m}{\partial V_a} = \frac{\alpha^2(1 - \beta_m)}{\alpha^2\beta_a V_a + V_\omega + V_m} \left[(1 - \beta_a)^2 \frac{dV_y}{dV_a} + \beta_a^2 \right] \quad (92)$$

So, after some algebra

$$\frac{dV_y}{dV_a} = \frac{-\alpha^2\beta_a^2(1 - \beta_m)^2}{1 - \alpha\rho [1 + (1 - \beta_m)(1 - \beta_a)]} \quad (93)$$

Note also that:

$$1 - \alpha\rho [1 + (1 - \beta_m)(1 - \beta_a)] = 1 - \alpha^2 [1 - (1 - \beta_m)^2(1 - \beta_a)^2] > 0 \quad (94)$$

Therefore, the effect of a change in V_a on V_y is:

$$\frac{dV_y}{dV_a} = \frac{-\alpha^2\beta_a^2(1 - \beta_m)^2}{1 - \alpha^2 [1 - (1 - \beta_m)^2(1 - \beta_a)^2]} < 0 \quad (95)$$

The effects of a change in V_a on β_a and β_m are:

$$\frac{d\beta_a}{dV_a} = -\frac{\beta_a(1 - \beta_a)}{V_a} \left[1 + \frac{\alpha^2\beta_a(1 - \beta_a)(1 - \beta_m)^2}{(1 - \alpha^2) + \alpha^2(1 - \beta_a)^2(1 - \beta_m)^2} \right] < 0 \quad (96)$$

$$= -\frac{\beta_a(1 - \beta_a)}{V_a} \left[\frac{1 - \alpha^2 [1 - (1 - \beta_a)(1 - \beta_m)^2]}{1 - \alpha^2 [1 - (1 - \beta_m)^2(1 - \beta_a)^2]} \right] \quad (97)$$

$$\frac{d\beta_m}{dV_a} = \frac{\alpha^2\beta_a^2(1 - \beta_m)^2}{V_m} \left[\frac{1 - \alpha^2}{1 - \alpha^2 [1 - (1 - \beta_m)^2(1 - \beta_a)^2]} \right] > 0 \quad (98)$$

The extent to which firms value and use the signal on background is not measured by β_a but by $\hat{\beta}_a = \alpha\beta_a(1 - \beta_m)$. A change in the precision of the advantage signal has the following effect on $\hat{\beta}_a$:

$$\frac{d\hat{\beta}_a}{dV_a} = \alpha(1 - \beta_m) \frac{d\beta_a}{dV_a} - \alpha\beta_a \frac{d\beta_m}{dV_a} < 0 \quad (99)$$

The effect of a change in V_a on ρ is:

$$\frac{d\rho}{dV_a} = \alpha(1 - \beta_m) \frac{d\beta_a}{dV_a} + \alpha(1 - \beta_a) \frac{d\beta_m}{dV_a} \quad (100)$$

$$= -\frac{\alpha\beta_a(1 - \beta_a)(1 - \beta_m)^2}{V_a V_m [1 - \alpha^2 [1 - (1 - \beta_m)^2(1 - \beta_a)^2]]} \quad (101)$$

$$\cdot \{ (1 - \alpha^2)(V_\omega + V_m) + \alpha^2(1 - \beta_a)(1 - \beta_m)V_m \} < 0 \quad (102)$$

□

H Proof of Result 7

Proof. Notice that:

$$\frac{d\beta_a}{dV_m} = \frac{\partial\beta_a}{\partial V_y} \frac{dV_y}{dV_m} = \frac{(1 - \beta_a)}{V_y + V_a} \frac{dV_y}{dV_m} \quad (103)$$

and

$$\frac{d\beta_m}{dV_m} = \frac{\partial\beta_m}{\partial\beta_a} \frac{d\beta_a}{dV_m} + \frac{\partial\beta_m}{\partial V_m} \quad (104)$$

$$= \frac{\alpha^2(1-\beta_a)^2(1-\beta_m)}{\alpha^2\beta_a V_a + V_\omega + V_m} \frac{dV_y}{dV_m} - \frac{\beta_m}{\alpha^2\beta_a V_a + V_\omega + V_m} \quad (105)$$

Then,

$$\frac{dV_y}{dV_m} = \frac{-\beta_m^2}{1 - \alpha\rho [1 + (1 - \beta_m)(1 - \beta_a)]} \quad (106)$$

As above:

$$1 - \alpha\rho [1 + (1 - \beta_m)(1 - \beta_a)] = 1 - \alpha^2 [1 - (1 - \beta_m)^2(1 - \beta_a)^2] \quad (107)$$

Therefore, the effect of a change in V_m on V_y is:

$$\frac{dV_y}{dV_m} = -\frac{\beta_m^2}{1 - \alpha^2 [1 - (1 - \beta_m)^2(1 - \beta_a)^2]} < 0 \quad (108)$$

The effects of a change in V_m on β_a and β_m are:

$$\frac{d\beta_a}{dV_m} = -\frac{(1 - \beta_a)^2}{V_a} \left[\frac{\beta_m^2}{1 - \alpha^2 [1 - (1 - \beta_m)^2(1 - \beta_a)^2]} \right] < 0 \quad (109)$$

$$\frac{d\beta_m}{dV_m} = -\frac{\beta_m(1 - \beta_m)}{V_m} \left[1 + \frac{\alpha^2\beta_m(1 - \beta_m)(1 - \beta_a)^2}{1 - \alpha^2 [1 - (1 - \beta_m)^2(1 - \beta_a)^2]} \right] < 0 \quad (110)$$

The effect of a change in V_m on ρ is:

$$\frac{d\rho}{dV_m} = \alpha(1 - \beta_m) \frac{d\beta_a}{dV_m} + \alpha(1 - \beta_a) \frac{d\beta_m}{dV_m} \quad (111)$$

$$= -\frac{\alpha\beta_m(1 - \beta_a)(1 - \beta_m)}{V_a V_m} \left\{ V_a + \frac{(1 - \beta_a)\beta_m V_m}{[1 - \alpha^2 [1 - (1 - \beta_m)^2(1 - \beta_a)^2]]} \right\} < 0 \quad (112)$$

Now turning to the effect on $\hat{\beta}_a$. A change in the precision of the ability signal has the following effect on $\hat{\beta}_a$:

$$\frac{d\hat{\beta}_a}{dV_m} = \alpha(1 - \beta_m) \frac{d\beta_a}{dV_m} - \alpha\beta_a \frac{d\beta_m}{dV_m} \quad (113)$$

$$= \frac{\alpha\beta_m(1 - \beta_m)}{V_a V_m [1 - \alpha^2 [1 - (1 - \beta_m)^2(1 - \beta_a)^2]]} \{ \beta_a V_a (1 - \alpha^2) - V_\omega (1 - \beta_m)(1 - \beta_a)^2 \} \quad (114)$$

This is $\leq$ zero if:

$$\beta_a V_a (1 - \alpha^2) \leq V_\omega (1 - \beta_m)(1 - \beta_a)^2 \quad (115)$$

There is a turning point at $\hat{V}_m$ where $\hat{V}_m$ gives values of β_a and β_m which solve the following equation:

$$\beta_a V_a (1 - \alpha^2) = V_\omega (1 - \beta_m)(1 - \beta_a)^2 \quad (116)$$

The left-hand side of the above equation is increasing in V_m while the right-hand side is decreasing. Therefore $\hat{V}_m$ is a unique turning point and a maximum:

$$\frac{d [\beta_a V_a (1 - \alpha^2)]}{dV_m} = V_a (1 - \alpha^2) \frac{d\beta_a}{dV_m} < 0 \quad (117)$$

$$\frac{d [V_\omega (1 - \beta_m) (1 - \beta_a)^2]}{dV_m} = -V_\omega (1 - \beta_a) \left\{ 2 (1 - \beta_m) \frac{d\beta_a}{dV_m} + (1 - \beta_a) \frac{d\beta_m}{dV_m} \right\} > 0 \quad (118)$$

For values of $V_m < \hat{V}_m$, $\beta_a V_a (1 - \alpha^2) > V_\omega (1 - \beta_m) (1 - \beta_a)^2$ and $\frac{d\hat{\beta}_a}{dV_m} > 0$. For values of $V_m > \hat{V}_m$, $\beta_a V_a (1 - \alpha^2) < V_\omega (1 - \beta_m) (1 - \beta_a)^2$ and $\frac{d\hat{\beta}_a}{dV_m} < 0$. $\square$

I Proof of Result 8

Proof. From equation 18 we can see that:

$$y_{t+1}^i = (1 - \beta_m) \left[\ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} + \alpha (1 - \beta_a) \mu_y \right] + \frac{\beta_m V_m}{2} \quad (119)$$

$$+ \alpha \beta_a (1 - \beta_m) a_{t+1}^i + \beta_m m_{t+1}^i \quad (120)$$

Substituting for a_t^i and m_t^i :

$$y_{t+1}^i = (1 - \beta_m) \left[\ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} + \alpha (1 - \beta_a) \mu_y \right] + \frac{\beta_m V_m}{2} \quad (121)$$

$$+ \alpha \beta_a (1 - \beta_m) (y_t^i + \varepsilon_{t+1}^{ia}) + \beta_m (h_{t+1}^i + \varepsilon_{t+1}^{im}) \quad (122)$$

Equation 6 then allows us to substitute for log human capital:

$$y_{t+1}^i = (1 - \beta_m) \left[\ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} + \alpha (1 - \beta_a) \mu_y \right] + \frac{\beta_m V_m}{2} \quad (123)$$

$$+ \alpha \beta_a (1 - \beta_m) (y_t^i + \varepsilon_{t+1}^{ia}) + \beta_m \left(\ln Z + \alpha \ln X_{t-1}^i - \frac{V_\omega}{2} + \omega_{t+1}^i + \varepsilon_{t+1}^{im} \right) \quad (124)$$

It is then just a matter of rearranging to find γ_0 , γ_1 , γ_2 and ε

$$y_t^i = \left\{ \ln Z + \alpha (1 - \beta_m) [(1 - \beta_a) \mu_y + \ln \lambda] + \frac{1}{2} [\beta_m V_m - V_\omega] \right\} \quad (125)$$

$$+ \alpha \beta_a (1 - \beta_m) y_t^i + \alpha \beta_m \ln X_t^i + \alpha \beta_a (1 - \beta_m) \varepsilon_{t+1}^{ia} + \beta_m (\omega_{t+1}^i + \varepsilon_{t+1}^{im}) \quad (126)$$

λ is found by substitution of γ_1 and γ_2 into equation 7. Note that $\gamma_1 + \gamma_2 = \rho < 1 + \delta$ as required to ensure B is positive. $\square$

J Proof of Results 10 and 11

Proof.

$$\lambda = \frac{\beta_m \alpha}{1 + \delta - (1 - \beta_m) \alpha \beta_a} \quad (127)$$

The partial derivatives of λ with respect to β_m and β_a are therefore:

$$\frac{\partial \lambda}{\partial \beta_m} = \frac{\alpha [(1 + \delta) - \alpha \beta_a]}{[1 + \delta - (1 - \beta_m) \alpha \beta_a]^2} > 0 \quad (128)$$

$$\frac{\partial \lambda}{\partial \beta_a} = \frac{\alpha^2 \beta_m (1 - \beta_m)}{[1 + \delta - (1 - \beta_m) \alpha \beta_a]^2} > 0 \quad (129)$$

So:

$$\frac{d\lambda}{dV_m} = \frac{\alpha [(1+\delta) - \alpha\beta_a]}{[1+\delta - (1-\beta_m)\alpha\beta_a]^2} \frac{d\beta_m}{dV_m} + \frac{\alpha^2\beta_m(1-\beta_m)}{[1+\delta - (1-\beta_m)\alpha\beta_a]^2} \frac{d\beta_a}{dV_m} < 0 \quad (130)$$

With respect to V_a :

$$\frac{d\lambda}{dV_a} = \frac{\alpha [(1+\delta) - \alpha\beta_a]}{[1+\delta - (1-\beta_m)\alpha\beta_a]^2} \frac{d\beta_m}{dV_a} - \frac{\alpha^2\beta_m(1-\beta_m)}{[1+\delta - (1-\beta_m)\alpha\beta_a]^2} \frac{d\beta_a}{dV_a} \quad (131)$$

$$= \frac{\alpha^2\beta_a(1-\beta_m)}{V_a V_m [1+\delta - (1-\beta_m)\alpha\beta_a]^2} \left\{ \frac{\alpha(1-\alpha^2)\beta_a V_a (1-\beta_m) [(1+\delta) - \alpha\beta_a]}{1 - \alpha^2 [1 - (1-\beta_m)^2 (1-\beta_a)^2]} \right\} \quad (132)$$

$$- \frac{\alpha^2\beta_a(1-\beta_m)}{V_a V_m [1+\delta - (1-\beta_m)\alpha\beta_a]^2} \left\{ \frac{(1-\beta_a)\beta_m V_m [1 - \alpha^2 [1 - (1-\beta_a)(1-\beta_m)^2]]}{1 - \alpha^2 [1 - (1-\beta_m)^2 (1-\beta_a)^2]} \right\} \quad (133)$$

This is positive if:

$$\alpha \left[\frac{1 - \alpha^2}{1 - \alpha^2 [1 - (1-\beta_m)^2 (1-\beta_a)^2]} \right] (1-\beta_m) > \frac{V_m}{V_w} \left[\frac{1 - \alpha^2 [\beta_a(1-\beta_m) + \beta_m]}{[(1+\delta) - \alpha\beta_a]} \right] \quad (134)$$

Since the left-hand side of this equation is less than 1, it follows that a sufficient (though not necessary) condition for $\frac{d\lambda}{dV_a} < 0$ is:

$$\frac{V_m}{V_w} \left[\frac{1 - \alpha^2 [\beta_a(1-\beta_m) + \beta_m]}{[(1+\delta) - \alpha\beta_a]} \right] > 1 \quad (135)$$

□

K Privilege and Meritocracy

In this section we present an extension to the model where the children of the rich get an advantage independently of how productive they are thought to be.

As in the main text, human capital is a function of parental investment in education, and parents invest a fixed percentage of their income in education. Consequently, we can write the process of human capital acquisition as:

$$H_{t+1}^i = A \times (Y_t^i)^\alpha \times e^{\epsilon_{h,t+1}^i} \rightarrow h_{t+1}^i = \ln A + \alpha y_t^i + \epsilon_{t+1}^{ih} \quad ; \quad \epsilon_{t+1}^{ih} \sim N(0, V_h) \quad (136)$$

Also as in the main text, firms receive a signal on the human capital of agents. But, like in section 5.4 we assume that firms have no signal on the agent's background (i.e., $V_a \rightarrow \infty$). Therefore, the information available to the firm when forming an expectation about the agent's human capital is $\Omega_{t+1}^{ip} = \{m_{t+1}^i, \mu_{y_t}, V_{y_t}\}$, where

$$m_{t+1}^i = h_{t+1}^i + \epsilon_{t+1}^{im} \quad (137)$$

The difference with the main text is that we assume that the pricing mechanism shows discrimination against the poor. Not statistical discrimination, but raw preference-based discrimination. Firms somehow *prefer* to hire children of the rich *even when conditioning on what one would rationally think about their human capital*. Formally, an agent with certain observable traits $m_t^i + 1$ and parental income Y_t^i gets as income:

$$Y_{t+1}^i = \left(E \left(H_{t+1}^i | \Omega_{t+1}^{ip} \right) \right)^{1-\pi} \times (Y_t^i)^\pi \times e^{\epsilon_{t+1}^{iu}} \quad ; \quad \epsilon_{t+1}^{iu} \sim N(0, V_u) \quad (138)$$

where π is a parameter measuring how much *privilege* the children of the rich enjoy. If the elasticity of income to parental income *when conditioning on expected human capital* equals one, the income process is exogenously determined to be a random walk, with maximal correlation between parental and children's income (and unconditional income variance approaching infinity). The polar case is the one that we study in the main text, where $\pi = 0$, and conditional on human capital expectations, parental income is irrelevant.

Of course parental income could be used to form those expectations (as we do in the main text).

The difference now is that there is an advantage above and beyond the news that parental income may provide on human capital. Firms, thus, are inefficient in the sense of not maximizing profits because they overpay employees who are the children of rich people. In Becker's parlance they "prefer" to employ the children of the rich, and are willing to pay them more for the same expected output. We do not explain how it turns out that these firms exist, and why they are not driven to extinction by competitive firms that do not prefer the children of the rich (like the ones in the main text). Notice that if there was any information on parental income, the firms in the main text would also pay more to the children of the rich, but not beyond their expected human capital.

Given this set up, it is straight forward to notice that if the distribution of log income among the parents generation is $y_t^i \sim N(\mu_{y_t}, V_{y_t})$, the prior (before observing m_{t+1}^i) of the log of human capital of a certain individual i is:

$$h_{t+1}^i \sim N(\ln A + \alpha \mu_{y_t}, \alpha^2 V_{y_t} + V_h) \quad (139)$$

and the posterior (after observing a certain realization of m_{t+1}^i ,

$$h_{t+1}^i | \Omega_{t+1}^{ip} \sim N \left(\beta_{t+1} m_{t+1}^i + (1 - \beta_{t+1})(\ln A + \alpha \mu_{y_t}) \quad , \quad \frac{(\alpha^2 V_{y_t} + V_h) V_m}{\alpha^2 V_{y_t} + V_h + V_m} \right) \quad (140)$$

with

$$\beta_{t+1} = \frac{\alpha^2 V_{y_t} + V_h}{\alpha^2 V_{y_t} + V_h + V_m} \quad (141)$$

Thus,

$$\log E \left(H_{t+1}^i | \Omega_{t+1}^{ip} \right) = \beta_{t+1} m_{t+1}^i + (1 - \beta_{t+1})(\ln A + \alpha \mu_{y_t}) + \frac{\beta_{t+1} V_m}{2} \quad (142)$$

and

$$y_{t+1}^i = \pi y_t^i + (1 - \pi) \left[\beta_{t+1} m_{t+1}^i + (1 - \beta_{t+1})(\ln A + \alpha \mu_{y_t}) + \frac{\beta_{t+1} V_m}{2} \right] + \epsilon_{t+1}^{iu} \quad (143)$$

which can be written in a Becker-Tomes form by substituting m_{t+1}^i for its value $h_{t+1}^i + \epsilon_{t+1}^{im}$ (and then h_{t+1}^i by $\ln A + \alpha y_t^i + \epsilon_{t+1}^{ih}$):

$$y_{t+1}^i = (1 - \pi) \left[\ln A + \alpha(1 - \beta_{t+1}) \mu_{y_t} + \frac{\beta_{t+1} V_m}{2} \right] \quad (144)$$

$$+ [\pi + (1 - \pi) \alpha \beta_{t+1}] y_t^i + (1 - \pi) \beta_{t+1} (\epsilon_{t+1}^{ih} + \epsilon_{t+1}^{im}) + \epsilon_{t+1}^{iu} \quad (145)$$

Equation 145 determines a law of motion for the variance of log income, the persistence of the income process and β :

$$V_{y_{t+1}} = (\rho_{t+1})^2 V_{y_t} + (1 - \pi)^2 (\beta_{t+1})^2 (V_h + V_m) + V_u \quad (146)$$

$$\rho_{t+1} = \pi + (1 - \pi) \alpha \beta_{t+1} \quad (147)$$

$$\beta_{t+1} = \frac{\alpha^2 V_{y_t} + V_h}{\alpha^2 V_{y_t} + V_h + V_m} \quad (148)$$

In this extension we solve numerically, not analytically, for the properties of the solution of this system.

It is easy to check numerically that

Result 12. *As in our main model, in steady state, an exogenous increase in the degree of meritocracy (a decrease of V_m) leads to a new steady state with a higher degree of intergenerational correlation, (ρ), greater weight given to the observable signal (β), and a higher degree of inequality (V_y).*

The reasons are the same as in our main model: more meritocracy increases the differences between agents, as they are paid more accurately according to their productive ability, and given that the children of the rich are indeed on average more capable, this effect only increases the inequalities inherited via the privilege channel. The increase in inequality increases the value that the market assigns to objective information in the determination of the wage, β , increasing further the degree of inequality, etc...

This is the main result of this section: nothing of interest changes by introducing "irrational" privilege, at least with respect to the effects of meritocracy. It still increases inequality and still decreases mobility.

It should be mentioned though, that the effects of privilege (π) are more complicated than the ones that advantages (rational advantages) had in the main text. In particular, an increase in the degree

of privilege (π) always increases the intergenerational persistence of income, but it may decrease the degree of inequality (and thus, β).

L A Model of Dynamic Statistical Discrimination

Assuming $V_m \rightarrow \infty$ shuts off the role of the human capital signal, since it is useless and completely disregarded by firms. The first interesting thing to note is that in the model with $V_m \rightarrow \infty$ the unique steady state requires V_y to be zero for any value of V_a .³¹

The reason is as follows. Remember that different agents will receive different realizations of ω_t^i , implying that some get more human capital per unit of investment made on them. Thus, even if parents are homogeneous in their income, their children won't be in their productivities. Nevertheless in the absence of a signal on talent (if $V_m \rightarrow \infty$), the market is unable of seeing these differences. It only observes something that belongs to the parents, and the inference that it makes on parental income is mean reverting (because the posterior is a convex combination of the signal with the prior). Thus, the only possible steady state has no inequality. In the only steady state everybody is expected to have the same education, and everybody is *expected* to have the same productive ability. There is no way for the more talented to show their worth.

This is not the case in the full model ($V_m < \infty$). There the differences in innate productive ability of workers is reflected in m_t^i , even for the same parental income. Moreover, the mistakes in the perception of workers, ϵ_m^t , actually differentiate between families. With the same backgrounds, even with the same talents, some of them are *perceived* as better, and thus, paid more. This extra real income has real effects: it produces more education. Thus, firms will be attentive to parental income, generating differences in income for their children.

Thus, to introduce meaningfully statistical discrimination while abstracting from the “meritocracy” aspect, the shorter route is to introduce some noise in the income process. We assume that parental income is:

$$Y_{t+1}^i = E [\exp \{h_{t+1}^i\} | \Omega_{t+1}^i] \exp \{u_{t+1}^i\} \quad (149)$$

with $u_{t+1}^i \sim N(0, V_u)$ and $V_u > 0$, and $\Omega_{t+1}^i = \{a_{t+1}^i, \mu_{y_t}, V_{y_t}\}$.

The stochastic noise u_t^i generates income variance, in the same manner that the signal noise in talent does in the full model. This noise is not related to the real productivity of the worker, but that does not matter. The important thing is that it produces different incomes, so the next generation of parents will invest different amounts in the education of their kids. Their kids, thus, will be expected to have different productive abilities when conditioning on parental income.

The law of motion of log income now given by the following equation, analogous to equation 19 in the general model:

$$y_{t+1}^i = \ln Z + \alpha \ln \lambda + \alpha (1 - \beta_a) \mu_y + \frac{\alpha^2 \beta_a V_a}{2} + \alpha \beta_a y_t^i + \alpha \beta_a \varepsilon_{t+1}^{ai} + u_{t+1}^i \quad (150)$$

Intergenerational mobility is then given by $\rho = \alpha \beta_a$, the weight of parental income in the agent's income.

Result 13. ³² *There exists a unique steady state for the system of differential equations defined by (150) and (16). This results from solving:*

$$V_y = \frac{V_u}{1 - \alpha^2 \beta_a} \quad (151)$$

$$\beta_a = \frac{V_y}{V_y + V_a} \quad (152)$$

Equations 151 and 152 completely define the steady state and show the feed back mechanism which we want to stress. Equation 151 indicates that the larger the role of background in the determination

³¹To see this note that in such a case $\beta_m = 0$

³²Proof in appendix L.1

of income, the larger the degree of inequality. From equation 152 the more inequality there is, the more that the market will rationally give extra advantages to people from good backgrounds. Both together indicate that inequality fosters further inequality because it increases the demand for information on people's background.

Consequently, any exogenous change that increases inequality will have a general equilibrium effect larger than its initial effect – an inequality multiplier. Furthermore, the degree of intergenerational mobility decreases as firms give more weight to the background of agents.

Result 14. ³³ *The full effect of an exogenous change in the parameters α , V_a or V_u is greater than the partial effect due to the feed back from income inequality to discrimination*

$$\left| \frac{dV_y}{d\chi} \right| > \left| \frac{\partial V_y}{\partial \chi} \right| ; \quad \left| \frac{d\beta_a}{d\chi} \right| > \left| \frac{\partial \beta_a}{\partial \chi} \right| \quad \text{where } \chi = \{\alpha, V_a, V_u\} \quad (153)$$

Any exogenous change in parameters that has a direct effect on either the degree of inequality or the use of information on background, will have its effect amplified via the other variable. For example, if there is an increase in the variance of income shocks, V_u , its effect on the variance of income is larger than the direct one on 151, because the increase in inequality feeds into an increase in the degree of discrimination, which feeds back into inequality, which feeds back into discrimination, and so on.

Moreover, increases in the precision of the signal on background increase inequality, the value of the signal and the persistence of income across generations:

Result 15. ³⁴ *The solution of (151) and (152) implies that V_y , ρ and β_a are all decreasing in V_a .*

Notice that the result is not obvious. Let's look again at equation 150 and consider an exogenous decrease in the amount of informational noise (a smaller value of V_a). For a given β_a the variance of income will fall in the subsequent generation. However, once we allow β_a to adjust endogenously we find just the opposite: the less informational noise we feed into the economy, the *larger* the variance of log income in the economy. The reason is precisely *because* the increase in precision translates into more discrimination, which itself translates into a larger intergenerational correlation ($\alpha\beta_a$). Thus, the income process becomes more persistent, and this is bound to increase the *unconditional* variance of income, *even if what we did in the first place was to decrease the conditional variance of income*.

L.1 Proof of Result 14

Proof. Let:

$$\Psi_a(V_y, \alpha, V_a, V_u) = \alpha^2 \frac{V_y^2}{V_y + V_a} + V_u \quad (154)$$

where the steady state is defined by $V_y = \Psi_a(V_y, \alpha, V_a, V_u)$.

The steady state is unique as $\Psi_a(V_y, \alpha, V_a, V_u)$ is increasing in V_y , but with a slope smaller than one. As $\Psi_a(V_y = 0, \alpha, V_a, V_u) = V_u > 0$, there exists a unique fixed point.

$$\frac{\partial \Psi_a(V_y, \alpha, V_a, V_u)}{\partial V_y} = \alpha^2 \left[1 - \frac{V_a^2}{(V_y + V_a)^2} \right] \in (0, 1) \quad (155)$$

As a result, we generate multiplier effects from a change in any of the parameters (α , V_a or V_u) which shift $\Psi_a(V_y, \alpha, V_a, V_u)$. If they cause a fall in $\Psi_a(V_y, \alpha, V_a, V_u)$, V_y will have to fall to bring the steady state equation back into equality, which will further reduce $\Psi_a(V_y, \alpha, V_a, V_u)$ and so on. Similarly, any exogenous change in the parameters which increases $\Psi_a(V_y, \alpha, V_a, V_u)$ will increase V_y , further increasing $\Psi_a(V_y, \alpha, V_a, V_u)$.

For example, a change in V_u has the partial effect:

$$\frac{\partial \Psi_a(V_y, \alpha, V_a, V_u)}{\partial V_u} = 1 \quad (156)$$

³³Proof in appendix L.1

³⁴Proof in appendix L.2

But the overall effect of the change on steady state V_y is:

$$\frac{d\Psi_a(V_y, \alpha, V_a, V_u)}{dV_u} = 1 + \frac{\partial\Psi_a(V_y, \alpha, V_a, V_u)}{\partial V_y} \frac{dV_y}{dV_u} \quad (157)$$

$$\frac{dV_y}{dV_u} = 1 + \alpha^2 \beta_a (1 + (1 - \beta_a)) \frac{dV_y}{dV_u} \quad (158)$$

$$\frac{dV_y}{dV_u} = \frac{1}{1 - \alpha^2 [1 - (1 - \beta_a)^2]} > \frac{\partial\Psi_a(V_y, \alpha, V_a, V_u)}{\partial V_u} \quad (159)$$

In a similar manner, if we call $B_a(\beta_a, \alpha, V_a, V_u) = \frac{V_u}{V_u + V_a(1 - \alpha^2 \beta_a)}$, we can show that:

$$\frac{d\beta_a}{dV_u} = \frac{V_a (1 - \alpha^2 \beta_a)}{[V_u + V_a (1 - \alpha^2 \beta_a)]^2 - \alpha^2 V_a V_u} > \frac{\partial B_a(\beta_a, \alpha, V_a, V_u)}{\partial V_u} \quad (160)$$

It is easy to show the partial and total effect from a change in α or V_a in exactly the same way, proving that that the total effect on steady state β_a or V_y from a change in the parameters is greater in absolute terms than the partial effect (or, alternatively, the long run effect is greater than the short run one). $\square$

L.2 Proof of Result 15

Proof. Comparative statics from a change in the precision of the signal a :

$$\frac{d\Psi_a(V_y, \alpha, V_a, V_u)}{dV_a} = \frac{\partial\Psi_a(V_y, \alpha, V_a, V_u)}{\partial V_a} + \frac{\partial\Psi_a(V_y, \beta_a, \alpha, V_u)}{\partial V_y} \frac{dV_y}{dV_a} \quad (161)$$

$$\frac{dV_y}{dV_a} = -\alpha^2 \beta_a^2 + \alpha^2 \beta_a (1 + (1 - \beta_a)) \frac{dV_y}{dV_a} \quad (162)$$

$$\frac{dV_y}{dV_a} = \frac{-\alpha^2 \beta_a^2}{1 - \alpha^2 [1 - (1 - \beta_a)^2]} < 0 \quad (163)$$

$$\frac{d\beta_a}{dV_a} = \frac{\partial\beta_a}{\partial V_a} + \frac{\partial\beta_a}{\partial V_y} \frac{dV_y}{dV_a} \quad (164)$$

$$= \frac{-\beta_a (1 - \beta_a)}{V_a} \left[\frac{1 - \alpha^2 \beta_a}{1 - \alpha^2 [1 - (1 - \beta_a)^2]} \right] < 0 \quad (165)$$

Since $\rho = \alpha\beta_a$ it follows that:

$$\frac{d\rho}{dV_a} = \frac{-\alpha\beta_a (1 - \beta_a)}{V_a} \left[\frac{1 - \alpha^2 \beta_a}{1 - \alpha^2 [1 - (1 - \beta_a)^2]} \right] < 0 \quad (166)$$

$\square$

M The Cost of Misallocation

One of the omissions from the original model is a cost associated with the misallocation of workers given that firms do not perfectly observe their human capital. In equation 11, a worker's output is assumed to equal his human capital. This was done for simplicity. However, if a firm cannot perfectly identify that human capital it is likely that he will not be able to realise his full productive potential, and all the moreso as the error in the firm's belief grows. In this appendix, the role of allocating workers to tasks is added into the model. It does not have any effect on the results presented in the body of the paper.

In order to capture the role of allocation in production, a term A is added to the production function. This allocative cost pulls the output of the worker down to below their full level of human capital. The

greater the error is firms' beliefs, the greater the degree of misallocation and the lower the level of output. Firms, as before, pay workers their expected output conditional on the available information but taking into account that their output will be lowered by misallocation.

$$\begin{aligned} Y_{t+1}^i &= E \left[A \exp \{ h_{t+1}^i \} | \Omega_{t+1}^i \right] \\ A &= \exp \left\{ -\frac{\theta}{2} (h_{t+1}^i - E(h_{t+1}^i))^2 \right\} \\ \theta &\in R^+ \end{aligned} \quad (167)$$

The allocative cost is given by the square of the difference between actual and expected human capital. An exogenous parameter θ can control the extent to which this misallocation impacts on production. If θ is set equal to zero, this returns us to the model in the text.

Intuitively, since the human capital of any individual worker is unknown to the firm they cannot make adjustments to the individual wages paid. They adjust everyone's income by the same amount depending on the aggregate amount of misallocation, which is itself an endogenous function of the income distribution. As a result, the additional of this allocative cost impacts on the mean of income, but not the intergenerational elasticity or inequality presented above. Thus the main results of the paper are preserved without any adjustment.

The additional of A alters result 4 in the following way.

Result 16. *Given the posterior belief about the log of human capital follows a normal distribution*

$$h_{t+1}^i | \Omega_{t+1}^i \sim N \left(\mu_{h_{t+1}^i | \Omega_{t+1}^i}, V_{h_{t+1}^i | \Omega_{t+1}^i} \right) \quad (168)$$

and income is given by equation (167), it follows that:

$$Y_{t+1}^i = \frac{1}{\sqrt{1 + \theta V_{h_{t+1}^i | \Omega_{t+1}^i}}} \exp \left\{ \frac{1}{2} \left(\frac{V_{h_{t+1}^i | \Omega_{t+1}^i}}{1 + \theta V_{h_{t+1}^i | \Omega_{t+1}^i}} \right) \right\} \exp \left\{ \mu_{h_{t+1}^i | \Omega_{t+1}^i} \right\} \quad (169)$$

Proof. We need to calculate

$$Y_{t+1}^i = E \left[\exp \left\{ h_{t+1}^i - \frac{\theta}{2} (h_{t+1}^i - \mu_{h_{t+1}^i | \Omega_{t+1}^i})^2 \right\} \right] \quad (170)$$

$$= \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi} \sqrt{V_{h_{t+1}^i | \Omega_{t+1}^i}}} \exp \left\{ h_{t+1}^i - \frac{\theta}{2} (h_{t+1}^i - \mu_{h_{t+1}^i | \Omega_{t+1}^i})^2 - \frac{1}{2} \frac{(h_{t+1}^i - \mu_{h_{t+1}^i | \Omega_{t+1}^i})^2}{V_{h_{t+1}^i | \Omega_{t+1}^i}} \right\} dh_{t+1}^i \quad (171)$$

After some manipulation, this becomes:

$$Y_{t+1}^i = \frac{1}{\sqrt{1 + \theta V_{h_{t+1}^i | \Omega_{t+1}^i}}} \exp \left\{ \frac{1}{2} \left(\frac{V_{h_{t+1}^i | \Omega_{t+1}^i}}{1 + \theta V_{h_{t+1}^i | \Omega_{t+1}^i}} \right) + \mu_{h_{t+1}^i | \Omega_{t+1}^i} \right\} \quad (172)$$

$$\cdot \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi} \sqrt{\frac{V_{h_{t+1}^i | \Omega_{t+1}^i}}{1 + \theta V_{h_{t+1}^i | \Omega_{t+1}^i}}}} \exp \left\{ -\frac{1}{2} \frac{\left(h_{t+1}^i - \left[\mu_{h_{t+1}^i | \Omega_{t+1}^i} + \left(\frac{V_{h_{t+1}^i | \Omega_{t+1}^i}}{1 + \theta V_{h_{t+1}^i | \Omega_{t+1}^i}} \right) \right] \right)^2}{\left(\frac{V_{h_{t+1}^i | \Omega_{t+1}^i}}{1 + \theta V_{h_{t+1}^i | \Omega_{t+1}^i}} \right)} \right\} dh_{t+1}^i \quad (173)$$

The second term is equal to one, therefore income of individual i is given by:

$$Y_{t+1}^i = \frac{1}{\sqrt{1 + \theta V_{h_{t+1}^i | \Omega_{t+1}^i}}} \exp \left\{ \frac{1}{2} \left(\frac{V_{h_{t+1}^i | \Omega_{t+1}^i}}{1 + \theta V_{h_{t+1}^i | \Omega_{t+1}^i}} \right) \right\} \exp \left\{ \mu_{h_{t+1}^i | \Omega_{t+1}^i} \right\} \quad (174)$$

□

By taking logs of equation (169) and substituting from result 3 we can find the log income of individual i with signals a_{t+1}^i and m_{t+1}^i :

$$\begin{aligned}
y_{t+1}^i = & (1 - \beta_m) \left[\ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} + \alpha (1 - \beta_a) \mu_y \right] \\
& - \frac{1}{2} \left[\ln (1 + \theta \beta_m V_m) - \frac{\beta_m V_m}{1 + \theta \beta_m V_m} \right] \\
& + \alpha \beta_a (1 - \beta_m) a_{t+1}^i + \beta_m m_{t+1}^i
\end{aligned} \tag{175}$$

This is equivalent to equation (18) in the main body of the paper. Given that a_{t+1}^i and m_{t+1}^i are both stochastic functions of y_t^i we can write the law of motion of the log of income:

$$\begin{aligned}
y_{t+1}^i = & \ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} + \alpha (1 - \beta_a) (1 - \beta_m) \mu_y \\
& - \frac{1}{2} \left[\ln (1 + \theta \beta_m V_m) - \frac{\beta_m V_m}{1 + \theta \beta_m V_m} \right] \\
& + \alpha [\beta_a (1 - \beta_m) + \beta_m] y_t^i + \alpha \beta_a (1 - \beta_m) \varepsilon_{t+1}^{ia} + \beta_m (\varepsilon_{t+1}^{im} + \omega_{t+1}^i)
\end{aligned} \tag{176}$$

This is equivalent to equation (19) in the text. It can clearly be seen that the only role for misallocation is in the intercept. As misallocation plays a larger role (θ is higher), mean log income is reduced, as would be expected. The steady state equations for β_a , β_m , V_y and ρ are the same as those given in result 5.

Result 17. *Steady State with misallocation.* Given the process of human capital accumulation in equation (10), income given by equation (167), and the information set given by equation (12), there exists a unique steady state which is globally stable. In the steady state log income is normally distributed with variance being characterized by the (unique) solution of the following system of equations:

$$V_y = \frac{\beta_m V_\omega}{1 - \alpha^2 [\beta_a (1 - \beta_m) + \beta_m]} \tag{177}$$

$$\beta_a = \frac{V_y}{V_y + V_a} \tag{178}$$

$$\beta_m = \frac{\alpha^2 \beta_a V_a + V_\omega}{\alpha^2 \beta_a V_a + V_\omega + V_m} \tag{179}$$

The steady state mean of log income and intergenerational correlation of income are given by:

$$\mu_y = \frac{\ln Z + \alpha \ln \lambda - \frac{V_\omega}{2} - \frac{1}{2} \left[\ln (1 + \theta \beta_m V_m) - \frac{\beta_m V_m}{1 + \theta \beta_m V_m} \right]}{1 - \alpha} \tag{180}$$

$$\rho = \alpha [\beta_a (1 - \beta_m) + \beta_m] \tag{181}$$

Misallocation will push down median income (e^{μ_y}) but will otherwise leave the steady state unaffected. All of the results from sections 5.3 and 5.4 are preserved.

Similarly, misallocation will have a fairly limited impact on the equilibrium presented in section 6. The γ_0 term will be impacted upon by the degree of misallocation, but otherwise nothing is affected.

Result 18. *Equilibrium with misallocation.* The equilibrium stochastic process of income as a function of parental income and investment is $Y_{t+1}^i = e^{\gamma_0} (Y_t^i)^{\gamma_1} (X_t^i)^{\gamma_2} e^{\varepsilon_{t+1}^i}$ with:

$$\gamma_0 = \ln Z + \alpha (1 - \beta_m) [(1 - \beta_a) \mu_y + \ln \lambda] - \frac{1}{2} \left[\ln (1 + \theta \beta_m V_m) - \frac{\beta_m V_m}{1 + \theta \beta_m V_m} + V_\omega \right] \tag{182}$$

$$\gamma_1 = \alpha \beta_a (1 - \beta_m) \tag{183}$$

$$\gamma_2 = \alpha \beta_m \tag{184}$$

$$\varepsilon_{t+1}^i = \alpha \beta_a (1 - \beta_m) \varepsilon_{t+1}^{ai} + \beta_m (\omega_{t+1}^i + \varepsilon_{t+1}^{mi}) \tag{185}$$

and, consequently, the equilibrium share of income invested in children's education is:

$$\lambda = \frac{\alpha\beta_m}{1 + \delta - \alpha\beta_a(1 - \beta_m)} \quad (186)$$

We now have a more realistic model of production which accounts for the degree of asymmetry of information between the worker and firm, and the impact this will have on worker's productivity. As such, it might be instructive to look at the impact on median income of increases in the degree of meritocracy or inherited advantage. That is the objective of the remainder of this appendix. Result 19 summarises the findings.

Result 19. *When misallocation is sufficiently costly (θ is above a threshold level), an increase in the accuracy of either signal will raise median income. The threshold value above which this occurs for the ability signal, $\bar{\theta}$, solves:*

$$\frac{2\alpha}{\lambda} \frac{\frac{d\lambda}{dV_m}}{\left[\beta_m + V_m \frac{d\beta_m}{dV_m}\right]} = \frac{[\bar{\theta}(1 + \bar{\theta}\beta_m V_m) - 1]}{(1 + \bar{\theta}\beta_m V_m)^2} \quad (187)$$

The threshold value above which this occurs for the background signal, $\hat{\theta}$, solves:

$$\frac{2a}{\lambda} \frac{\frac{d\lambda}{dV_a}}{V_m \frac{d\beta_m}{dV_a}} = \frac{[\hat{\theta}(1 + \hat{\theta}\beta_m V_m) - 1]}{(1 + \hat{\theta}\beta_m V_m)^2} \quad (188)$$

Proof. Median income is e^{μ_y} so the log of median income is μ_y .

$$\mu_y = \frac{\ln Z + \alpha \ln \lambda - \frac{1}{2} \left[\ln(1 + \theta\beta_m V_m) - \frac{\beta_m V_m}{1 + \theta\beta_m V_m} + V_\omega \right]}{(1 - \alpha)} \quad (189)$$

The derivative with respect to V_m is given by:

$$\frac{d\mu_y}{dV_m} = \frac{\partial \mu_y}{\partial V_m} + \frac{\partial \mu_y}{\partial \lambda} \frac{d\lambda}{dV_m} + \frac{\partial \mu_y}{\partial \beta_m} \frac{d\beta_m}{dV_m} \quad (190)$$

$$= -\frac{[\theta(1 + \theta\beta_m V_m) - 1]}{(1 + \theta\beta_m V_m)^2} \frac{\beta_m}{2(1 - \alpha)} + \frac{\alpha}{(1 - \alpha)} \frac{d\lambda}{\lambda dV_m} \quad (191)$$

$$- \frac{[\theta(1 + \theta\beta_m V_m) - 1]}{(1 + \theta\beta_m V_m)^2} \frac{V_m}{2(1 - \alpha)} \frac{d\beta_m}{dV_m} \quad (192)$$

$$= - \left[\frac{\beta_m + V_m \frac{d\beta_m}{dV_m}}{2(1 - \alpha)} \right] \left\{ \frac{\theta(1 + \theta\beta_m V_m) - 1}{(1 + \theta\beta_m V_m)^2} - \frac{2\alpha \frac{d\lambda}{dV_m}}{\lambda \left[\beta_m + V_m \frac{d\beta_m}{dV_m} \right]} \right\} \quad (193)$$

Note that

$$\beta_m + V_m \frac{d\beta_m}{dV_m} = \frac{(1 - \alpha^2) \beta_m^2}{(1 - \alpha^2) + \alpha^2(1 - \beta_m)^2(1 - \beta_a)^2} > 0 \quad (194)$$

The derivative $\frac{d\mu_y}{dV_m}$ is negative if:

$$\frac{2\alpha \frac{d\lambda}{dV_m}}{\lambda \left[\beta_m + V_m \frac{d\beta_m}{dV_m} \right]} < \frac{\theta(1 + \theta\beta_m V_m) - 1}{(1 + \theta\beta_m V_m)^2} \quad (195)$$

The left-hand side is independent of θ . The right-hand side is increasing in θ :

$$\frac{d}{d\theta} \left[\frac{\theta(1 + \theta\beta_m V_m) - 1}{(1 + \theta\beta_m V_m)^2} \right] = \frac{1 + \beta_m V_m (\theta + 2)}{(1 + \theta\beta_m V_m)^3} > 0 \quad (196)$$

So if there is a sufficient cost of misallocation, then an increase in the accuracy of the signal on ability leads to higher median income. As the left-hand side of equation (195) is negative, a sufficient condition for $\frac{d\mu_y}{dV_m}$ to be negative is:

$$\theta(1 + \theta\beta_m V_m) - 1 > 0 \quad (197)$$

which holds if $\theta \geq 1$. $\theta \geq 1$ is therefore a sufficient condition for a fall in V_m to raise median income.

The derivative of μ_y with respect to V_a is given by:

$$\frac{d\mu_y}{dV_a} = \frac{\partial\mu_y}{\partial\lambda} \frac{d\lambda}{dV_a} + \frac{\partial\mu_y}{\partial\beta_m} \frac{d\beta_m}{dV_a} \quad (198)$$

$$= \frac{a}{(1-\alpha)\lambda} \frac{d\lambda}{dV_a} - \frac{[\theta(1+\theta\beta_m V_m) - 1]}{(1+\theta\beta_m V_m)^2} \frac{V_m}{2(1-\alpha)} \frac{d\beta_m}{dV_a} \quad (199)$$

$$= \frac{V_m}{2(1-\alpha)} \frac{d\beta_m}{dV_a} \left[\frac{2a}{\lambda V_m} \frac{\frac{d\lambda}{dV_a}}{\frac{d\beta_m}{dV_a}} - \frac{[\theta(1+\theta\beta_m V_m) - 1]}{(1+\theta\beta_m V_m)^2} \right] \quad (200)$$

This is negative if:

$$\frac{2a}{\lambda V_m} \frac{\frac{d\lambda}{dV_a}}{\frac{d\beta_m}{dV_a}} < \frac{[\theta(1+\theta\beta_m V_m) - 1]}{(1+\theta\beta_m V_m)^2} \quad (201)$$

The left-hand side of equation 201 is also independent of θ . Thus, a sufficiently large θ will again ensure that the inequality holds and that median income is increasing in the accuracy of the background signal. $\square$

Misallocation makes an important contribution to ensuring that improvements in the information available to firms lead to increases in median income. It is one of the channels through which this can occur (another being parental investment) and, provided misallocation has a sufficient impact on productivity, ensures that it does. That is, if the problem of asymmetric information between the worker and firm is acute, improvements in the precision of information available to the firm will lead to higher median income. Furthermore, the addition of misallocation does not alter the results set out in the body of the paper above.

N From rank-rank to intergenerational income correlation.

- Online Data Table 5 has measures of intergenerational correlation. Need to make some assumptions that map these into the intergenerational correlation of income coefficient. We have "RM" = "Relative Mobility" = Rank-Rank Slope, β_1 ; under some model, $\beta_1 \implies \rho$
- Have an income distribution

$$\ln Y \sim N(\bar{y}, V_y)$$

and given large populations, we can assume perfect sampling from this distribution. The rank of a draw is therefore 1 minus the probability that any other draw is less than that i.e. let y be a draw from $\ln Y$, then the rank of y , $r(y) = 1 - F(y)$. This can be sampled across generations to obtain samples $(\{y_i^t, y_i^{t+1}\}, \{r_i^t, r_i^{t+1}\})$.

- Then we want to estimate $0 < \hat{\rho} < 1$ from

$$\begin{aligned} y_i^{t+1} &= \bar{y} + \rho(y_i^t - \bar{y}) + \varepsilon_i^y \\ &= (1 - \rho)\bar{y} + \rho y_i^t + \varepsilon_i^y \end{aligned}$$

but we have only the estimate $\hat{\beta}_1$ from

$$r_i^{t+1} = \beta_0 + \beta_1 r_i^t + \varepsilon_i^r$$

in the data.

- It is clear that $R \sim U[0, 1]$, so

$$P(R > r) = 1 - r$$

We also have

$$1 - r(y) = F(y) = \Phi\left(\frac{y - \bar{y}}{\sqrt{V_y}}\right)$$

i.e.

$$P(R > r(y)) = P(\ln Y < y) = \Phi\left(\frac{y - \bar{y}}{\sqrt{V_y}}\right)$$

- Also have

$$\text{Corr}(r_i^{t+1}, r_i^t) = \beta_1 \frac{\text{stdev}(r_i^{t+1})}{\text{stdev}(r_i^t)} = \beta_1$$

$$\text{Corr}(y_i^{t+1}, y_i^t) = \rho \frac{V_y^{\frac{1}{2}}}{V_y^{\frac{1}{2}}} = \rho$$

and

$$\begin{aligned} y^t &= \bar{y} + aZ_1 + bZ_2 \\ y^{t+1} &= \bar{y} + cZ_1 + dZ_2 \\ \begin{bmatrix} y^t \\ y^{t+1} \end{bmatrix} &= \begin{bmatrix} \bar{y} \\ \bar{y} \end{bmatrix} + \begin{bmatrix} a & b \\ c & d \end{bmatrix} \begin{bmatrix} Z_1 \\ Z_2 \end{bmatrix} \\ y &= \bar{y} + AZ \end{aligned}$$

where

$$\begin{aligned} Z_1, Z_2 &\sim N(0, 1) \text{ are independent} \\ a^2 + b^2 &= V_y = c^2 + d^2 \\ E[y^t y^{t+1}] &= E[(\bar{y} + aZ_1 + bZ_2)(\bar{y} + cZ_1 + dZ_2)] \end{aligned}$$

$$\begin{aligned} \text{Cov}[y^t, y^{t+1}] &= E[y^t y^{t+1}] - E[y^t] E[y^{t+1}] \\ &= E[(\bar{y} + aZ_1 + bZ_2)(\bar{y} + cZ_1 + dZ_2)] - \bar{y}^2 \\ &= ac + bd = \rho V_y \end{aligned}$$

i.e.

$$\begin{aligned} \text{Let } b &= 0 \\ \text{then } a &= V_y^{\frac{1}{2}} \\ c &= \rho V_y^{\frac{1}{2}} \\ d &= \sqrt{V_y(1 - \rho^2)} \end{aligned}$$

$$AA' = \begin{bmatrix} a & b \\ c & d \end{bmatrix} \begin{bmatrix} a & c \\ b & d \end{bmatrix} = V_y \begin{bmatrix} 1 & \rho \\ \rho & 1 \end{bmatrix}$$

allows us to determine a, b, c, d in terms of ρ (which we don't know yet).

- For each value of ρ , $\bar{y}$ and V_y however, we can simulate many realisations of Z_1 and Z_2 , and therefore of y_i^t and y_i^{t+1} and therefore of r_i^t and r_i^{t+1} and so, if this appears to be independent of $\bar{y}$ and V_y , then we can produce a map between ρ and β_1 . The relationship does seem to be independent of $\bar{y}$ and V_y , and it looks like very close relationship between ρ and β_1 , as it can be seen in figure 8

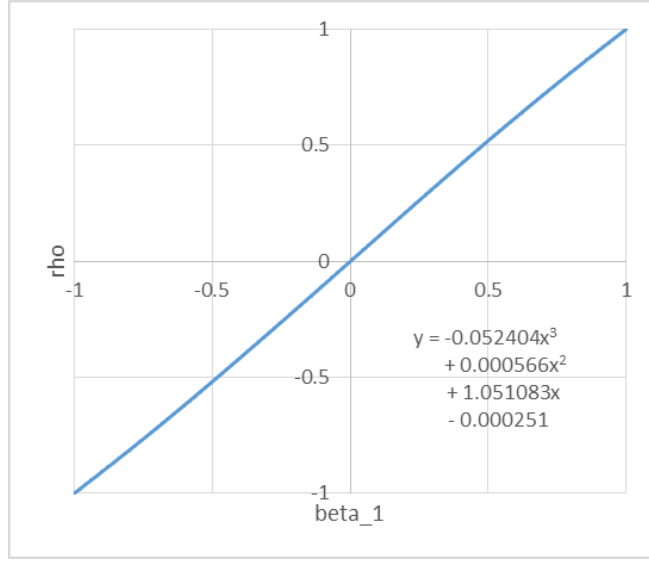


Figure 8: Mapping from rank-rank correlation (β_1 in the x axis) into intergenerational income correlation (ρ in the y axis)

O Calibrationa Procedure Appendix

O.1 Basic Idea

- Model Equations:
for location i

$$V_y(i) = \frac{\beta_m(i) V_\omega}{1 - \alpha \rho_i}$$

$$\rho_i = \alpha (\beta_a(i) (1 - \beta_m(i)) + \beta_m(i))$$

$$\lambda_i = \frac{\alpha \beta_m}{1 + \delta - \alpha \beta_a(i) (1 - \beta_m(i))}$$

with

$$\beta_m(i) = \frac{V_y(i) (1 - \alpha \rho_i)}{V_\omega}$$

$$\beta_a(i) = \frac{\rho_i - \alpha \beta_m(i)}{\alpha (1 - \beta_m(i))} = \frac{\rho_i / \alpha - \beta_m(i)}{1 - \beta_m(i)}$$

$$V_a(i) = V_y(i) \frac{1 - \beta_a(i)}{\beta_a(i)}$$

$$V_m(i) = (\alpha^2 \beta_a(i) V_a(i) + V_\omega) \frac{1 - \beta_m(i)}{\beta_m(i)}$$

where:

- λ_i is total parental investment in next generation as a share of income,
- $V_y(i)$ is the variance of log income, and
- ρ_i is the intergenerational elasticity
- Calibration strategy is
 - Given a set of "global" parameters $\{\alpha, V_\omega, \delta\}$, and data $T_i^D = \{\hat{\lambda}_i, \hat{V}_y(i), \hat{\rho}_i\}$, find then values for $\{V_a(i), V_m(i)\}$ such that $T_i^M = \{\lambda_i, V_y(i), \rho_i\}$ best matches T_i^D

- For each location i , this is an overidentified system of 3 equations in 2 unknowns which therefore cannot be matched exactly i.e. there will be some aggregate error value $E [|T_i^M - T_i^D|, \forall i]$ (α, V_ω) that is a function of the global parameters
- The "global" parameters α, V_ω & δ are determined by minimising the value of this aggregate error.

Implementation

The data moments that we use are: $\{\hat{V}_y(i), \hat{\rho}_i, \hat{\kappa}_i : \forall i\}$ where $\hat{\lambda}_i = \mu + \sigma \hat{\kappa}_i$, as described in the text.

- To generate initial conditions to start the calibration, we set:

$$\begin{aligned}\mu &= 0, \sigma = 1 \\ \delta &= 1.01^{30} - 1 \text{ (i.e. assume } \delta \equiv \text{ann discount rate of 1\%)} \\ \alpha &= 0.5 \text{ (we test that the result is not sensitive to this initial condition)} \\ V_a(i) &= V_m(i) = V_\omega = 1, \forall i\end{aligned}$$

- We do not vary δ in the calibration.
- Then $V_a(i) = V_m(i) = V_\omega, \forall i$, so that $(\rho_i, V_y(i))^M = (\bar{\rho}, \bar{V}_y)^D, \forall i$
- Then keeping these initial global parameters constant, we choose, for each i , $(V_a(i), V_m(i))$ to minimise $(\ln \rho_i - \ln \hat{\rho}_i)^2 + (\ln V_y(i) - \ln \hat{V}_y(i))^2$. This ensures that correlations between data and model ρ 's and V_y 's are well defined rather than being divide by zero errors. i.e. in each location we have an initial local calibration that matches local inequality and mobility but pays no attention to investment.
- The main part of the calibration routine is then as follows:

1. Let $tol = 100$
2. For $i = 1$ to 701, choose $(V_a(i), V_m(i))$ to minimise

$$\begin{aligned}Obj &= (1 - Corr[\hat{\rho}, \rho])^2 + \left(1 - Corr[\hat{V}_y, V_y]\right)^2 + \left(1 - Corr[\hat{\lambda}, \lambda]\right)^2 \\ &+ (\ln \bar{\rho}^D - \ln \bar{\rho}^M)^2 + (\ln \bar{V}_y^D - \ln \bar{V}_y^M)^2\end{aligned}$$

3. Repeat 2. until Obj is stable, then save the matrix $V_{MAT}^{agg} = \{(V_a(i), V_m(i)) : \forall i\}$
4. Choose μ & σ so that the mean and standard deviation of the $\hat{\lambda}$ matches the modelled equivalents, subject to the constraint that $\min \{\hat{\lambda}_i\} \geq 0$. NB Always the case, independent of μ & $\sigma > 0$, that $Corr[\hat{\kappa}, \hat{\lambda}] = 1$.
5. For $i = 1$ to 701, if $Obj_i > tol$, choose $(V_a(i), V_m(i))$ to minimise

$$Obj_i = (\ln(\hat{V}_y(i)) - \ln(V_y(i)))^2 + (\ln(\hat{\rho}_i) - \ln(\rho_i))^2 + (\ln(\hat{\lambda}_i) - \ln(\lambda_i))^2$$

then save the matrix $V_{MAT}^{out} = \{(V_a(i), V_m(i)) : \forall i\}$

6. Count the number of rows in which $V_{MAT}^{agg} \neq V_{MAT}^{out}$, repeat from 2. until this number is constant.
7. Let $tol = \frac{1}{10} tol$, then repeat from 2. unless $tol < 1$
8. For $i = 1$ to 701, choose $(V_a(i), V_m(i))$ to minimise

$$\begin{aligned}Obj &= (1 - Corr[\hat{\rho}, \rho])^2 + \left(1 - Corr[\hat{V}_y, V_y]\right)^2 + \left(1 - Corr[\hat{\lambda}, \lambda]\right)^2 \\ &+ (\ln \bar{\rho}^D - \ln \bar{\rho}^M)^2 + (\ln \bar{V}_y^D - \ln \bar{V}_y^M)^2\end{aligned}$$

9. Repeat 8. until Obj is stable

10. Let $Inc = 10\%$

11. Let $V_{MAT} = \{(V_a(i), V_m(i)) : \forall i\}$; $Params = \{\alpha, V_\omega\}$

12. Let $\{(V_a(i), V_m(i)) : \forall i\} = V_{MAT}$; No change to $Params$; For $i = 1$ to 701, choose $(V_a(i), V_m(i))$ to minimise

$$\begin{aligned} Obj(0) &= (1 - Corr[\hat{\rho}, \rho])^2 + \left(1 - Corr[\hat{V}_y, V_y]\right)^2 + \left(1 - Corr[\hat{\lambda}, \lambda]\right)^2 \\ &\quad + (\ln \bar{\rho}^D - \ln \bar{\rho}^M)^2 + (\ln \bar{V}_y^D - \ln \bar{V}_y^M)^2 \\ V_{MAT}(0) &= \{(V_a(i), V_m(i)) : \forall i\} \end{aligned}$$

13. Let $\{(V_a(i), V_m(i)) : \forall i\} = V_{MAT}$; $Params = Params$ except $\alpha = \min\{0.9, \max\{0.1, \alpha \times (1 + Inc)\}$
For $i = 1$ to 701, choose $(V_a(i), V_m(i))$ to minimise

$$\begin{aligned} Obj(1) &= (1 - Corr[\hat{\rho}, \rho])^2 + \left(1 - Corr[\hat{V}_y, V_y]\right)^2 + \left(1 - Corr[\hat{\lambda}, \lambda]\right)^2 \\ &\quad + (\ln \bar{\rho}^D - \ln \bar{\rho}^M)^2 + (\ln \bar{V}_y^D - \ln \bar{V}_y^M)^2 \\ V_{MAT}(1) &= \{(V_a(i), V_m(i)) : \forall i\} \end{aligned}$$

14. Let $\{(V_a(i), V_m(i)) : \forall i\} = V_{MAT}$; $Params = Params$ except $\alpha = \min\{0.9, \max\{0.1, \alpha / (1 + Inc)\}$
For $i = 1$ to 701, choose $(V_a(i), V_m(i))$ to minimise

$$\begin{aligned} Obj(2) &= (1 - Corr[\hat{\rho}, \rho])^2 + \left(1 - Corr[\hat{V}_y, V_y]\right)^2 + \left(1 - Corr[\hat{\lambda}, \lambda]\right)^2 \\ &\quad + (\ln \bar{\rho}^D - \ln \bar{\rho}^M)^2 + (\ln \bar{V}_y^D - \ln \bar{V}_y^M)^2 \\ V_{MAT}(2) &= \{(V_a(i), V_m(i)) : \forall i\} \end{aligned}$$

15. Let $\{(V_a(i), V_m(i)) : \forall i\} = V_{MAT}$; $Params = Params$ except $V_\omega = \max\{0.001, V_\omega \times (1 + Inc)\}$;
For $i = 1$ to 701, choose $(V_a(i), V_m(i))$ to minimise

$$\begin{aligned} Obj(3) &= (1 - Corr[\hat{\rho}, \rho])^2 + \left(1 - Corr[\hat{V}_y, V_y]\right)^2 + \left(1 - Corr[\hat{\lambda}, \lambda]\right)^2 \\ &\quad + (\ln \bar{\rho}^D - \ln \bar{\rho}^M)^2 + (\ln \bar{V}_y^D - \ln \bar{V}_y^M)^2 \\ V_{MAT}(3) &= \{(V_a(i), V_m(i)) : \forall i\} \end{aligned}$$

16. Let $\{(V_a(i), V_m(i)) : \forall i\} = V_{MAT}$; $Params = Params$ except $V_\omega = \max\{0.001, V_\omega / (1 + Inc)\}$;
For $i = 1$ to 701, choose $(V_a(i), V_m(i))$ to minimise

$$\begin{aligned} Obj(4) &= (1 - Corr[\hat{\rho}, \rho])^2 + \left(1 - Corr[\hat{V}_y, V_y]\right)^2 + \left(1 - Corr[\hat{\lambda}, \lambda]\right)^2 \\ &\quad + (\ln \bar{\rho}^D - \ln \bar{\rho}^M)^2 + (\ln \bar{V}_y^D - \ln \bar{V}_y^M)^2 \\ V_{MAT}(4) &= \{(V_a(i), V_m(i)) : \forall i\} \end{aligned}$$

17. If $\min\{Obj(j) : j = 0, \dots, 4\} = Obj(0)$ then

No change to $Params$, let $Inc = \frac{1}{10}Inc$ and let $V_{MAT} = V_{MAT}(0)$; goto 12. so long as $Inc > 0.01\%$

Else if $\min\{Obj(j) : j = 0, \dots, 4\} = Obj(k), k \neq 0$

Make appropriate change to $Params$ and let $V_{MAT} = V_{MAT}(k)$; goto 12.

18. For $i = 1$ to 701, choose $(V_a(i), V_m(i))$ to minimise

$$\begin{aligned} Obj &= (1 - Corr[\hat{\rho}, \rho])^2 + \left(1 - Corr[\hat{V}_y, V_y]\right)^2 + \left(1 - Corr[\hat{\lambda}, \lambda]\right)^2 \\ &\quad + (\ln \bar{\rho}^D - \ln \bar{\rho}^M)^2 + (\ln \bar{V}_y^D - \ln \bar{V}_y^M)^2 \end{aligned}$$

19. Repeat 18. until Obj is stable

20. Choose μ & σ so that the mean and standard deviation of the $\hat{\lambda}$ matches the modelled equivalents, subject to the constraint that $\min \left\{ \hat{\lambda}_i \right\} \geq 0$.