

Identification in a Search and Bargaining Model with OTJ Training

Christopher Adjaho

May 6 2021

Abstract

In this paper, I ask the following two questions: (i) How important are wage and employment transition data for recovering reliable estimates of the worker ability and worker-firm productivity distributions? (ii) How can we measure the sensitivity of estimated structural parameters to calibrated parameters? Using a search model with on-the-job training as in Flinn, Gemici & Laufer (2017) I find that data characterising the wage distribution of job-movers is relatively more important for recovering precise estimates.

1 Introduction

This paper asks the following two questions: (i) what is the relative importance of data characterising the wage distribution to data characterising employment transitions for recovering precise estimates of the worker-firm and worker ability distributions? (ii) how sensitive are estimated structural parameters and counterfactual policies to calibrated parameters? Using a search and bargaining model that allows for human capital accumulation on-the-job, I try to understand the sensitivity of estimates to different moments and find that absent worker-firm productivity data, moments characterising the wage distribution of job-movers are important for recovering reliable estimates. The challenge of identifying parameters of a distribution within search models in the absence of such data is not uncommon. For example, Karahan, Ozkan, and Song (2019) study the importance of lifetime earnings growth on lifetime inequality in the U.S within a search and bargaining model with OTJ search and heterogeneity in worker ability. In their estimation, a key challenge is identifying the distribution of firm productivity using moments related to the wage growth distribution and employment transitions over the life-cycle. While there have been suggestions for dealing with this problem such as Bagger and Lentz (2019) who propose using poaching patterns between firms to rank them with respect to their productivity, the problem tackled in this paper does not feature firm productivity but instead is

concerned with the joint worker-firm match productivity. The model is found in Flinn, Gemici, and Laufer (2017), hereafter FGL. Unique to their paper, the authors use job training data to facilitate estimation of the worker-firm and worker ability distributions and conclude that employment transition data alone is not sufficient for recovering these parameters, but that one requires data related to the distribution of wages both within and across job spells.

Within the econometrics literature, there has been a series of efforts to better understand sensitivity of parameter estimates within dynamic structural models using moments-based estimation. Andrews, Gentzkow, and Shapiro (2017) study asymptotic bias in the presence of misspecified moments, while Honoré, Jørgensen, and de Paula (2020) study the effect on asymptotic variance to included moments. As interest is often in counterfactual policy, sensitivity of such results has also been of concern. Christensen and Connault (2019) study sensitivity to assumptions on unobserved distributions, while Jørgensen (2020) studies the effect of parameters fixed in estimation. Employing some of the methods developed in this literature may help in solving problems for users of these structural models, such as those mentioned in the previous paragraph.

Section 2, outlines the partial equilibrium model of FGL (2017). Section 3 briefly describes the data used in estimating the model. Section 4 discusses estimation and model fit to the data and proximity to the results presented in FGL (2017). In Section 5, I discuss two sensitivity measures that might shed light on which moments are most informative in estimation, and hence whether the authors' claims are verified. Section 6 considers sensitivity of the structural parameter estimates and a minimum wage policy to the calibrated bargaining power parameter, and Section 7 concludes.

2 Model

The model is set in continuous time. Individuals are characterized in terms of a (general) ability level a , with which they enter the labor market:

$$0 < a_1 < \dots < a_M.$$

When an individual of type a_i encounters a firm, he draws a value θ from the discrete distribution over the K values of match productivity θ given by

$$0 < \theta_1 < \dots < \theta_K.$$

Denote the cdf of θ by G and define $p_j := \Pr(\theta = \theta_j)$. The flow productivity of a match is

$$y(i, j) = a_i \theta_j.$$

Both general ability and match productivity can be changed through investment on the job. Indeed, training times allocated to general training, τ_a , and match-specific training, τ_θ , is determined coop-

eratively using a Nash bargaining process. The only cost of either type of training is time, so total productivity is

$$(1 - \tau_a - \tau_\theta)y(i, j).$$

The production technology for increasing general productivity is

$$\varphi_a(i, \tau_a) = \delta_a^0 \times a_i^{\delta_a^1} \times (\tau_a)^{\delta_a^2}$$

and it is assumed that $\varphi_a(M, \tau_a) = 0$ for all τ_a so that an individual of the highest ability cannot experience further increase. Similarly, we have

$$\varphi_\theta(j, \tau_\theta) = \delta_\theta^0 \times \theta_j^{\delta_\theta^1} \times (\tau_\theta)^{\delta_\theta^2}$$

with $\varphi_\theta(K, \tau_\theta) = 0, \forall \tau_\theta$. An individual may also experience skill depreciation: when $i > 1$, he is subject to Poisson shocks that arrive at exogenous rate $\tilde{\varphi}_a^-$. When such a shock arrives, his new skill level is a_{i-1} . Similarly, when $j > 1$ he is subject to $\tilde{\varphi}_\theta^-$ and his new skill level is reduced to θ_{j-1} .

2.1 Bargaining problem

The model allows for on-the-job search and it is assumed that there are no opportunities for the worker and firm to renegotiate the employment contract once an outside offer is received. That is, the outside option for the worker is always unemployment. This gives

$$\begin{aligned} \max_{w, \tau_a, \tau_\theta} \quad & \left(\tilde{V}_E(i, j; w, \tau_a, \tau_\theta) - V_U(i) \right)^\alpha \left(\tilde{V}_F(i, j; w, \tau_a, \tau_\theta) - 0 \right) \\ \text{s.t.} \quad & \tau_a + \tau_\theta \leq 1. \end{aligned} \tag{1}$$

It is also assumed the value to the firm of an unfilled vacancy is 0. The value to an employed worker with contract (w, τ_a, τ_θ) given state variables (a_i, θ_j) is

$$\tilde{V}_E(i, j; w, \tau_a, \tau_\theta) = \frac{N_E(w, \tau_a, \tau_\theta; i, j)}{D(\tau_a, \tau_\theta; i, j)} \tag{2}$$

where

$$\begin{aligned} N_E(w, \tau_a, \tau_\theta; i, j) = & w + \lambda_E \sum_{s=j+1} p_s V_E(i, s) \\ & + \varphi_a(i, \tau_a) Q(i+1, j) + \varphi_\theta(j, \tau_\theta) V_E(i, j+1) \\ & + \tilde{\varphi}_a^- Q(i-1, j) + \tilde{\varphi}_\theta^- Q(i, j-1) \\ & + \eta V_U(i) \end{aligned} \tag{3}$$

and

$$D(\tau_a, \tau_\theta; i, j) = \rho + \lambda_E \tilde{G}(\theta_j) + \varphi_a(i, \tau_a) + \varphi_\theta(j, \tau_\theta) + \tilde{\varphi}_a^- + \tilde{\varphi}_\theta^- + \eta. \tag{4}$$

with $\tilde{G}(\theta_j) = 1 - G(\theta_j)$. The term $Q(i+1, j) = \max\{V_E(i+1, j), V_U(i+1)\}$ and reflects the fact that at a higher level of a , a_{i+1} , the match value θ_j may no longer be in the acceptance set of employment contracts, in which case the worker becomes an unemployed searcher. There may also be endogenous separation as reflected by $Q(i, j-1)$. Finally, the match can also be terminated at the exogenous rate η .

The value to the firm conditional on the wage and investment decisions is given by

$$\tilde{V}_F(i, j; w, \tau_a, \tau_\theta) = \frac{N_F(w, \tau_a, \tau_\theta; i, j)}{D(\tau_a, \tau_\theta; i, j)}, \quad (5)$$

where

$$\begin{aligned} N_F(w, \tau_a, \tau_\theta; i, j) = & (1 - \tau_a - \tau_\theta)y(i, j) - w + \varphi_a(i, \tau_a)Q_F(i+1, j) \\ & + \varphi_\theta(j, \tau_\theta)V_F(i, j+1) \\ & + \tilde{\varphi}_a^- Q_F(i-1, j) \\ & + \tilde{\varphi}_\theta^- Q_F(i, j-1), \end{aligned} \quad (6)$$

with $Q_F(i+1, j) = V_F(i+1, j)$ if $V_E(i+1, j) > V_U(i+1)$ and equals 0 otherwise, and $Q_F(i, j-1) = V_F(i, j-1)$ if $j-1 > r^*(i)$ where $r^*(i)$ is the reservation match value for a worker of ability a_i .

Finally, the value of unemployed search is

$$V_U(i) = \frac{ba_i + \lambda_U \sum_{j=r^*(i)+1} p_j V_E(i, j)}{\rho + \lambda_U \tilde{G}(\theta_{r_i^*})} \quad (7)$$

where ba_i is the flow value of unemployment to an individual of type a_i .

3 Data

I use the same data and sample FGL use in estimation: a sub-sample of the National Longitudinal Survey of Youth 1997 (NLSY97). The sample consists of males between ages 26-32 in the latest 2011 survey round who have completed schooling. There is an unbalanced panel of 1,994 individuals and 742,528 person-week observations. The proportion of high school graduates is 37 percent, while the proportion of those with some college and those with a college degree are 30 and 33 percent respectively.

While weekly employment and training information is used, wage observations are taken from respondents' reported current wage at the time of the interview. Finally, the training data used is taken from the NLSY97 training roster where respondents are asked about what types of training they receive over the survey year and about the start and end dates of training periods by source of training. However, FGL (2017) do not make assumptions about the type of human capital acquired.

4 Estimation

4.1 Implementation

To remain close to the authors, I set $M = K = 24$ and make the support of the distributions for a and θ identical. The grid points are spaced logarithmically from 2.5 standard deviations below the mean to 3.5 standard deviations above, and at my estimates shown below, moving up one grid point corresponds to a 7 percent increase in productivity. I also allow training times to be chosen from a grid with 26 points.

Furthermore, workers in the model almost always spend some fraction of their time training, which is not the case in the data. To relate the two measures, I also assume that a worker who spends a fraction of time τ engaged in training is observed to receive training in that period with probability

$$\Pr(\text{Training observed}|\tau) = \Phi(\beta_0 + \beta_1\tau) \quad (8)$$

where Φ is the cdf for the normal distribution and $\tau = \tau_a^{ave} + \tau_\theta^{ave}$ is the average fraction of time spent training in a 13-week window or over the entire job spell if it lasts less than 13 weeks.

As in FGL (2017), there are 23 parameters, of which 19 are estimated and 4 are fixed in estimation. The latter include the bargaining parameter α which is set to 0.5, and the standard deviation of log wage measurement error, σ_w set to 0.15. All rate parameters are expressed at a weekly frequency, with $\rho = 0.00203$ which accounts for both a time discount rate of four percent and a death shock that produces an average career length of 45 years. Finally, β_1 is set to 1.

4.2 Identification

Estimation of the 19 remaining parameters is achieved by Simulated Method of Moments (SMM). I use a total of 235 moments and to stay close to the authors, almost all of the moments are employed in FGL (2017). Some of these moments are listed in Appendix A. We can think of identification as containing two parts: The first is based on proving the moments used in estimation allow recovery of the parameters of interest ¹. Given the overlap between the moments I've chosen and those of the authors, I will forgo this part. The second is more practical. Indeed, because the problem of identification of the parametric structural model is also very much a numerical and data-driven issue, I try to build some confidence in the estimates reported in Table 1. I use the following procedure: Given a parameter iteration in the Nelder-Mead algorithm, Γ^k ,

1. Solve the model

¹This will most likely involve an informal discussion rather than a mathematical proof, as is often the case in moment-based estimation of dynamic models, and as the authors make clear.

2. Simulate a panel data of weekly histories as in the actual data, using the same sample-selection criteria
3. Evaluate the objective function.

I then repeat until the objective function value has stabilized – although not necessarily converged – and re-run the best stabilized point as the initial vector in the Nelder-Mead process. Having arrived at the estimates $\hat{\Gamma}$ in Table 1, I then generate 2000 uniform quasi-random Sobol points in the parameter subspace $[0.5\hat{\Gamma}_j, 1.5\hat{\Gamma}_j]$ as well as 100 Sobol points in the subspace $[0.95\hat{\Gamma}_j, 1.05\hat{\Gamma}_j]$ and 100 in $[0.98\hat{\Gamma}_j, 1.02\hat{\Gamma}_j]$, each time comparing the objective function value of the best point in each subspace to its value at $\hat{\Gamma}$. My reason for using Sobol points is to provide a “more uniform” search over the parameter space for local minima. Indeed, consider Figure 1 in the Appendix. In this simple example of searching over a $[0, 1]^2$ parameter space, I generate 1000 random Uniform(0,1) points and 1000 Sobol points. As can be seen, random number generation can have a tendency to cluster and leave parts of the space unexplored.

Let m_k , $k = 1, \dots, K$ denote an empirical moment and let $m_k(\Gamma)$ be the corresponding model moment that is simulated for a given parameter vector Γ . To deal with issues relating to the large variation in the scales of the moments, I minimize the arc-percentage deviation of model-implied moment $m_k(\Gamma)$ from target m_k . That is, my SMM estimator is defined as

$$\hat{\Gamma} = \arg \min_{\Gamma} \mathbf{F}(\Gamma)' W \mathbf{F}(\Gamma), \quad \mathbf{F}(\Gamma) = [F_1(\Gamma), \dots, F_K(\Gamma)], \quad (9)$$

and

$$F_k(\Gamma) = \frac{m_k(\Gamma) - m_k}{0.5(|m_k(\Gamma)| + |m_k|)}. \quad (10)$$

The reason for using arc-percent deviation instead of percent deviation is because some of the empirical moments used in estimation are zero or close to zero, and so using (10) eliminates this problem. As seen in Appendix A, we can categorize the moments into three broad groups: employment, wages, and training moments. The weighting matrix W I’ve used is diagonal and (given moments are independent of scale) assigns 10% weight to the training data moments, and 90% to both employment and wage data moments². I’ve done this based on my subjective belief about the importance of each moment and because the choice of weighting matrix should not affect consistency of the estimator – only precision. I have, however, estimated the model using the simple deviation of model moments from target and with a weighting matrix where the weights are the inverse of the variance of the empirical moments and obtain the same result. But the point here is that given some of the most important moments also happen to have the most noise (discussed in the next section) and one is trying to learn

²More precisely, with 4 training moments, each is given weight 0.1/4 and with a remaining 231 moments for employment and wages, each is given weight 0.9/231.

of which moments to use for identification, it is not necessarily the case that one would like to use “optimal” weighting.

Table 1 compares the estimates, $\hat{\Gamma}$, I obtain and the ones reported in FGL (2017). Despite some differences, the qualitative story is the same: (i) the job offer rate for the unemployed is higher than for the employed, (ii) general training becomes less productive the higher is worker ability and (iii) the mean initial ability estimates is increasing in education. In addition, and as seen in Figure 2, the fraction of time spent training decreases with time, and initially increases for match-specific training because $\hat{\delta}_\theta^1 > 0$. However, a noticeable difference is the contact rate of workers with firms. My estimates suggests that unemployed workers receive offers on average once every two weeks ($1/0.4585$) while employed workers receive offers on average once every three weeks ($1/0.373$). The estimate of β_0 I obtain is similar to the authors’ and suggests that an individual who spends his whole time training is observed to be receiving training with probability 0.03.

4.3 Model Fit

Table 2 compares the fit of the model at both mine and the authors’ parameter estimates. Panel A suggests the estimates I obtain fit employment transition data levels quite well for each education level, but doesn’t capture the positive relationship between job-stayer probability and education, and the negative relationship between job-mover probability and education. Turning to Panel B and C, my estimates, like FGL (2017), are able to capture the effect of an intervening non-employment spell on the wage growth of job-movers. In the data, wage growth is 0.03 for those with non-employment (EUE) and 0.12 for those without (EE), while my estimates are 0.03 and 0.25 respectively. As in FGL (2017), one can see that fitting the levels of these wage moments for job-to-job transitions is challenging, which is problematic because the authors believe wage moments to be essential for identifying the productivity distributions. It turns out that the wage moments relating to job-movers have the fewest observations and high variance³. Furthermore, let g_n denote the vector of moments with expectation g_0 , and let $f_R(\Gamma)$ denote the simulated vector so that under correct specification $g_0 = f(\Gamma_0)$ when $R \rightarrow \infty$. We know that

$$\sqrt{n}(g_n - f_R(\Gamma_0)) \rightarrow^d N\left(0, \left(1 + \frac{1}{R}\right)S\right), \quad (11)$$

where S is the variance-covariance matrix of the empirical moments. Given asymptotic normality, we know that for such noisy wage moments, we can build confidence intervals that are likely to be wide. Thus, while it is challenging fitting the levels for these moments, the model-implied estimates are not

³Indeed, take wage growth for those who completed college and experienced an intervening non-employment spell in Panel B. The data suggests wage growth of 0.23, but there are only 22 observations, and the variance is 0.16, implying a mean-to-standard deviation of 0.59.

Parameter estimates					
PARAMETERS FOR EMPLOYMENT TRANSITION		Authors	Mine	Authors	Mine
flow value of unemployment	b	6.4705	9.8965	(0.857)	(0.942)
job offer rate - unemployed	λ_u	0.1375	0.4585	(0.019)	(0.068)
job offer rate - employed	λ_e	0.0685	0.3730	(0.023)	(0.050)
exogenous job separation rate	η	0.0036	0.0007	(0.001)	(0.0003)
PARAMETERS OF INVESTMENT FUNCTIONS					
general ability investment TFP	$\delta_{a,0}$	0.0285	0.0430	(0.005)	(0.004)
firm-specific investment TFP	$\delta_{\theta,0}$	0.0229	0.0824	(0.003)	(0.019)
state-dependence of general ability investment	$\delta_{a,1}$	-0.1353	-0.1697	(0.144)	(0.041)
state-dependence of firm-specific investment	$\delta_{\theta,1}$	0.4426	0.6822	(0.098)	(0.124)
curvature of general ability investment	$\delta_{a,2}$	0.2826	0.3177	(0.084)	(0.057)
curvature of firm-specific investment	$\delta_{\theta,2}$	0.4426	0.6150	(0.062)	(0.112)
rate of decrease in general ability	$\tilde{\varphi}_a^-$	0.0014	0.0069	(0.0003)	(0.002)
rate of decrease in match quality	$\tilde{\varphi}_\theta^-$	0.0113	0.1119	(0.003)	(0.018)
PARAMETERS OF INITIAL ABILITY DISTRIBUTION					
mean of initial general ability - High school	$\mu_{a,1}$	0.8844	0.8921	(0.159)	(0.239)
mean of initial general ability -Some college	$\mu_{a,2}$	1.1515	1.1426	(0.112)	(0.138)
mean of initial general ability - BA or higher	$\mu_{a,3}$	1.4489	1.3728	(0.083)	(0.1313)
variance of initial ability	σ_a	0.2041	0.2586	(0.086)	(0.056)
PARAMETERS OF JOB OFFERS					
mean of match quality distribution	μ_θ	1.4078	1.5013	(0.040)	(0.088)
variance of match quality distribution	σ_θ	0.2895	0.2589	(0.034)	(0.030)
PARAMETERS GOVERNING TRAINING OBSERVATION					
intercept for training observation	β_0	-2.7331	-2.8198	(0.039)	(0.902)
coefficient on τ for training observation	β_τ	1.0000	1.0000	—	—

Table 1: Parameter estimates

unreasonable and often would fall within such intervals.

Table 3 shows the fit of the model to the training data. In total, the model is able to provide a good fit to both the level and direction, with incidence of training decreasing in education.

Finally, Table 4 shows the fit of the model against untargeted moments – that is, moments I did not use in estimation. This table was not reported in FGL (2017) and so I can only compare my estimates to the data. However, given the similarities up to now of my estimates to theirs, it might serve as some approximation. Since we use wage observations at each interview date, one can define the length of a job spell by the number of interview dates it spans. Table 4 shows that, as in the data, the length of a job spell in the model is increasing in education. We see that the proportion of job spells spanning 3 or more interviews is 9, 12 and 15 percent for HS, SC and College respectively. In the model it is 18, 20, 24 percent.

5 Informative moments

Up to now, I have discussed in detail how I obtained my estimates in order to build some confidence that they might be considered a minimum, and have compared the fit of the model to what is reported in FGL (2017). The goal of this exercise was to build confidence that what I obtained might serve as a close-enough approximation to what the authors report. Accepting this, we can now discuss the relative importance of the different groups of moments for recovering precise estimates of the structural parameters. In particular, we recall the following from FGL (2017):

“by utilizing some training data, albeit of questionable quality, we can get some additional information on the values of (a, θ) ... That is, our model produces a mapping from the state variables $(a, \theta) \rightarrow (\tau_a, \tau_\theta)$. While this mapping is not, in general, invertible, it still conveys information on the set of values (a, θ) consistent with the reported training time of the individual.” (274).

In this section, I will focus on two parts of the preceding extract: (i) the quality of the training data, and (ii) its informativeness with respect to (a, θ) . Indeed, the training data used is likely to be noisy and a rough approximation to the truth. Following Honoré et al. (2020) consider the following question: How much precision would we lose if the k -th moment is subject to a little additional noise? Recall the SMM standard errors:

$$\begin{aligned}\Sigma &= \left(1 + \frac{1}{R}\right) (D'WD)^{-1} D'WSWD (D'WD)^{-1} \\ &= \left(1 + \frac{1}{R}\right) M_1 S M_1',\end{aligned}\tag{12}$$

Model fit: annual labor turnover rates and wage growth

	HS	Some college	College or more
Panel A: Employment-to-Employment (EE) transitions btw $t - 1$ and t			
% of EE transitions <i>with no</i> job-change			
.....Data	81%	85%	88%
.....Authors	74%	74%	76%
.....Mine	88%	84%	82%
% of EE transitions <i>with</i> job change (job-to-job transitions)			
.....Data	19%	15%	12%
.....Authors	25%	26%	24%
.....Mine	12%	16%	18%
% of job-to-job transitions <i>with</i> non-employment btw $t - 1$ and t			
.....Data	4%	3%	2%
.....Authors	8%	8%	8%
.....Mine	4%	7%	8%
% of job-to-job transitions <i>with no</i> non-employment btw $t - 1$ and t			
.....Data	15%	12%	10%
.....Authors	17%	17%	16%
.....Mine	7%	9%	10%
Panel B: Wage Growth btw $t - 1$ and t			
EE transitions with no job change			
.....Data	0.08	0.08	0.09
.....Authors	0.08	0.08	0.08
.....Mine	0.13	0.13	0.15
EE transitions with job change			
.....Data	0.11	0.15	0.20
.....Authors	0.09	0.11	0.11
.....Mine	0.17	0.15	0.11
job-to-job transitions with non-employment btw $t - 1$ and t			
.....Data	0.06	0.06	0.23
.....Authors	-0.10	-0.06	-0.02
.....Mine	0.03	0.05	0.03
job-to-job transitions with no non-employment btw $t - 1$ and t			
.....Data	0.12	0.17	0.20
.....Authors	0.18	0.19	0.17
.....Mine	0.25	0.22	0.18
Panel C: % of Negative Wage Growth btw $t - 1$ and t			
EE transitions <i>with no</i> job change			
.....Data	17%	19%	23%
.....Authors	38%	37%	37%
.....Mine	40%	40%	38%
EE transitions <i>with</i> job change			
.....Data	32%	32%	28%
.....Authors	37%	35%	34%
.....Mine	40%	39%	44%
job-to-job transitions with non-employment btw $t - 1$ and t			
.....Data	39%	47%	27%
.....Authors	61%	57%	50%
.....Mine	51%	47%	49%
job-to-job transitions with no non-employment btw $t - 1$ and t			
.....Data	30%	28%	28%
.....Authors	26%	24%	26%
.....Mine	34%	34%	39%

Table 2: Annual labor turnover rates and wage growth (Targeted)

Incidence of training				
	All	HS	Some col- lege	College or more
% who got training at least once				
.....Data	15%	18%	13%	13%
.....Authors	17%	21%	16%	12%
.....Mine	12%	16%	11%	9%
% who got training at the start of job spell				
.....Data	6%	10%	5%	3%
.....Authors	5%	6%	4%	3%
.....Mine	4%	5%	5%	3%

Table 3: Incidence of training (Targeted)

Proportion of job spells by the number of interview dates they span				
		HS	Some college	College or more
0	Data	61%	55%	53%
	Mine	64%	61%	55%
1	Data	23%	24%	21%
	Mine	12%	12%	14%
2	Data	7%	9%	11%
	Mine	6%	6%	7%
3	Data	3%	4%	6%
	Mine	5%	4%	5%
4	Data	2%	3%	3%
	Mine	4%	3%	3%
5	Data	1%	2%	3%
	Mine	3%	2%	2%
6	Data	3%	3%	3%
	Mine	6%	11%	14%

Table 4: Length of job spell (Untargeted)

where R is the number of replications used in simulating the model – which I’ve set to 10. The $K \times p$ matrix $D := \mathbb{E} \left[\frac{\partial g(\Gamma)}{\partial \Gamma'} \middle| \Gamma = \Gamma_0 \right]$ is the Jacobian of the moment function $g(\Gamma)$ with respect to the structural parameters evaluated at the true value Γ_0 , and S is, as mentioned in (11), the variance matrix of the empirical moments. Informally, we can see then that when the variance of the moments S increases by 1, the asymptotic variance of the structural parameter estimates increases by M_1^2 . More formally, we can define

$$\mathcal{E}^{(j,k)} := \underbrace{\frac{\partial \Sigma^{j,j}}{\partial S^{k,k}}}_{M_1 O_{kk} M_1'} \frac{S^{k,k}}{\Sigma^{j,j}} \quad (13)$$

where O_{kk} is a $K \times K$ matrix with a 1 in entry (k, k) and zeros elsewhere. The measure (13) is an elasticity of the percentage change in the precision of the j -th parameter, σ_j^2 , to a percentage change in the variance of the k -th moment. Table 5 shows the moments each parameter is most responsive to according to the sensitivity measure (13) and computation of the standard errors can be found in Appendix B. It is immediate that moments characterising the distribution of wages for job-movers are important for the precision of parameter estimates. There is one exception: the parameter governing training observation, $\hat{\beta}_0$. We see that a one percent increase in the noise of the moment “fraction receiving training in current job” leads to a 70% increase in its asymptotic variance.

An individual is classed as ‘EE’ if they were employed at both interview dates $t - 1$ and t , with a different employer at t and have spent less than five weeks unemployed between interview dates. An individual is ‘EUE’ if they have spent more than five weeks unemployed between interview dates.

The moment “ w_{t-1} , 1 change, EE, ed=1” refers to the mean wage at interview date $t - 1$ among those individuals who (i) have completed high school (ii) had a job-to-job transition with an intervening unemployment spell of less than five weeks and (iii) had exactly one job change from $t - 1$ to t .

The moment “ w_t , > 1 change, EE, ed=1” refers to the mean wage at interview date t among individuals who (i) have completed HS (ii) are EE (iii) changed jobs more than once from $t - 1$ to t .

The moment “ w_1 , two wage obs., ed=2” refers to the mean of the first wage across job spells having exactly two wage observations for individuals with some college.

The moment “ w , 6-8yrs, ed=3” refers to the mean wage among individuals who have been in the labor market for 6-8 years and completed college.

Consider now the following question: How does the asymptotic variance change from completely excluding the k -th moment? Let

$$M_{2,k} = \tilde{\Sigma}_k - \Sigma, \quad (14)$$

Sensitivity 1: Additional noise		
	Moment	Percent
b	$w_t, > 1$ change, EE, ed=1	8.63
λ_u	$w_{t-1}, > 1$ change, EE, ed=2	5.15
λ_e	$w_t, > 1$ change, EE, ed=1	9.00
η	w_1^2 , two wage obs., ed=2	4.32
$\delta_{a,0}$	w_{t-1}^2 , 1 change, EE, ed=1	4.53
$\delta_{\theta,0}$	w_t^2 , 1 change, EUE, ed=3	5.83
$\delta_{a,1}$	$w_{t-1}, > 1$ change, EE, ed=2	6.94
$\delta_{1,\theta}$	w , 6-8yrs, ed=3	4.18
$\delta_{2,a}$	$w_{t-1}^2, > 1$ change, EUE, ed=1	9.84
$\delta_{2,\theta}$	$w_t^2, > 1$ change, EUE, ed=2	10.9
$\tilde{\phi}_a^-$	$w_t^2, > 1$ change, EUE, ed=3	4.58
$\tilde{\phi}_\theta^-$	w_{t-1}^2 , 1 change, EUE, ed=1	8.13
$\mu_{a,1}$	w_{t-1}^2 , 1 change, EUE, ed=1	5.05
$\mu_{a,2}$	w_1 two wage obs., ed=2	3.49
$\mu_{a,3}$	$w_{t-1}^2, > 1$ change, EE, ed=2	4.42
σ_a	w_{t-1}^2 , 1 change, EUE, ed=2	6.10
μ_θ	$w_{t-1}, > 1$ change, EE, ed=2	6.76
σ_θ	$w_{t-1}, > 1$ change, EE, ed=2	8.95
β_0	fraction receiving training in current job	70.0

Table 5: Which moment is each parameter most responsive to?

where

$$\tilde{\Sigma}_k = \left(1 + \frac{1}{R}\right) \left(D' \tilde{W}_k D\right)^{-1} D' \tilde{W}_k S \tilde{W}_k D \left(D' \tilde{W}_k D\right)^{-1} \quad (15)$$

$$\tilde{W}_k = W \odot (\iota_k \iota_k'), \quad (16)$$

ι_k is a $K \times 1$ vector with ones in each entry but zero in the k -th entry, and $\odot$ denotes element-wise multiplication. Essentially, the new weighting matrix $\tilde{W}_k$ excludes the k -th moment by keeping all entries of W but setting each element in row and column k to zero. It is clear that we need $K > p$ for the model to remain identified. Similar to before, we can make our measure scale-invariant by focusing on percentage changes:

$$\Delta^{(j,k)} = \frac{M_{2,k}^{j,j}}{\Sigma^{j,j}}, \quad (17)$$

which measures the percentage change in the asymptotic variance of the j -th parameter from excluding moment k . Table 6 shows the moments each parameter is most responsive to according to the sensitivity measure (17). Firstly, it is worth noting that the moments listed are almost identical to Table 5, and so it is still the case that moments characterising the wage distribution of job-movers are important for precision. Secondly, we see just how important training data is for β_0 : Excluding the training moment now leads to a 341% increase in its variance.

Sensitivity 2: Excluding moment		
	Moment	Percent
b	$w_t, > 1$ change, EE, ed=1	12.8
λ_u	$w_{t-1}, > 1$ change, EE, ed=2	9.14
λ_e	$w_t, > 1$ change, EE, ed=1	13.3
η	w_t^2 , two wage obs., ed=2	5.50
$\delta_{a,0}$	$w_{t-1}, > 1$ change, EE, ed=2	7.75
$\delta_{\theta,0}$	w_t , 1 change, EUE, ed=3	7.83
$\delta_{a,1}$	$w_{t-1}, > 1$ change, EE, ed=2	12.3
$\delta_{1,\theta}$	w , 6-8yrs, ed=3	5.56
$\delta_{2,a}$	$w_{t-1}^2, > 1$ change, EUE, ed=1	15.0
$\delta_{2,\theta}$	$w_t^2, > 1$ change, EUE, ed=2	15.6
$\tilde{\phi}_a^-$	$w_{t-1}, > 1$ change, EE, ed=1	7.00
$\tilde{\phi}_\theta^-$	w^2 before 1 change, EUE, ed=1	12.6
$\mu_{a,1}$	w_{t-1}^2 , 1 change, EUE, ed=1	7.84
$\mu_{a,2}$	w_{t-1}^2 , 1 change, EUE, ed=1	5.23
$\mu_{a,3}$	$w_{t-1}^2, > 1$ change, EE, ed=2	6.26
σ_a	w_{t-1}^2 , 1 change, EUE, ed=2	8.41
μ_θ	$w_{t-1}, > 1$ change, EE, ed=2	12.0
σ_θ	$w_{t-1}, > 1$ change, EE, ed=2	15.9
β_0	fraction receiving training in current job	341

Table 6: Which moment is each parameter most responsive to?

Now consider the effect of adding noise to a group of moments which, given similarities between the two tables, might indicate the importance of excluding the same group from estimation. The extract from FGL (2017) at the beginning of this section suggested that the inclusion of training data facilitates estimation of the (a, θ) distribution. The paper, however, does not provide the reader with how this might be so; instead, relying on the theoretical mapping from the parameters of interest to the training data. We can use the previous sensitivity measures to formalise this. The total differential of the asymptotic variance w.r.t moments $k, \dots, k+l$ is:

$$d\Sigma^{j,j} = \frac{\partial \Sigma^{j,j}}{\partial S^{k,k}} dS^{k,k} + \dots + \frac{\partial \Sigma^{j,j}}{\partial S^{k+l,k+l}} dS^{k+l,k+l} \quad (18)$$

which implies

$$\begin{aligned} \frac{d\Sigma^{j,j}}{\Sigma^{j,j}} &= \frac{\partial \Sigma^{j,j}}{\partial S^{k,k}} \frac{dS^{k,k}}{\Sigma^{j,j}} + \dots + \frac{\partial \Sigma^{j,j}}{\partial S^{k+l,k+l}} \frac{dS^{k+l,k+l}}{\Sigma^{j,j}} \\ &= \mathcal{E}^{(j,k)} + \dots + \mathcal{E}^{(j,k+l)}. \end{aligned} \quad (19)$$

In estimation, I have used a total of four training moments. Consider the effect of adding noise to each these moments on the precision of $\hat{\mu}_\theta$ and $\hat{\sigma}_\theta$. Table 7 shows that the total effect on the mean and standard deviation is 0.01% and 0.05% respectively. Therefore, it seems that such moments have very little effect on the the precision of our estimates of the worker-firm and worker ability distributions.

Sensitivity to groups of moments		
	$\hat{\mu}_\theta$	$\hat{\sigma}_\theta$
Training	0.01%	0.05%
EE ($w_t - w_{t-1}$)	20.1%	21.3%
EUE ($w_t - w_{t-1}$)	9.95%	4.60%
EE	0.431%	0.599%
EUE	1.68%	1.90%

Table 7: Sensitivity of θ distribution to groups of moments

On the other hand, the total effect from the group of moments that make up wage growth between interview dates for EE individuals is 20% for $\hat{\mu}_\theta$ and 21% for $\hat{\sigma}_\theta$. Finally, the total effect from the group of moments relating to the fraction of EE workers is 0.431% and 0.599% respectively. It seems then, and as mentioned in FGL (2017), that inclusion of moments relating to wage growth more so than employment transition, are important for recovering reliable estimates of the distribution parameters.

6 Sensitivity of counterfactual

In this section, we consider the sensitivity of the unemployment rate when a minimum wage is binding to the calibrated bargaining parameter. As in FGL (2017), I introduce a \$15 per hour minimum wage (in 2014 dollars corresponding to \$10.17 in the 1994 dollars used in the analysis). In the data, the mean unemployment rate at interview dates is 14% while in the model it is 9.2%. The introduction of the minimum wage increases this to an overall rate of 9.8%. As in FGL, the minimum wage has the largest effect on employment for individuals who have just entered the labor market. Indeed, the unemployment rate among individuals with 0-2 years in the labor force increases by 8.63% relative to baseline, while it increases by a mere 1.41% for those with 3-5 years and decreases by 3.19% relative to baseline for those who have been in the labor force for 6-8 years. Conditional on employment, average wages increases by 0.83% for new entrants, and falls by 0.11% and 0.05% for the other two groups respectively. Figure 3 shows the effect of the minimum wage policy on the acceptable match values. Consistent with FGL (2017), we see that the effect is concentrated at the lower end of the worker ability scale.

Let

$$h(\Gamma, \alpha) = \mathbb{1}[\text{emp}(\Gamma, \alpha) = 0 | \text{interview date}] \quad (20)$$

be an indicator for whether an individual in the model is unemployed at an interview date, which of course depends on the structural parameters Γ and the calibrated bargaining parameter α . The

moment of interest is then the sample average of (20). Jørgensen (2020) defines the sensitivity of the moment to the calibrated parameter as

$$H = A + BZ \quad (21)$$

where $A = \mathbb{E} \left[\frac{dh(\Gamma, \alpha)}{d\alpha} \right]$ is 1×1 , and $B = \mathbb{E} \left[\frac{\partial h(\Gamma, \alpha)}{\partial \Gamma'} \Big|_{\Gamma=\Gamma_0} \right]$ is the $1 \times p$ Jacobian, and $Z = \frac{\partial \Gamma}{\partial \alpha}$ is a $p \times 1$ vector of the sensitivity of the estimated structural parameters to the calibrated parameter. To gain intuition for (21), consider the scalar case and notice that H is just an application of the chain rule once one recalls that $h(\Gamma, \alpha) = h(\Gamma(\alpha), \alpha)$. The direct effect of α on $h(\cdot)$ is A , while the indirect effect is BZ . It is straightforward to show that if one assumes $\alpha = \alpha_0$ or, more generally, is a consistent estimator of α_0 , then Z takes the form

$$Z = M_1 G, \quad (22)$$

where as before $M_1 = -(D'WD)^{-1}D'W$, while $G = \mathbb{E} \left[\frac{\partial g(\Gamma|\alpha)}{\partial \alpha} \Big|_{\Gamma=\Gamma_0} \right]$ is the $K \times 1$ Jacobian of the moment function w.r.t the calibrated parameter⁴. The proof is given in Appendix C. Table 8 shows the vector Z , but instead transformed into an elasticity. The parameter most responsive to a one percent increase in $\alpha = 0.5$ is $\hat{\eta}$, which decreases by 56% and governs exogenous separations between the worker and firm. The signs in Table 8 are mostly consistent with those reported in the appendix of FGL (2017), and there are a few similarities such as investment technology parameters being particularly responsive to changes in α . However, one should keep in mind that the sensitivity measure (22) holds within a neighbourhood of $\alpha = 0.5$ and we know from FGL (2017), who consider $\alpha = 0.2$ and 0.8 , that the effects are nonlinear. Nevertheless, to the extent that the reader is interested in local changes around the calibrated parameter not reported in the paper, (22) serves as a useful approximation.

We can make (21) an elasticity by multiplying by $\alpha/h(\cdot)$. At the estimated values, $\hat{\Gamma}$, and $\alpha = 0.5$, the elasticity of the unemployment rate at interview dates when the minimum wage is binding is 8.51%. This implies that a 1 percent increase in the bargaining power α implies the overall unemployment rate at the interview dates will be 10.64%. Repeating this exercise for average wages for those who have been in the labor force for 0-2 years, we get an elasticity of -1.43% . This implies that at $\hat{\Gamma}$ and $\alpha = 0.5$, a one percent increase in bargaining reduces the mean wage of new entrants by -0.603% relative to baseline instead of increasing it by 0.83% as previously found. The sensitivity measure therefore shows how the qualitative result can possibly change with variation in the calibrated parameter. It is worth noting, however, that the quantitative results are small for possibly two reasons. First, a minimum wage of \$10.17, as the authors make clear and as illustrated in Figure 3, is small and only really affects the lower end of the ability distribution. Second, the object (20) is an indicator function that we take

⁴One advantage of this approach to calculating Z (and H) is that it is numerically feasible. Indeed, M_1 is constructed from the earlier standard error stage, while G involves solving and simulating the model twice if using forward-backward finite difference. A brute-force approach would involve *re-estimating* the model to compute the change in $\hat{\Gamma}$.

Sensitivity to bargaining parameter		
	Percent	Value
b	1.083	10.004
λ_u	-32.95	0.3074
λ_e	4.768	0.3908
η	-56.47	0.0003
$\delta_{a,0}$	6.223	0.0457
$\delta_{\theta,0}$	-7.334	0.0764
$\delta_{a,1}$	-30.75	-0.1175
$\delta_{\theta,1}$	15.04	0.7848
$\delta_{a,2}$	-15.84	0.2674
$\delta_{\theta,2}$	-11.79	0.5425
$\tilde{\phi}_a^-$	35.24	0.0094
$\tilde{\phi}_\theta^-$	7.509	0.1203
$\mu_{a,1}$	12.79	1.0062
$\mu_{a,2}$	0.599	1.1495
$\mu_{a,3}$	2.795	1.4112
σ_a	-2.080	0.2532
μ_θ	2.451	1.5381
σ_θ	-12.94	0.2254
β_0	-4.599	-2.6901

Table 8: Sensitivity of each parameter to $\alpha = 0.5$

the (numerical) derivative of when computing our sensitivity measure. In the absence of smoothing, we might think of our results as lower bounds.

7 Conclusion

In this paper, I revisited the identification problem in Flinn et al. (2017). More precisely, I asked the question of whether the authors' claim of wage moments within and across job spells being essential for estimating the distribution of the worker-firm match productivity and the worker ability distributions could be verified. To help formalise some of the claims, I used sensitivity measures that have been recently proposed in the econometrics literature seeking to become more aware and minimise sensitivity of parameter estimates in dynamic structural models. Using these measures, I was able to verify the claims of the authors, and go further: not only are moments characterising the wage distribution of job-movers important for recovering reliable estimates of the distributions, but moments characterising employment transitions seem to be far less important while those related to worker training seem to be negligible. This latter result counters the authors' intuition for the importance of training data for the distribution parameters.

One advantage to the proposed sensitivity measures is that they do not go beyond the researcher's

information set, but instead ask what more could have been learned about primitive parameters. This is important because in the absence of matched employer-employee data, the identification problem becomes more challenging and one has to carefully choose which moments to include. However, a drawback of these measures is that they are local. Indeed, each measure can hold as a good approximation close to the true Γ_0 , and this is a consequence of the fact that they all depend on the asymptotic normality of the SMM estimator, $\hat{\Gamma}$, which is derived from a local expansion around Γ_0 . Ideally, one wishes to better understand their SMM objective function over the entire parameter space given their chosen moments. Identification of the parameters is a necessary condition for consistency of the asymptotic normal estimator. This then implies that we should treat the measures in Section 5 as helping the researcher choose the minimum sufficient number of moments that identify the model, as opposed to which moments to include in the first place.

References

- Altonji, J. G. and L. M. Segal (1996). Small-sample bias in gmm estimation of covariance structures. *Journal of Business & Economic Statistics* 14(3), 353–366.
- Andrews, I., M. Gentzkow, and J. M. Shapiro (2017). Measuring the sensitivity of parameter estimates to estimation moments. *The Quarterly Journal of Economics* 132(4), 1553–1592.
- Bagger, J. and R. Lentz (2019). An empirical model of wage dispersion with sorting. *The Review of Economic Studies* 86(1), 153–190.
- Christensen, T. and B. Connault (2019). Counterfactual sensitivity and robustness. *arXiv preprint arXiv:1904.00989*.
- Flinn, C., A. Gemici, and S. Laufer (2017). Search, matching and training. *Review of Economic Dynamics* 25, 260–297.
- Honoré, B. E., T. H. Jørgensen, and Á. de Paula (2020). The informativeness of estimation moments. *Journal of Applied Econometrics* 35(7), 797–813.
- Jørgensen, T. (2020). Sensitivity to calibrated parameters.
- Karahan, F., S. Ozkan, and J. Song (2019). Anatomy of lifetime earnings inequality: Heterogeneity in job ladder risk vs. human capital. *FRB of New York Staff Report* (908), 1–45.

Appendices

A Moments used in Estimation

In addition to those moments shown in Tables 5 and 6, some additional moments used in estimation are:

Employment transition

1. EU probability
2. unemployment rate at interview
3. EE transition probability
4. fraction EUE, 1 change, ed=1,2,3
5. fraction EUE, 2 changes, ed=1,2,3
6. fraction EUE, > 2 changes, ed=1,2,3
7. fraction EE, 1 change, ed=1,2,3
8. fraction EE, 2 changes, ed=1,2,3
9. fraction EE, > 2 changes, ed=1,2,3
10. fraction same employer, ed=1,2,3

Wages

1. mean wage after 0-2 years, ed=1,2,3
2. mean wage after 3-5 years, ed=1,2,3
3. mean wage after 6-8 years, ed=1,2,3

same employer

4. mean w_2 , 2 wage obs., ed=1,2,3
5. mean $w_1 \times w_2$, 2 wage obs., ed=1,2,3
6. mean $w_2 - w_1$, 2 wage obs., ed=1,2,3

7. mean w_3 , 3 wage obs., ed=1,2,3
8. mean w_3^2 , 3 wage obs., ed=1,2,3
9. mean $w_1 \times w_3$, 3 wage obs., ed=1,2,3
10. mean $w_3 - w_1$, 3 wage obs., ed=1,2,3
11. mean w_4 , > 3 wage obs., ed=1,2,3
12. mean $w_4 - w_1$, > 3 wage obs., ed=1,2,3

job-to-job transition

13. mean w_{t-1} , 1 change, EUE, ed=1,2,3
14. mean w_t , 1 change, EUE, ed=1,2,3
15. mean w_{t-1} , > 1 change, EUE, ed=1,2,3
16. mean w_t , > 1 change, EUE, ed=1,2,3
17. mean w_{t-1} , 1 change, EE, ed=1,2,3
18. mean w_t , 1 change, EE, ed=1,2,3
19. mean w_{t-1} , > 1 change, EE, ed=1,2,3
20. mean w_t , > 1 change, EE, ed=1,2,3

As well as second-moments of job-to-job transitions.

Training

1. fraction receiving training in current job
2. fraction receiving training in current job by education ed=1,2,3

B Computation of standard errors

The SMM variance-covariance matrix used throughout the analysis is computed using a diagonal weighting matrix with elements being the inverse of the variance of the empirical moments. Moreover, the variance-covariance matrix of the empirical moments is also diagonal. The reasons are twofold: (i) I follow the method used in FGL (2017) with respect to S and (ii) Altonji and Segal (1996)

highlight the small-sample bias present when using the full optimal weighting matrix. Of course, the computation of S has the implicit assumption that the empirical moments are uncorrelated, which is unlikely to be true. Therefore, the variance matrix takes the form

$$\Sigma = \left(\hat{D}' \hat{W} \hat{D} \right)^{-1}. \quad (23)$$

To compute the Jacobian matrix, I do the following:

$$\hat{D}_j = \frac{g(\hat{\Gamma} + \epsilon_j) - g(\hat{\Gamma} - \epsilon_j)}{0.05\Gamma_j} \quad (24)$$

where $\hat{D}_j$ is the j -th column of the $K \times p$ Jacobian matrix $\hat{D}$, and ϵ_j is a vector of zeros with one positive element at the j -th position equal to 2.5% of the parameter value $\hat{\Gamma}_j$. The $K \times 1$ vector $g(\Gamma)$ is the simulated moment vector. Thus, with 19 parameters estimated, and with forward-backward finite-differencing used, obtaining $\hat{D}$ and thus, obtaining Σ , requires solving the model 38 times in order to evaluate the simulated moments. In total, this procedure takes 1.5 hours.

C Proofs

In this section, I provide the proof to the sensitivity of the estimated structural parameters, $\hat{\Gamma}$, to the calibrated bargaining parameter α , which is given in (22). The proof resembles that found in Jørgensen (2020) but I give a little more detail. Consider the minimum-distance objective function

$$Q_n(\Gamma|\alpha) = \frac{1}{2} g_n(\Gamma|\alpha)' \hat{W} g_n(\Gamma|\alpha). \quad (25)$$

Suppose $Q(\Gamma|\hat{\alpha})$ and $Q_n(\Gamma|\hat{\alpha})$ and the parameter space satisfy the requirements for $\hat{\Gamma}$ to be a consistent estimator of Γ_0 , when $p \lim_{n \rightarrow \infty} \hat{\alpha} = \alpha_0$. Then the approximate first-order condition holds:

$$\begin{aligned} o_p(n^{-\frac{1}{2}}) &= \frac{\partial Q_n(\hat{\Gamma}|\hat{\alpha})}{\partial \Gamma} \\ &= -D_n(\hat{\Gamma}|\hat{\alpha})' \hat{W} g_n(\hat{\Gamma}|\hat{\alpha}). \end{aligned} \quad (26)$$

If we take a mean-value expansion of the term $g_n(\hat{\Gamma}|\hat{\alpha})$ around Γ_0 then

$$o_p(n^{-\frac{1}{2}}) = -D_n(\hat{\Gamma}|\hat{\alpha})' \hat{W} (g_n(\Gamma_0|\hat{\alpha}) + D_n(\tilde{\Gamma}|\hat{\alpha})(\hat{\Gamma} - \Gamma_0)), \quad (27)$$

where $D_n(\hat{\Gamma}|\hat{\alpha}), D_n(\tilde{\Gamma}|\hat{\alpha}) \rightarrow^p D(\Gamma_0|\alpha)$. We can rearrange (27) and represent it as (suppressing α dependence)

$$\begin{aligned} \hat{\Gamma} &= \Gamma_0 - \left(D_n(\hat{\Gamma})' \hat{W} D_n(\hat{\Gamma}) \right)^{-1} D_n' \hat{W} g_n(\Gamma_0) + o_p(n^{-\frac{1}{2}}) \\ &= \Gamma_0 + \hat{M}_1 g_n(\Gamma_0) + o_p(n^{-\frac{1}{2}}) \end{aligned} \quad (28)$$

which converges in probability to

$$\hat{\Gamma}(\alpha) = \Gamma_0 + M_1(\alpha)g(\Gamma_0|\alpha). \quad (29)$$

where $p \lim_{n \rightarrow \infty} = \alpha$. The Jacobian of the estimated structural parameters w.r.t α is then

$$\frac{d\hat{\Gamma}}{d\alpha} = \frac{dM_1(\alpha)}{d\alpha}g(\Gamma_0|\alpha) + M_1(\alpha)\frac{dg(\Gamma_0|\alpha)}{d\alpha}. \quad (30)$$

If $\alpha = \alpha_0$ then $g(\Gamma_0|\alpha_0) = 0$ which implies

$$Z = M_1 G. \quad (31)$$

D Figures

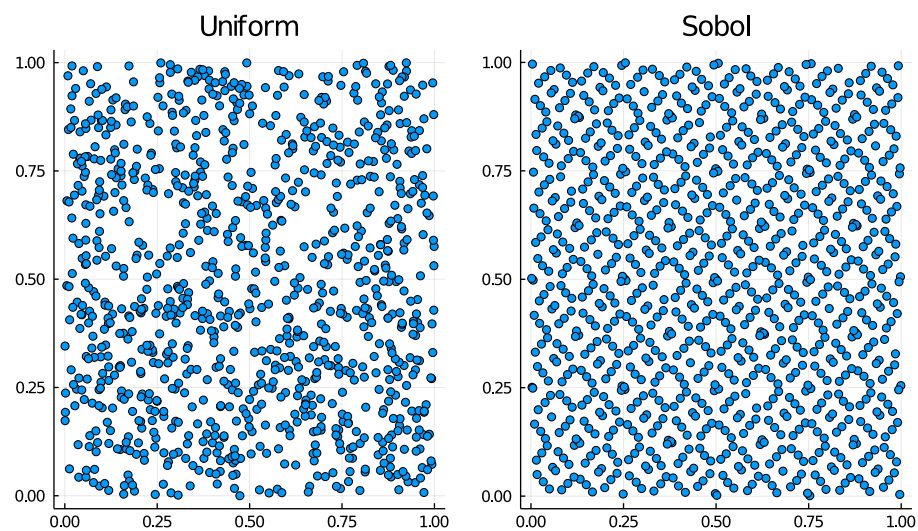


Figure 1: Uniform vs. Sobol sequences

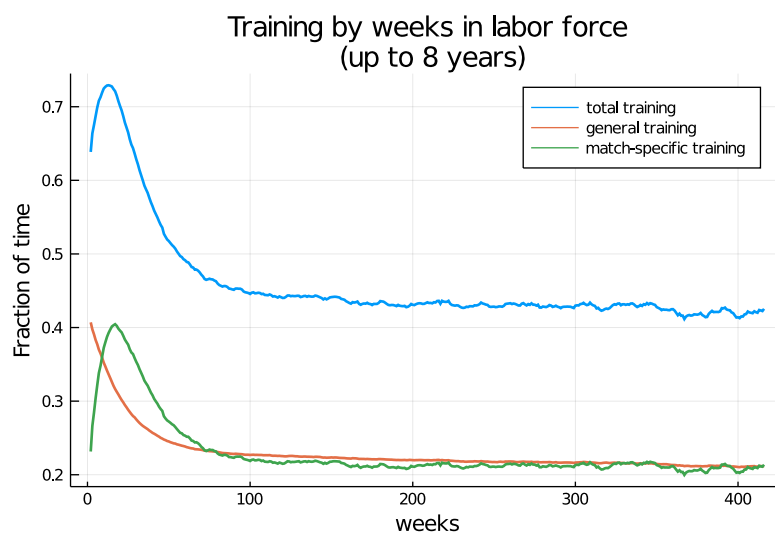


Figure 2: Simulated Training

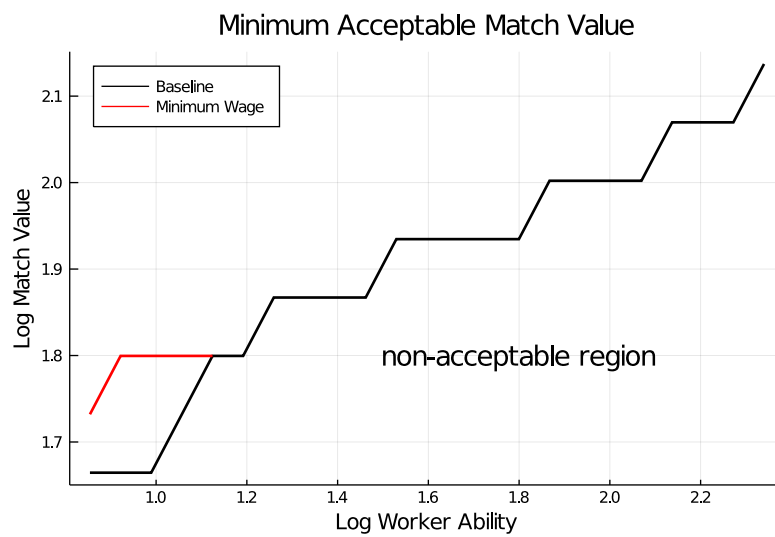


Figure 3: Minimum acceptable match value