

Introductory Labor Economics: A Modern Approach with Applications

Jeremiah Richey

Contents

1	Introduction	1
	Some General Economic Issues to Keep in Mind	2
1.1	Positive vs. Normative	2
1.2	Optimality	2
1.3	Scarcity and Rationality	3
1.4	The Use of Models	3
	Some Basics About Labor Markets	4
2.1	Long Run Trends	4
2.2	Payments	6
2	Crash Course in Regression Analysis and Optimization	7
	Optimization	7
1.1	Basic Idea	7
1.2	Constrained Optimization	8
	Regression	9

2.1	Basics	9
2.2	Omitted Variable Bias	12
3	Labor Supply 1: Labor-Leisure Model	16
	The Budget Constraint	16
	Utility and Optimization	18
2.1	Corner Solution and Reservation Wage	19
	Changes in the Wage Rate	20
	Non Labor Income	21
	Program Evaluations	22
5.1	Income Replacement Programs	23
5.2	Traditional Welfare	24
5.3	Earned Income Tax Credit	25
	Extensive vs Intensive Margins	26
4	Labor Supply 2: Household Production and Life-Cycle Models	28
	Household Production	28
1.1	General Model	28
1.2	Simplified Model	29
	Household Specialization	31
2.1	Simple Model	32

2.2	Recessions and Labor Supply	33
	Life Cycle Considerations	34
3.1	Life Profiles	35
3.2	Retirement	35
5	Education and Human Capital	41
	Some Facts	41
	Choosing Education and Selection	42
2.1	A Simple Model	42
2.2	Selection	43
2.3	Why Higher Wages	45
	Human Capital	45
3.1	General Training/Education	45
3.2	Specific Training	46
	Education as a Signal	47
4.1	Extensions	49
	Other Issues	51
6	Labor Demand	53
	Basic Profit Maximization	53
	Scale vs. Substitution Effect	56

Labor Demand Function	57
Government Regulations and Labor Demand	59
4.1 Minimum Wage	59
4.2 Taxes and Mandated Benefits	63
4.3 Employment Protection Laws	64
7 Compensated Wage Differentials	67
Basics	67
Application to Regulation	69
Application to Pricing	71
8 Discrimination	72
Some Data for Korea	72
Theory	74
2.1 Taste Based Discrimination	75
2.2 Statistical Discrimination	77
Decompositions	79
9 Search Theory and Unemployment	80
Basics of Search	80
Types of Unemployment	82
Efficiency Wages and The Wage Curve	83

10 Unions	86
Monopoly-Union Model	87
Efficient Contracts	87
Issues and Evidence	90
3.1 The Demand Curve	90
3.2 Effect of Unions	91

Chapter 1

Introduction

This is a course in ‘Modern Labor Economics’ as opposed to ‘Institutional Labor Economics’ (see “Labor Economics” by Boyer and Smith (1998) in *Industrial relations at the dawn of the new millennium* for a discussion of the history and evolution of the two). Thus we will view labor economics is basically just an applied microeconomics field (though there are some areas of labor economics that are ‘macro-labor’ in a sense, but our focus is be in the micro setting). When you think about it, labor, like any other commodity, is something that is owned by one and sold to another. But of course there are some key differences between looking at a market for labor verses a market for socks. For one labor is rented and not sold. Also, the ‘sellers/suppliers’ of labor care about how their labor is used. That is, you do not rent out our labor, have it returned at the end of the day with no concern for for what happened in the mean time. Another issue is that the ‘quality’ of labor supplied has some interesting twists. Maybe it is not observed precisely, maybe the working conditions themselves affect the quality. Also, the quality of labor can likely be improved through training or eduction. And lastly, pay is not the only issue for supplying labor. There are usually other ‘non-pecuniary’ matters that come into the decision. But as long as we keep some of these caveats in mind we will proceed just as an applied micro setting. But labor economics is not a field just about the supply and demand of labor. It also encompasses related issues such as family, marriage and fertility, education, training, and discrimination. This text will cover these related topics as well.

Some General Economic Issues to Keep in Mind

1.1 Positive vs. Normative

When discussing economic problems, especially ones that can seem personal such as supplying ones own labor, we want to make sure we distinguish between *normative* and *positive* arguments or statements. Positive statements are in the form of ‘what is’. It is a statement of fact. For example we may have a model and we may say that if some parameter is changed then the level of labor supply will go up. This is a positive statement. A normative statement, on the other hand, is in the form of ‘how it should be’. It is a statement regarding how one thinks things should be. For example ‘the government should raise the retirement level because...’. All of our models will be based on positive economics, though once we decide what they mean for the real world we may many times want to delve into the normative world. This is, of course, OK. But, make sure you are aware when you make that transition.

1.2 Optimality

Of course regularly in economic discussions there will be normative statements used, and it is normally based on an understood value system, and that value system is optimality or efficiency. The most basic form of optimality is *pareto optimality*. This is a result which comes about once all three of the following transactions have occurred:

1. Transactions that benefit both parties
2. Transactions that benefit some, but nobody loses
3. Some benefit, some lose, but the winners compensate the losers

This last one, #3, gets then into the world of cost-benefit analysis - because we need some how to measure how *much when* discussing wins and loses. Of course many times people vote to do non-pareto efficient policies. Are they crazy? No, it is just that they are basing their decision on a different value system. But it does seem that at least efficiency should be a good starting point for any policy discussion. And that is in the realm of economics.

You may many times hear the issue of *equity* vs. *efficiency*. And it is most likely that people are willing to give up some efficiency for more equity, but that is a decision to be made by voters. The role of the economist is to highlight exactly what that trade-off would look like.

1.3 Scarcity and Rationality

We will be working with two basic assumptions concerning the world and the people living in it. The first is that there is scarcity. That is there is not enough for everyone to have everything. Somehow goods and services must be rationed. Just how that happens is what we will be looking at, and we will be using markets as are starting point. The second thing we will be assuming is that people are rational. We will assume that people maximize their utility (see a standard microeconomic text to see a in depth discussion of what exactly we mean by utility eg. “Intermediate Microeconomics: A Modern Approach” by Varian). This does not mean they don’t care about other people. If you see someone giving to charity then an economist will simply say that other people’s consumption must be part of that persons utility function. It also does not mean the person knows everything. Many models will be based on uncertainty. Similarly we will assume firms maximize profits. Though there are some firms that explicitly do not maximize profits, eg. non-profits, they will not be our concern here.

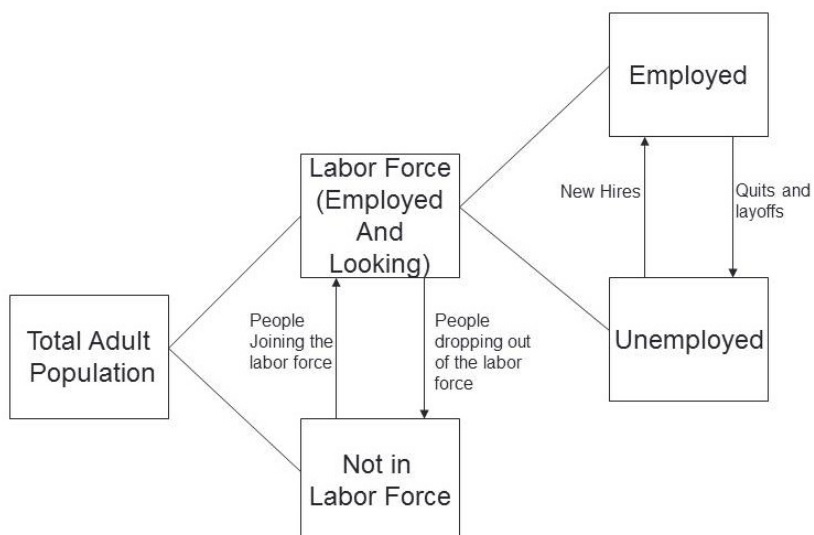
1.4 The Use of Models

We will be looking at many models in this course. None of these models exactly match reality. But we already have a model that exactly matches reality - the real world! What we want is a simplified version of reality that will allow us to understand certain aspects of it. In order to do this we will have to sacrifice some realism. But as long as this lack of realism does not interfere with the model and its role in helping us understand the issue at hand then the simplification is a good thing. Please do not disregard a model because it is not ‘real’. But *do* look closely at the assumptions and try and understand where they come into play in the results of the model. This is what makes a good economist.

Some Basics About Labor Markets

The demand side of the labor market are the firms in the economy (employers). The supply side of the labor market is the population (employees). The *labor force* is the part of the population that is either working or actively seeking work - the labor force participation rate is this number divided by the ‘working age’ population (15-65). So if someone is not even looking for work they are not considered in the labor force. The *unemployed* are those in the labor force but not employed - so these are those actively seeking but not working. And there are constant dynamics of people moving in and out of the labor force and also moving in and out of employment.

Figure 1.1: Breakdown of Population



Currently Korea has a labor force participation rate of about 61.5% and an employment rate at about 59.5% (about 71% for men, 49% for women, 40% for youths (15-29), and 65% for older workers (55-64)).

2.1 Long Run Trends

We also tend to see long run trends in most economies regarding labor force. We generally see female participation rates increase as economies develop while the male participation rate tends to stay rather constant (though it tends to dip during harsh recessions). We can see these trends in Korean economy:

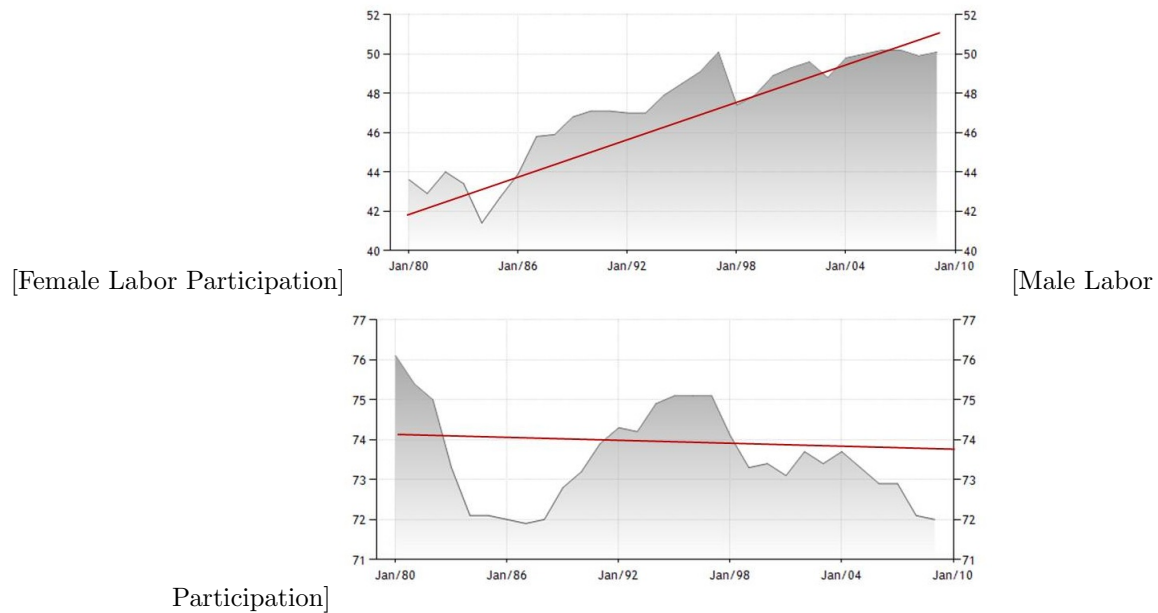


Figure 1.2: Labor Participation Rates for both Sexes.

We also tend to see similar trends across countries as they develop regarding the sectors. Traditionally we would divide an economy up into three sectors: (1) agricultural, (2) industrial, (3) service. And what we tend to see is the labor force is initially concentrated in the agricultural sector. Then as the economy grows the industrial sector becomes much more important. And as an economy becomes 'industrialized' (or rich) the service sector becomes the dominant area. Today in Korea the agricultural sector accounts for about 6% of labor, the industrial about 24% and the service sector about 70%. Moreover when can see the trends over time.

Why do we see these trends? Agricultural output experience rapid growth in labor productivity, but has a relatively low income elasticity of demand. So less people are needed to grow the food, but the increase in wealth does little to boost demand for the product of the sector. The industrial sector also experienced growth in labor productivity, but it does have a moderate to high income elasticity, so there is a balancing out at some point. Services on the other hand have experience very small gains in labor productivity, and there is a high income elasticity of demand - thus the labor required increases. And of course depending on the size of the economy exports can play a large role - especially for the industrial sector for Korea.

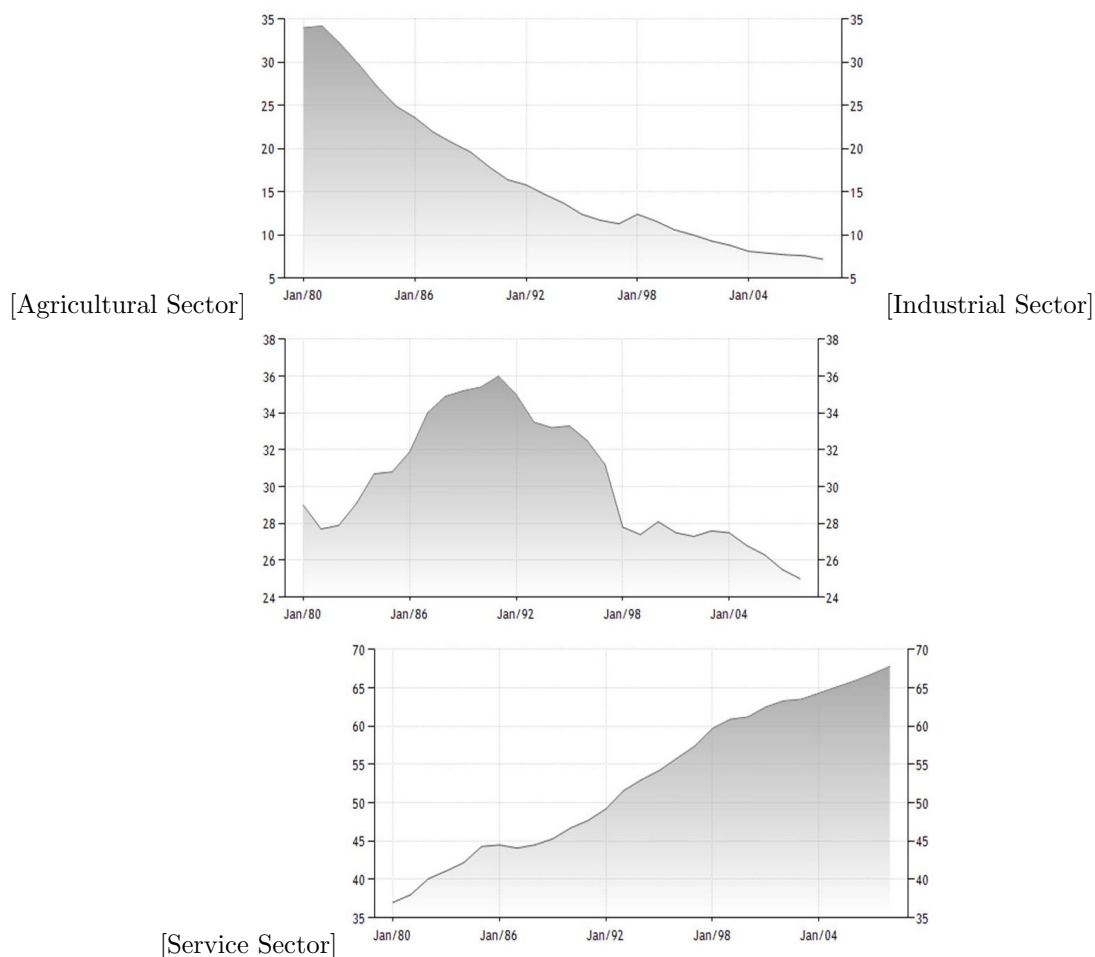


Figure 1.3: % of Labor Force per Sector

2.2 Payments

We should make some distinctions regarding the different measurements regarding how labor is paid. We could be talking about hourly or monthly wages or yearly earnings. This makes a difference. Maybe some people have much higher yearly earnings because they work more hours. This is quite different than having higher hourly wages. Also there is a distinction between real and nominal wages. Nominal wages are the actual amount deposited into your bank account. Real wages are what those wages are actually ‘worth’ - ie. what can they buy. To move from the former to the latter we need to adjust for some level of prices. This is generally done using some form of a Consumer Price Index (CPI). This is based on the cost of a standard ‘basket of goods’. Tracking how this cost changes over time gives us an index of how average prices (or some measure of them) changes over time. We need to keep this in mind when deciding if wages actually go up or down.

Chapter 2

Crash Course in Regression Analysis and Optimization

Optimization

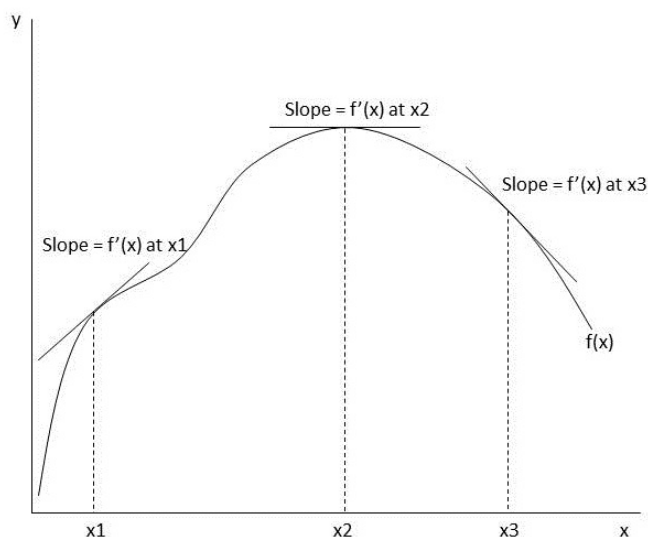
1.1 Basic Idea

We will be looking at some simple models. To understand some of these you need a basic understanding of optimization. Optimization, in general, means picking a value of some variable that gives the highest (or possibly lowest) value of a function based on that variable. This is not an optimization course and so I will not really go into details of sufficient conditions or existence/uniqueness. We will assume these hold for all of our models. We just want to make sure you get the ‘down and dirty basics’. The main thing you need to understand is what a function is and what a derivative is.

Given a function $y = f(\cdot)$ of a variable x we write the derivative as either $\frac{df}{dx}$, $\frac{dy}{dx}$, or simply $f'(x)$. Both of these answer the question ‘how much does y change when I changes x a very small amount at a particular value of x ?’ That’s it. What do we mean by ‘a very small amount’? Technically we mean as the change in x approaches zero, but lets not get sidetracked with details, we just want the big idea. So basically, in a

simple function like $y = f(x)$ the derivative is just the slope at some value of x . Note however that $f'(x)$ is a function and not a number. It gives a different value for different values of x . Of course we can evaluate $f(x)$ at particular points. The standard notation for evaluating a function at a given point is: $f'(x) |_{x=\hat{x}}$. Given the follow graph we see that the slope depends on where we are on the function.

Figure 2.1: A Function and What a Derivative Is.



If we were trying to maximize our function $f(x)$ we would want to locate the top of the hump. And this is characterized by the derivative being equal to zero. Now if there are multiple humps we may have problems, but all of our applications will be ‘well behaved’. That’s pretty much it, at least at the simplest level. If you want to find the maximum of a function, just take the derivative and figure out where it equals zero. This basic procedure will be used repeatedly in our models.

1.2 Constrained Optimization

Sometimes we will want to maximize something in the face of some ‘constraint’. For example we will be looking at maximizing utility in the face of an income constraint. In all of our application we will generally be able to transform a constrained optimization into an unconstrained one by simply rewriting our problem. So this is not a big concern for us and we will deal with it when we see it in applications.

Regression

In this text we will not only be looking at the theory of labor economics, but also applied papers that have attempted to actually give quantitative results to the qualitative ones coming out of the theory. To properly understand these papers and findings you need to be at least somewhat familiar with regression analysis at its most basic level. This section will do that for you. Please do not think this in anyway is a statistics or econometrics chapter in any rigorous sense - it is just what it says - a “crash course”. For a very good introduction to econometrics (statistics for economic problems) see “Introductory Econometrics: A Modern Approach” by Wooldridge.

2.1 Basics

Say we want to answer the question “Do people who attain more education receive higher wages?” How might we go about doing this¹? Essentially we are interested in the relationship between two variables: (1) *years of schooling* and (2) *wages*. Assume we have some data, maybe from some survey, regarding people’s education levels and their wages. Lets see what that might look like. Lets look at it in a *scatter plot* (Figure 2.2).

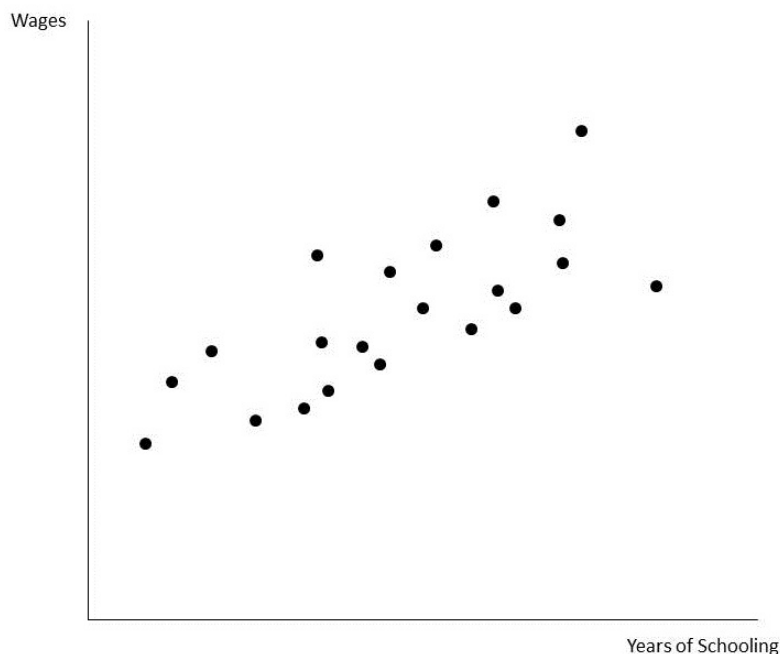
If you are like most people it seems like a good first guess regarding the relationship between these two variables would be a line ‘through the middle’ of the points. Lets call this the ‘best fit line’. Now recall that a line can be written in a general form: $y = a + b \cdot x$. So what we have here is a best guess about the relationship between wages and education that looks something like:

$$wages = \beta_0 + \beta_1 \cdot education. \quad (2.1)$$

We call *wages* our outcome or explained variable, and we call *education* our explanatory variable. The intercept of this line - β_0 - is what someone would make if they had zero years of education. The slope - β_1 - is the relationship of interest. This tells us on average how much wages go up with education levels in the data we have. Of course this relationship doesn’t hold exactly for all people (probably not any) but is rather an ‘average’ relationship. To reflect this we note that there is some ‘error term’ in this relationship written

¹Be careful and note that the question was not “Does attaining more education *cause* people to receive higher wages?” This is an important distinction that we will come to shortly.

Figure 2.2: Some data on wages and years of schooling.



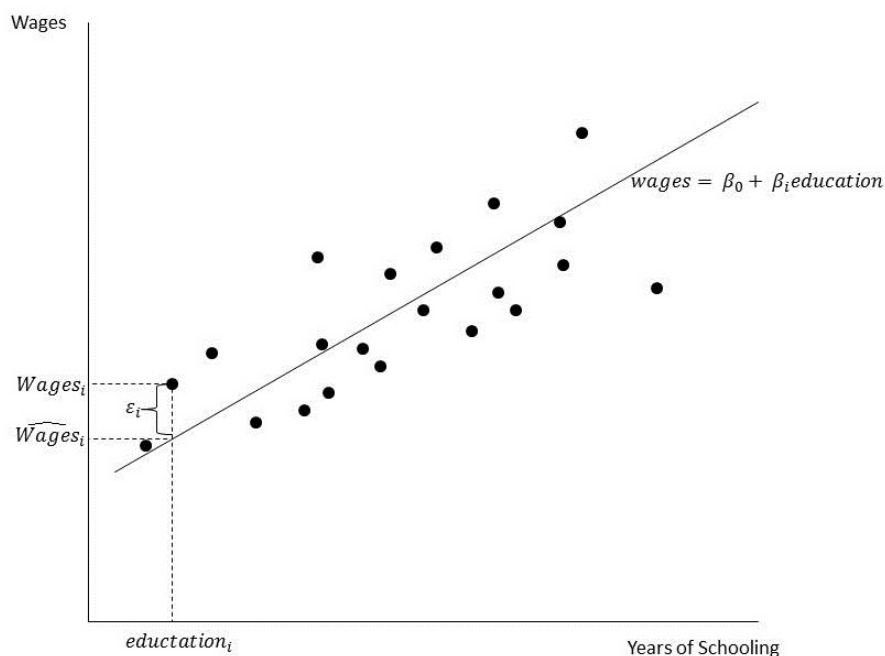
above. We do this by noting:

$$wages_i = \beta_0 + \beta_1 \cdot education_i + \epsilon_i. \quad (2.2)$$

Where the ϵ catches everything else that affects wages, and the subscript reflects the fact that there are many people (observations) and each has a different wage, education level, and error term. What we assume about that error term will be important. Furthermore, our ‘predicted value’ is what our equation gives us for *wages* for a specific value of *education*. That is $\hat{wages}_i = \beta_0 + \beta_1 \cdot education_i$. Now just to make sure we are on the same page, lets add these things to our picture (Figure 2.3).

Now when we actually ‘run a regression’ (the algebraic process of getting numbers for our line) and get actual numbers for β_0 and β_1 they are estimates, meaning they are not exactly the *true* relationship, but rather an estimate of it which is a result of a random process (so we are trying to get an *average relationship*, but end up with an *estimate of the average relationship*). And along with estimates of β_1 we also calculate measures of the uncertainty behind that estimate which is called the *standard error*. And when we report our results the standard error is normally in parenthesis. And what matters is the relationship between our estimate and its standard error. Now normally what we are doing is testing whether we actually observed a real relationship, or we just got our estimate by chance. So we are testing whether or not the real β_1 is just zero - ie. no relationship. In that case a good rule of thumb is of 2 - your estimate should be twice

Figure 2.3: Some data on wages and years of schooling and best fit line.



as large as your standard error. This is about right at what is called the 5% significance level (we will not get into the statistics behind this here). Moreover, many times authors report these ratios - they are called *t-statistics*. So you might see a set of results such as in Table 2.1.

Table 2.1: Results: Outcome is Wages

Dep. Variable	Estimate	T-Stat
Education	139.50** (66.10)	2.11
Intercept	197.6*** (60.80)	3.28

Robust standard errors are in parenthesis.

So here we see our relationship between education and wages is significant (some authors add stars as here where one indicates ‘significant at the 10% level’, two indicates at the 5% and three at the 1% - so more is better). And when we interpret these results they mean people with an extra year of schooling have (on average) wages that are higher by \$139.50 a month (if wages were measured in monthly earnings).

2.2 Omitted Variable Bias

OK, so that's all fine, but we probably don't really care about seeing if people with higher education levels make more money, we probably really care about whether higher education levels *cause* people to make more money - and this is a trickier question. What matters for this question - whether we can answer it - is the relationship between our explanatory variable and our error term. The key thing is we don't want them related (statistically this is incorrect, but OK for our needs). Perhaps the best way to think about a possible problem we want to avoid is by thinking about the error term as 'all other variables that affect wages'. And if one of those is related to education we have a problem. Specifically we have an *omitted variable problem*. This is usually the big hurdle in estimating interesting effects in labor economics.

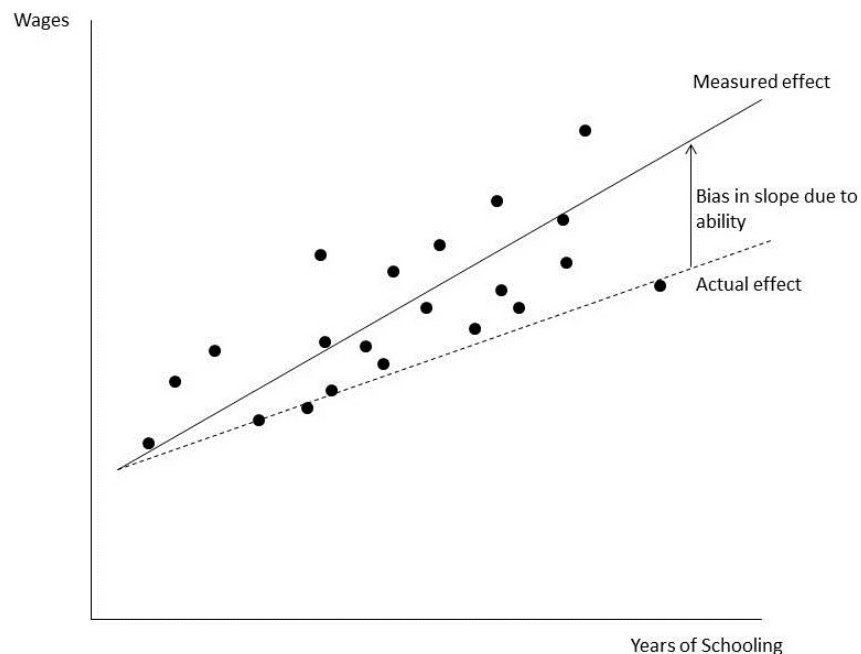
So is there a problem here in our example? Well can you think of a variable that is related to education levels and also directly related to earnings? How about natural ability, or IQ? Maybe people with high education levels would earn more even without a lot of schooling because they are naturally more talented (we will actually see such a model later). If this is the case then our estimate β_1 , though correct as far as the relationship between education and wages, is a *biased* estimate of the *causal effect* of education on wages. Biased just, for our concerns, means wrong. We can see graphically what is going on below in Figure 2.4. As education increases there is a causal effect (the dotted line) but as we move to the right these people are also on average smarter so they get higher wages also because of that, but since we don't measure ability we lump them together in our measured effect and wrongly overstate the causal effect of education.

Now if we have data on ability then we can just include it in our equation and run a 'multiple regression analysis'. It would look something like:

$$wages_i = \beta_0 + \beta_1 \cdot education_i + \beta_2 \cdot ability_i + \epsilon_i. \quad (2.3)$$

But in some instances we cannot get data on an important variable and so have a problem. In the papers we will be looking at this is a very common problem and researchers have developed many ways to overcome this problem and estimate causal effects of interest.

Figure 2.4: Omitted Variable Bias.



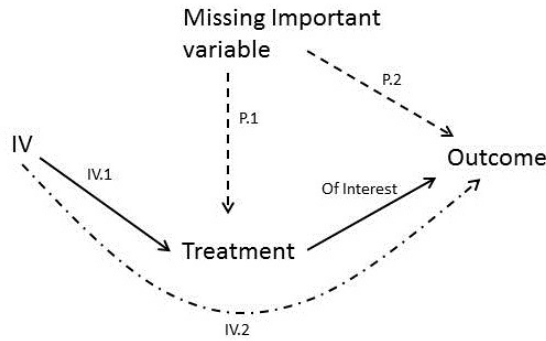
Possible ‘Fixes’ for OVB

There are several ways economists go about trying to ‘fix’ the problem of omitted variable bias. Here I will briefly explain the idea of two of them that are common in the literature. I will not explain mathematically how they work, but intuitively. You will need to get the basic idea of these to understand the papers we will be looking at.

Fix 1: Instrumental Variables (IVs): The idea of IVs are pretty simple. The problem of OVB is that there is something potentially causing both our causal variable and our outcome of interest, so just looking at the relationship of the two will confound the effect of this missing variable. An IV is a variable that is related to the causal variable of interest, but *only related to the outcome of interest through the effect of the causal variable*. If we can find one of these we can figure out the causal effect we want to estimate. This is because any relationship between this IV and the outcome must be through the effect of our causal effect of interest. See Figure 2.5. So if we divide the effect IV.2 by IV.1 this should give us our effect of interest, because IV.2 only works through IV.1 $\rightarrow$ ‘effect of interest’. Of course finding a good IV is a much harder thing then it seems. Some papers will argue they have a valid IV that helps them estimate their effect of interest.

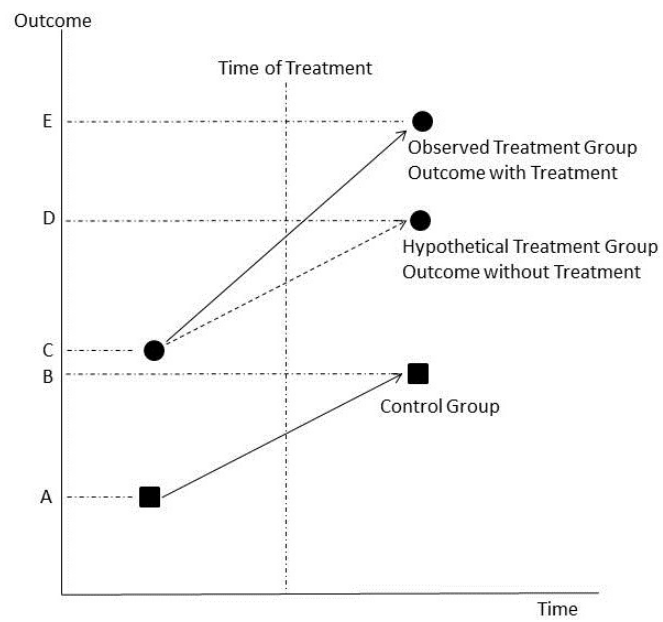
Fix 2: Difference-in-Difference (DnD): The idea of difference and difference is basically based on some ‘control

Figure 2.5: Understanding IV.



group'. Say we want to measure the effect of some treatment, but we worry about OVB. So we want to know some effect of a treatment and there is a before treatment and and after treatment. If we have some population, related to the group that received the treatment, but who did not receive the treatment, then we can treat them as a control group. So we look at how each group changed before and after (difference) and compare these two changes (difference). Thus we are interested in the difference between groups in the difference between before and after - hence the term difference-in-difference. Look at the following figure. If we just measure treatment as $E - C$ we see we will be overstating the effect. The key assumption is that the treatment and control groups would have changed in the same way were there no treatment ($D - C = B - A$). But if we measure the treatment effect as $(E - C) - (B - A)$ with our assumption regarding relation between treatment and control this equals $(E - C) - (D - C) = E - D$ which is the actual treatment effect.

Figure 2.6: Understanding Diff-in-Diff.



Chapter 3

Labor Supply 1: Labor-Leisure Model

The Budget Constraint

In this model there are only two uses of time: labor or leisure (work or play). Each individual selects the combination of hours in work and leisure that maximizes his or her utility. The individual faces a given wage rate w . Thus this wage rate is also the opportunity cost¹ of leisure. They also have a fixed amount of available time T .

Thus they face a *time constraint*. If we denote H the amount of time working and L the amount of time in leisure then their time constraint is²:

$$H + L = T \tag{3.1}$$

They also face a *goods constraint*. We ignore what type of goods they actually consume and just for simplification assume they ‘consume’ income. Thus they have a goods constraint of:

$$Y = wH \tag{3.2}$$

Note this implicitly puts the price of the good to be equal to one. Of course this does not matter much

¹Recall the opportunity cost of an activity is the value of the next best thing. Since there are only two things to do in this model, more leisure means less work and work receives w .

²Technically this should be a $\leq$, but individuals will use all of their time so we disregard this nuance.

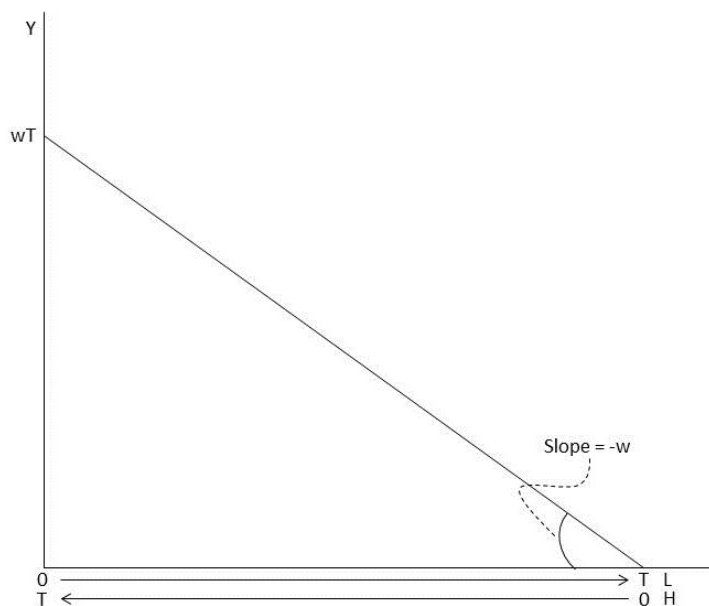
because if we allow it to be ‘ p ’ then we can just reinterpret the wage as a real wage: w/p .

Now if we put these two constraints together we get a budget constraint (solve 3.1 for H and substitute it into 3.2 and solve for wT):

$$wT = Y - wL \quad (3.3)$$

The interpretation is that they have some total possible income (wT) which they can either use to consume goods ($Y = wH$) or to consume leisure (wL) - of course to consume goods they must work to receive wages, thus the wH . We can see graphically what our budget constraint looks like:

Figure 3.1: Agent’s Budget Constraint.



Notice that the intercept on the goods axis (Y) is wT . This is because if they chose to consume only goods this is the most they could consume. And the intercept on the leisure axis (L) is T - this is how much time they have available to consume as leisure. Also the slope is $-w$. For each unit of time taken as leisure they must give up w times that in the consumption of goods. So we see here again that the wage rate is the opportunity cost of leisure. Furthermore note that we can label the x-axis as either leisure L in which case it goes from $0 \rightarrow T$, or we could label it as labor H in which case it goes from $T \rightarrow 0$.

Utility and Optimization

The agent maximizes utility based on their consumption of goods (income) and leisure.

$$utility = U(Y, L) = U(w(T - L), L) \quad (3.4)$$

Note the second equality substitutes the relation between income and leisure. The agent maximizes their utility subject to their budget constraint. And this can be easily solved by maximizing over their choice of L (note they only have one choice to make - once leisure is decided then so is consumption). Taking the derivative of the utility function with respect to leisure (using the chain rule) and setting it to zero we have:

$$\frac{\partial U}{\partial Y} \cdot (-w) + \frac{\partial U}{\partial L} = 0 \quad (3.5)$$

Now we can rearrange this to solve for the marginal rate of substitution (MRS) between consumption and leisure:

$$MU_L / MU_Y = w \quad (3.6)$$

or similarly we can write it as:

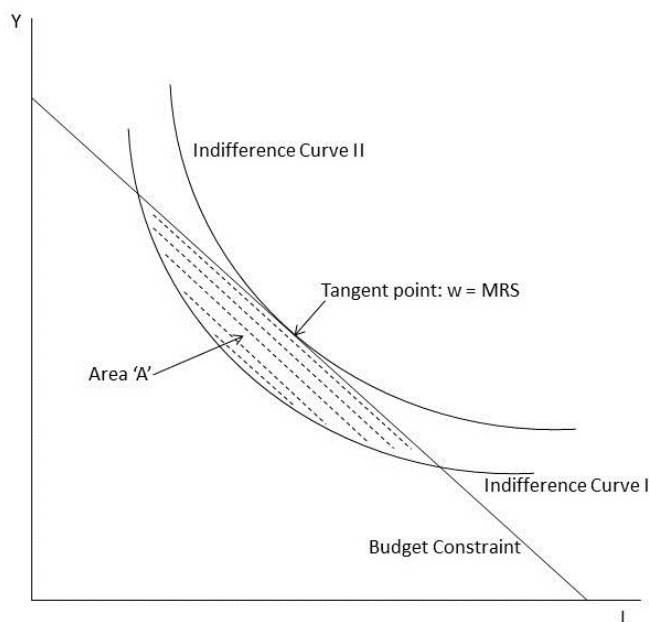
$$\frac{dY}{dL} \Big|_u = MRS = -w \quad (3.7)$$

Now to tie this together lets look at this graphically by introducing the agent's indifference curve. Recall the indifference curve is just a graphical representation of all bundles of Y and L that give the same level of utility, so it is what is called a 'level curve' of the utility function. And the slope of the indifference curve is the MRS - it is how much the agent is willing to give up of income to gain another unit of leisure while maintaining the same level of utility³ - see Figure 3.2.

Indifference curve I does not yield max utility as the agent can move to a higher indifference curve and still afford it. These 'higher utility bundles' are in Area 'A'. So we see this can continue until we reach Indifference curve II which is tangent with the budget constraint, and at this point the slope of the indifference curve and budget constraint are the same: $w = MRS$.

³This can also be seen by totally differentiating the utility function and noting that if you are on the same indifference curve (same level of utility) then this must be zero: $MU_L \Delta L + MU_Y \Delta Y = \Delta U = 0 \rightarrow MRS = \frac{\Delta Y}{\Delta L} = -\frac{MU_L}{MU_Y}$.

Figure 3.2: Utility Maximization.

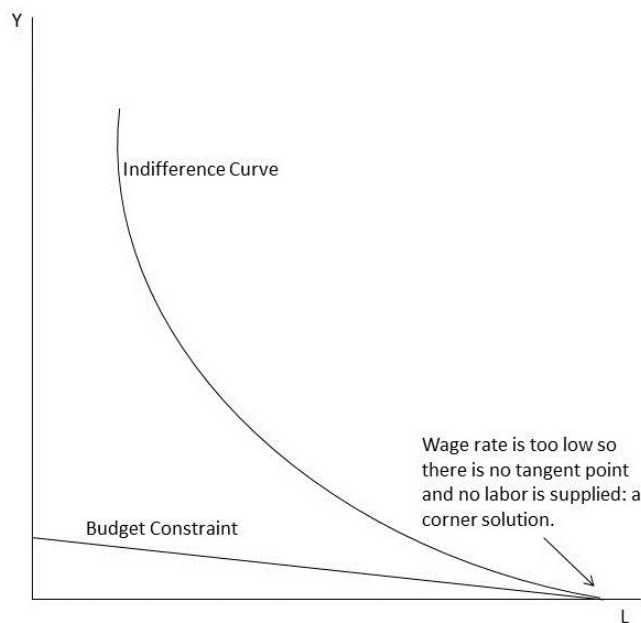


2.1 Corner Solution and Reservation Wage

The shape of the indifference curves in Figure 3.2 are of a standard shape, but note that the agent can always choose to work zero (thus consume zero goods) and consume only leisure - this is a 'corner solution'. This might happen if the wage rate is very low. The point of at which the indifference curve touches the x-axis will have a slope, and that slope represents what we call the *reservation wage*. This is the wage below which the agent will not work. Once the wage breaks this level the agent will choose to supply labor. Lets see a corner solution graphically in Figure 3.3.

Here the wage is too low to induce any labor to be supplied. The wage must break the slope of the indifference curve (MRS) at the x-axis intercept to induce labor - it must be higher than the reservation wage. (*Practice*: there could also be a reservation time of working. Think about someone who has to commute an hour to work. What would their budget constraint look like and what would the reservation wage/time look like? *Practice*: Instead of a commute time what if there was a cost to working, ex: buying materials, how would this budget constraint look. How do these two twists change what happens when the wage rate crosses the reservation wage?)

Figure 3.3: Corner Solution and Zero Labor Supplied.



Changes in the Wage Rate

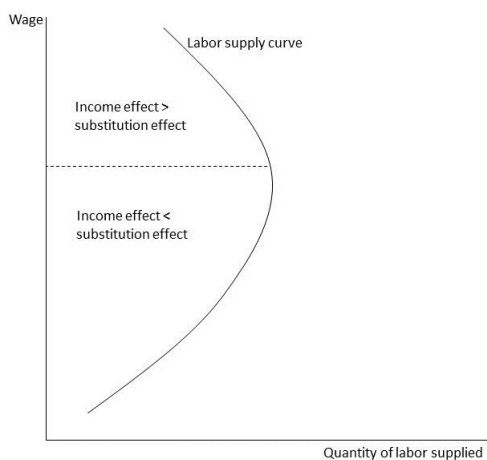
Lets see what happens when the wage rate changes. You may recall from your micro class that you will get two effects when the price of a good changes: an income effect and a substitution effect. Specifically lets look at the case when the wage rate increases.

Substitution effect: as the wage increases the opportunity cost of leisure increases. As leisure becomes more costly less of it is consumed and thus the agent works more.

Income effect: as the wage increases the individuals total income increases - they are 'richer' - and so will consume more of all normal goods. If leisure is a normal good (which we will in general assume) then they will consume more of it.

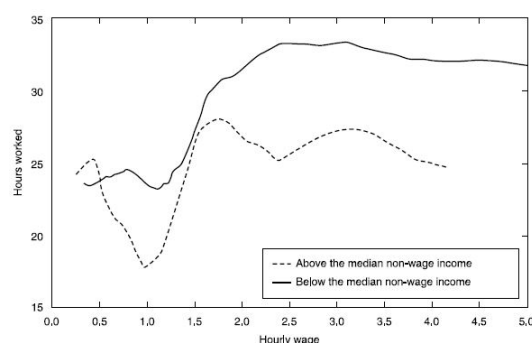
So we see that we have to opposite effects. What matters then is the net effect of the two. And this depends on which effect is larger. It is likely that at low wage/income levels the substitution effect dominates but once wages/incomes hit some level the income effect begins to dominate. This will lead to a 'backward bending' supply curve of labor as in Figure 3.4 (a). Moreover it also shows up in data on British single mothers - Figure 3.4 (b). If you ignore the very low wage levels there is a clear hump for those with high non-labor

income. And those with low non-labor income supply more labor. (*Practice:* Can you, using the budget constraint and indifference curves, isolate the income and substitution effect? First isolate the substitution effect - keep income constant but let wage change - then isolate the income effect - keeping the wage rate constant and letting income change.)



[Backward Bending Labor Supply]

[British Single Mother's



Labor Supply (source: Blundell et al (1992))

Figure 3.4: Bending Labor Supply Curve.

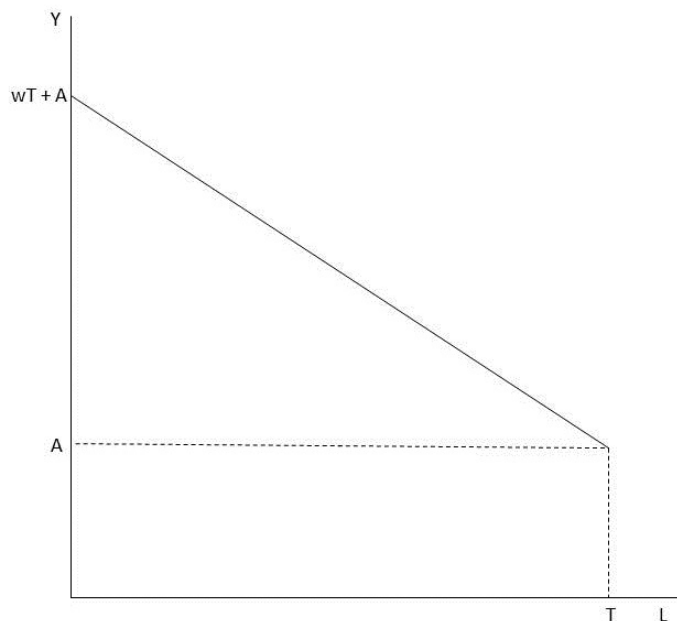
Non Labor Income

So far we have assumed if the agent does not supply any labor then they do not consume any goods - that is all income must come from labor. But it is very easy to incorporate non-labor income. For example maybe you receive earnings from investments or you have a spouse who works. This leads to a shift in the budget constraint. Lets denote the amount of non-labor income as A , then the new budget constraint becomes:

$$wT + A = Y - wL \quad (3.8)$$

So all this does is shift the budget constraint up. However note (Figure 3.5) that the total amount of time has not changed, so there is no longer an intercept with the x-axis.

Figure 3.5: Budget Constraint with Non-Labor Income.



Note that, if leisure is a normal good, then the introduction of non-labor income necessarily increases leisure. This is because since the wage rate does not change the introduction of A causes a pure income effect.

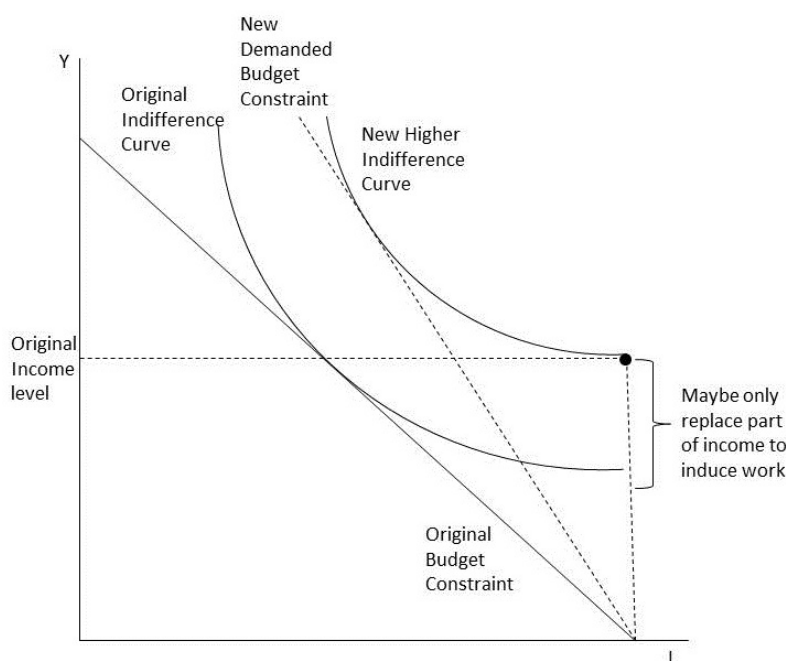
Program Evaluations

We can use this simple model of labor supply to think about some possible government/insurance programs and what the effect might be for labor supplied. Specifically we will look at income replacement insurance programs, traditional ‘income floor’ welfare programs, and more modern ‘earned income tax credit’ forms of income assistance.

5.1 Income Replacement Programs

Consider a simple form of unemployment or disability insurance. Say that if you lose your job or become unable to work (possibly temporarily) the program replaces all of your income while you are out of work but pays nothing once you return. What might this look like and how might this affect one's decision to return to work if they recover or possibly find a potential job. What we want to do is focus on the effect on the budget constraint - it is through the budget constraint that programs affect people's decisions. Figure 3.6 depicts the change to one's budget constraint from such a program. In this we assume the possible wage rate (from returning to work) is the same as before one was injured/fired. Note that there is a spike at zero hours of work - because once you start working the benefit goes away.

Figure 3.6: Income Replacement Program.

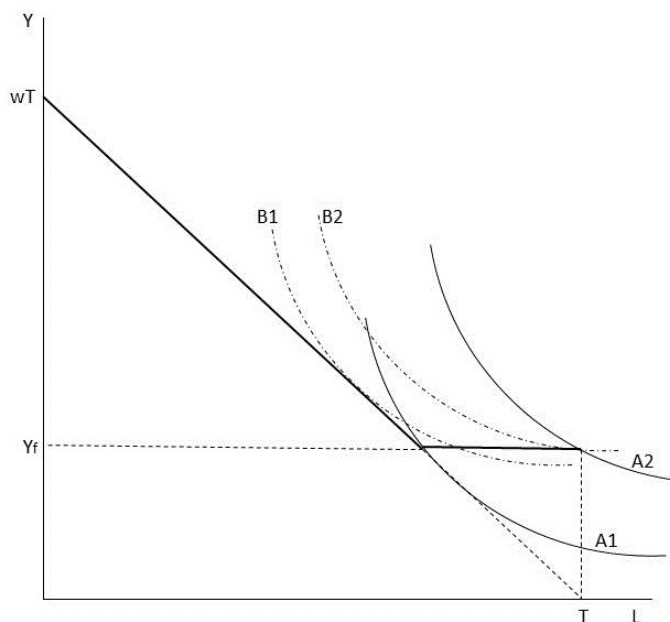


This spike increases utility from where it previously was (note the spike is at the height of the original income earned by now at zero hours worked). Thus to be induced to work again the agent will demand a higher wage. If we want to change these incentives so that there is still an incentive to work at the previous wage rate the program would need to only replace part of the lost income. So you can see there is already a conflict between 'helping people' and maintaining an incentive for them to work.

5.2 Traditional Welfare

Traditional welfare takes the form of an income floor. Basically nobody can earn less than some amount per month, and if they do the government fills in the gap. Say we denote the income floor Y_f , then the change in the budget constraint will look like Figure 3.7 where the budget line never falls under Y_f regardless of time worked:

Figure 3.7: Traditional Welfare Program.

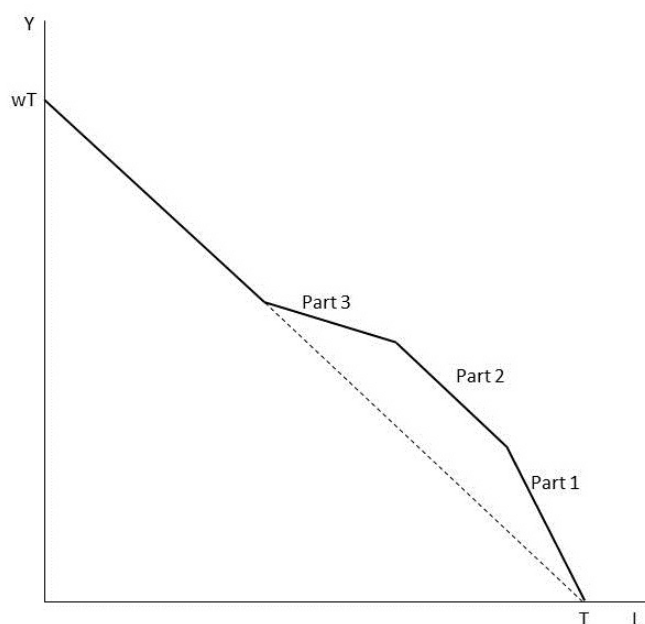


So in the range where the income floor dictates the income level of the agent there is a *zero marginal wage*. For individuals who previously were maximizing utility by working and receiving under Y_f , as indicated by indifference curve $A1$, they can now increase their utility by working zero hours and moving to indifference curve $A2$. Moreover even for some people who were earning above Y_f , as in indifference curve $B1$, they might still increase their utility by dropping out of working and moving to $B2$. So we see again there is a strong disincentive for (some) individuals to work. One possible way to overcome this is to have a work requirement such as requiring them to meet with an employment specialist for example.

5.3 Earned Income Tax Credit

Realizing the work disincentive inherent in traditional income floor welfare programs governments have sought alternative income assistance programs that (possibly) get around these issues. This has led to the implementation of what is known as earned income tax credits (EITCs). EITCs were introduced in Korea in 2008. The program works through the tax policy, as mentioned in the name. The idea is that at low income levels the government gives you a tax rebate dependent on your earnings (so it essentially increases your wage rate). Of course this can not continue for all income levels so it must taper off at some point. This is generally done in two parts: (1) for some income range, once the rebate has hit the highest level, you keep your rebate but do not get anymore with increased income (so in this range your wage rate is that of the market), then (2) at some point, once an income level is reached, the government begins to withdraw the rebate (this essentially causes your wage rate to fall). So this creates a modified budget constraint that can be seen in Figure 3.8.

Figure 3.8: Traditional Welfare Program.



In part 1 the agent's wage is higher than the market, in part 2 it is equal to the market, in part 3 it is less than the market, and above that the program has no effect on the budget constraint. So in all sections there is a positive income effect (because the program increases the agent's income) implying the program will lead to lower levels of work if leisure is a normal good. However, the substitution effect differs based on what section we are at. In section 1 the substitution effect leads to greater work, in section 2 there is no

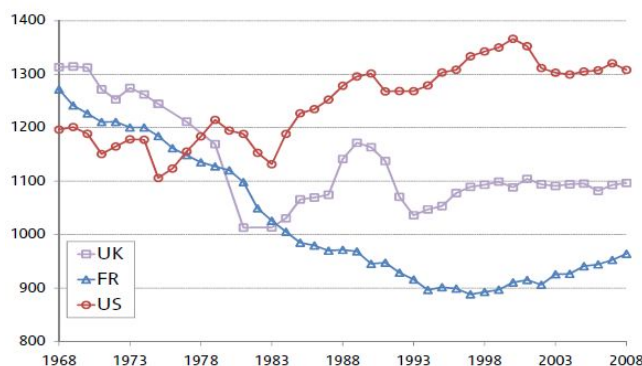
substitution effect, and in section 3 the substitution effect leads to less work.

So what are the net effects? In section 1 the effects work in opposite directions, so it depends on the size of each effect. However we do know that the wage rate is increased so this should induce more people to work (those who's reservation wage is above the market wage but below the program adjusted wage. And in section 2 and 3 we would expect a decrease in labor.

Extensive vs Intensive Margins

Since we are discussing labor supply decisions and how they relate to wages we can obviously discuss and try to estimate the supply elasticity of labor - and this is just a topic that occupies many labor economists. And though we will not discuss this in depth we will note a distinction - between an intensive and extensive margin of labor supply. The intensive margin is in regards to how much the average workers work - hours worked per worker. The extensive margin is how many people work. To see how important this might be for policy - such as income tax and unemployment benefits etc - lets look at some data for France, the U.K., and the U.S. First lets look at overall labor in terms of average hours for the whole population.

Figure 3.9: Average Annual Hours (population average).



We can see that up until about 1980 they were fairly similar with U.K. at the top, France in the middle, and U.S. in the bottom. After which there are big changes with the trends spreading out with U.S. on the top, U.K. in the middle and France on the bottom. So what caused this? Well lets look at this slightly different. Lets look at average hours per worker and the employment rates.

We can see here stark differences and also how U.K. reacts much like the U.S. in terms of employment rate,

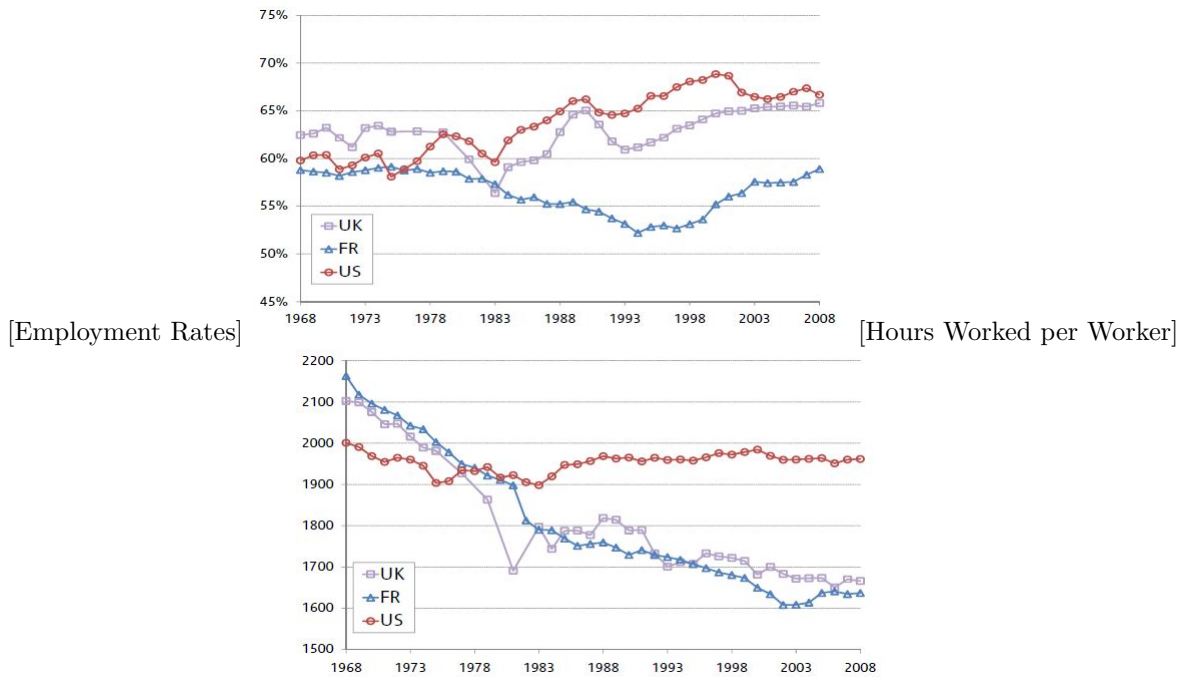


Figure 3.10: Measures of Labor Supply.

but much like France in terms of hours per worker. Of course an exact decomposition of average population hours into intensive/extensive margins is not simple (how much of those hours are due to ‘new workers’ and how much due to changes in ‘old workers’ - once someone starts working are their hours all part of the extensive or intensive margin?) For more on this (and the source of these charts) see “Labour Supply and the Extensive Margin” by Blundell, Bozio, and Laroque in *The American Economic Review: Papers and Proceedings*.

Chapter 4

Labor Supply 2: Household Production and Life-Cycle Models

Household Production

In the previous chapter we considered only two options: leisure or working and consuming some good. In reality time not spent in the labor market is many times spent ‘producing’ goods to be consumed, eg. cooking a meal. However, as an alternative, one could simply call Pizza Hut and pay for a fully prepared meal. So there is a deeper decision going on than just consumption vs. leisure. Here we will look more in depth at that decision. The following model is fairly simple, though may seem a bit complex for some. But once we lay it out in a formal way, we will quickly turn to looking at it graphically and in separate parts. So if the formal model and the math in it is a bit too complex do not worry, just skip the math if you need.

1.1 General Model

The agent still consumes only a good y and leisure l , so their utility is still $utility = u(y, l)$. They still have an option to work h_m units of time at wage rate w . The twist is that y must be ‘produced’ in the home with some mix of purchased products x ($x = wh_m$) and with time spent making it h_h (note again that we

assume the price of x is one). The production function is $y = f(x, h_h)$. The rest is pretty much the same as before. They still have T units of time, so they have a time constraint: $T = h_m + h_h + l$. So substituting in the time constraint for leisure and the market income for x we have our agent maximization problem:

$$\max_{h_m, h_h} u[f(wh_m, h_h), T - h_m - h_h] \quad (4.1)$$

Our two first order conditions yield:

$$\begin{aligned} h_m &: \partial u / \partial y \cdot \partial y / \partial x \cdot w = \partial u / \partial l \\ h_h &: \partial u / \partial y \cdot \partial y / \partial h_h = \partial u / \partial l \end{aligned}$$

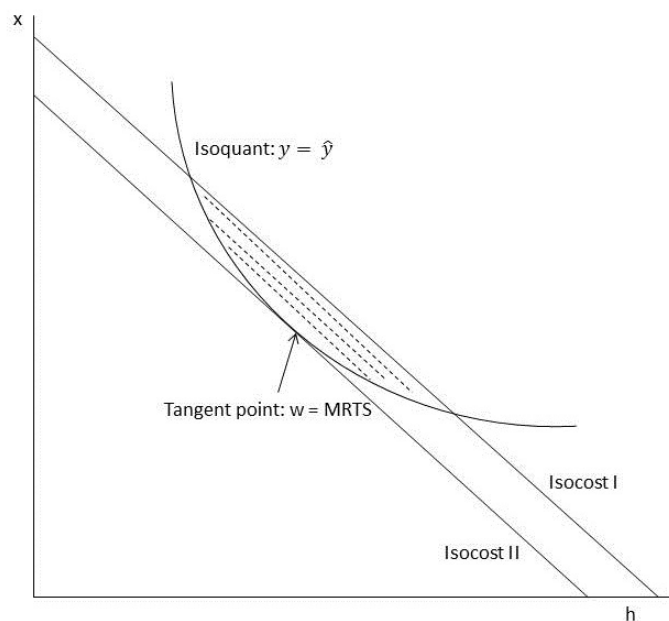
We can see this should make sense. Rewriting our first condition we see the ratio of the marginal utility of leisure to the marginal utility of the purchased product (x) should equal the wage rate (note the ∂y cancels). This is just what we had in our last model. The difference here is that marginal utility of the purchased good inherently contains the household production function (we do not consume x but use it to produce y). So if we are at some optimal choice set and an individual uses of inputs x becomes more productive in home production (in formal terms $\partial f / \partial x$ increases) then I will want my marginal utility of leisure to increase to purchase more inputs. And quite naturally, as in the second condition, the ratio of my marginal utilities of leisure and household production time should be equal.

1.2 Simplified Model

The above is the formal and correct way to think about the problem. And this can be extended to multiple y s that need to be produced before consumed. But we can think about a much simpler problem in which we omit leisure and assume the consumer only consumes one good y (but keep in mind they consume many different y s) and can produce that only consumption good according to the production function $y = f(x, h_h)$ and there is outside work for which they receive a wage w . So though our diagram will look identical to that in the previous chapter the interpretation is slightly different (see Figure 4.1).

Here we are minimizing cost subject to a fixed amount of output $y = \hat{y}$. For example what is the cheapest way to cook dinner given I can use time or work and purchase goods. So there are no indifference curves here, but rather ‘isoquant’ curves - mixes of the two inputs to production that yield the same output. And

Figure 4.1: Minimize cost subject to $y = \hat{y}$.

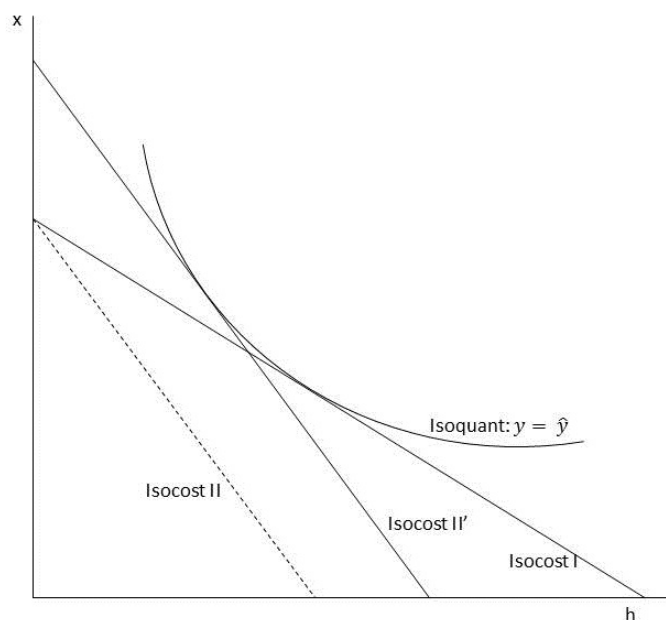


the slope of an isoquant is the marginal rate of technical substitution (MRTS). And rather than a budget constraint we refer to an isocost curve - every point on it reflects a bundle of time and inputs the cost the same amount, and the slope of the isocost is the ratio of prices (which is just w here since the price of x is set to one). So the resulting cost minimization point has all the same relations as maximizing utility. Isocost I does not minimize cost, the shaded area are mixes of inputs that can produce $\hat{y}$ that are cheaper. But once we reach Isocost II we have our cost minimization point. This is a simple model that helps explain the rise in fast food industry, why convenience stores are more expensive than grocery stores, why many people don't use coupons etc.

We can again think about what happens when the wage rate increases and analyse it graphically. First the isocost rotates around the y-axis because the price of x did change, only the cost of time. Thus for the same cost (definition of isocost) one can 'purchase' less time. But now to produce that same $\hat{y}$ output the cost is higher so we have to find a new minimization point that minimizes costs (which will have to be a higher cost because a price of one of the inputs went up). And we see that the new optimal point uses less time and more inputs to make our dinner (see Figure 4.2).

Now we mentioned that there are actually many goods being consumed. So what happens to the relative price of these goods after this change in the wage rate? Well time intensive goods become more relatively

Figure 4.2: Effect of change in wage on cost minimization.



more expensive. So not only will the agent change how they make these consumption goods, but they will shift to consuming less time intensive goods and more goods intensive goods. Both of these factors will act to increase labor.

We could also think about how changes in the production function might alter time use. Consider how the production might have changed for the average household over the last 60 years with the multitude of appliances that now fill the home. This has likely made it much easier to substitute purchased goods for labor. This implies the production function has changed shape. Think about how the slope at any given point would be different, what would this mean about the time spent on home production verses in the labor market? Similarly what happens when the price of inputs to home production fall? How might this effect the labor supply of women and the trend we saw in Chapter 1?

Household Specialization

We can also think about a model where a married couple decides together how much each should work and the level of leisure and consumption. As you may have guessed this model has several issues to examine. Is there a single utility function (a household utility) or two personal utilities? And if there is two separate

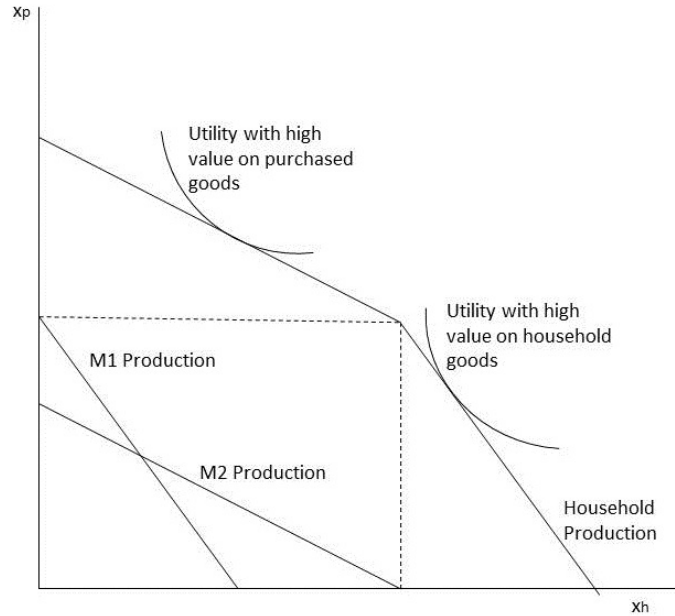
utilities should your utility enter my utility function? What should we incorporate in our utility function, leisure, purchased goods, home produced goods? If the latter what about production? Do they produce the good separately or can they produce together in some sense? All of these are important and interesting questions, though these are in general more than we will hope to do.

2.1 Simple Model

What we want is a very simple model that might serve as a basic view point on household specialization. Lets assume both members can work and purchase a purchased good x_p or stay home and make a household good x_h - note here, unlike in household production, we are assuming one does not need a purchased good to make the household good (it is purely a function of labor). Furthermore assume that the household function is linear for both members (though potentially different), thus, since the wage rate is constant, the production possibility frontier of each will be linear. Moreover lets assume one member receives a higher wage while the other is more productive in household production, and they have a single unified utility function: $u(x_p, x_h)$. We can easily depict their unified production frontier as in Figure 4.3. We can see if the household puts a lot of weight on household goods then member 2 will specialize in household production and member 1 will spend sometime in household production and some time in the labor market. If they put more weight in purchased goods then member 1 will specialize in the labor market and member 2 will divide their time between household production and the labor market. Note it is the person with the steepest production frontier (opportunity cost) that specializes in the labor market, not necessarily the one with the highest wage.

Here we have made two simplifications. One is regarding the linearity of the production of the household good. The other is the single 'household utility'. A more general model, and one that has been shown necessary for more in-depth analysis regarding taxation etc., is one in which each individual has their own utility function. The literature on this grew from labor-leisure with unitary utility, to labor-leisure with dual utility, to a model with dual utility and household production. But these are added complexities we will not address here.

Figure 4.3: Household Utility Maximization with Two Members.



2.2 Recessions and Labor Supply

We can take the above model and note some possible consequences for labor supply changes when there are wage shocks or possibly job losses by one of the members. We might see an *additional worker effect*. Consider a household where only one member works. Then there is a significant pay cut, or possibly a job loss and the earner receives a much smaller unemployment payment as long as the full time actively search for a job (we will ignore the details of this latter case and treat it as a pay cut). This could cause the member who previously did not work to join the labor force. This is represented below.

Now this implies that lower wages on average induce more to work, this is at odds with what we found in our most basic labor-leisure model. And in fact empirically this additional worker effect is small and we see labor force participation falling in recessions. Though we do not have any time aspect in our model here we can think about someone who is unemployed but searching for a job - they are in the labor force. Furthermore consider they are searching (a concept we will consider further in later chapters) on the basis of an *expected wage*. So their search is based on the wage they will get if they find a job and the probability of finding a job: $E(w) = \pi w$ where π is the probability of finding a job. A recession likely lowers both of these. So we would expect some people to drop out of the labor force. This is the effect that is largest in driving the employment data we see during recessions.

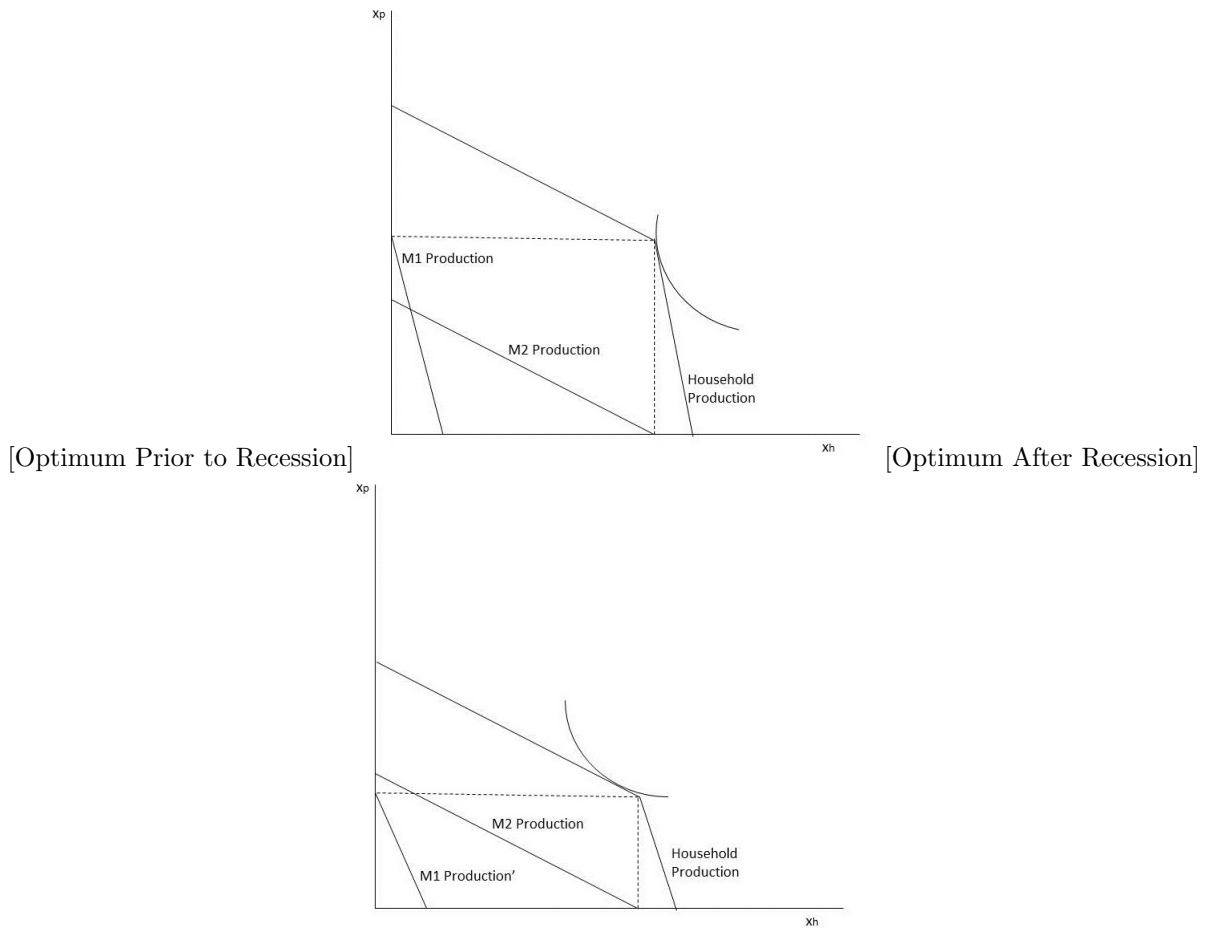


Figure 4.4: Additional Worker Effect from Wage Shock.

Life Cycle Considerations

Up until now we have been looking at ‘static models’ in that there was no time in the sense of multiple periods in which decisions were made. We might instead consider an individual who lives T periods and maximizes utility based on consumption and leisure over these periods. For example $U = u_1(c_1, l_1) + u_2(c_2, l_2) + \dots + u_T(c_T, l_T)$ where there is some interest rate (thus the individual can borrow or save money) and some wage rate. But these simple models run into problems matching the data we see, in particular if they are to be the bases of why we see large changes in employment levels in response to macro shocks. For one it would need to explain how small changes in wages lead to large changes in labor supply - this requires parameters for utility functions that don’t match empirical studies (of the intertemporal elasticity of substitution of leisure). Furthermore, the model would predict that in times of high consumption (good times) we would see high leisure (low labor force participation). But we actually see the opposite. In recessions we see consumption

and labor force participation rate drop. More complex models, ones with some type of labor market frictions, are needed. We will not consider these here but simply look at some general stylized facts and why we might expect them.

3.1 Life Profiles

Wage profiles and labor force participation rates then take the form depicted in Figure 4.5(a). We see wages increase in early life, as do participation rates, then flatten off. This is primarily the result of both education and training. At the same time we see participation rates peak in middle age and then drop off near retirement ages (possibly the disutility of working increases in late age). For some countries/cohorts we see a different shape for women in their participation rates. Some cohorts have a dip in mid-life as in Figure 4.5(b). This is due primarily to children and their associated peak in marginal product of household production.

We can see in Figure 4.6(a) how women's participation profiles have changed over time in the U.S. This change can be rationalized due to the increased wage profile of women from latter cohorts as well as changes in the household production function (see previous sections). Figure 4.6(b) depicts the participation profile of women for several countries. In particular note that Korea has the lowest profile and a relatively strong dip in the mid-30s. This is in stark contrast to Sweden where we see high rates and no dip at all.

3.2 Retirement

Retirement is an important part of people's lives and a big driver of savings rates. Some countries have mandatory retirement ages, some allow companies to dictate their own firm-specific retirement age, and other countries expressly disallow any mandatory retirement age based on anti-discrimination laws. Furthermore most countries have some national pension scheme which pays some percentage of lifetime average earnings (usually 40-60% assuming a minimum years of working) once an individual reaches a certain age. These 'percentages' are usually actually some complex formula relating years one contributed, contributions, and age of retirement. Many companies also have their own additional pension scheme.

In the U.S. mandatory retirement is illegal, and individuals are eligible to collect national pension payments

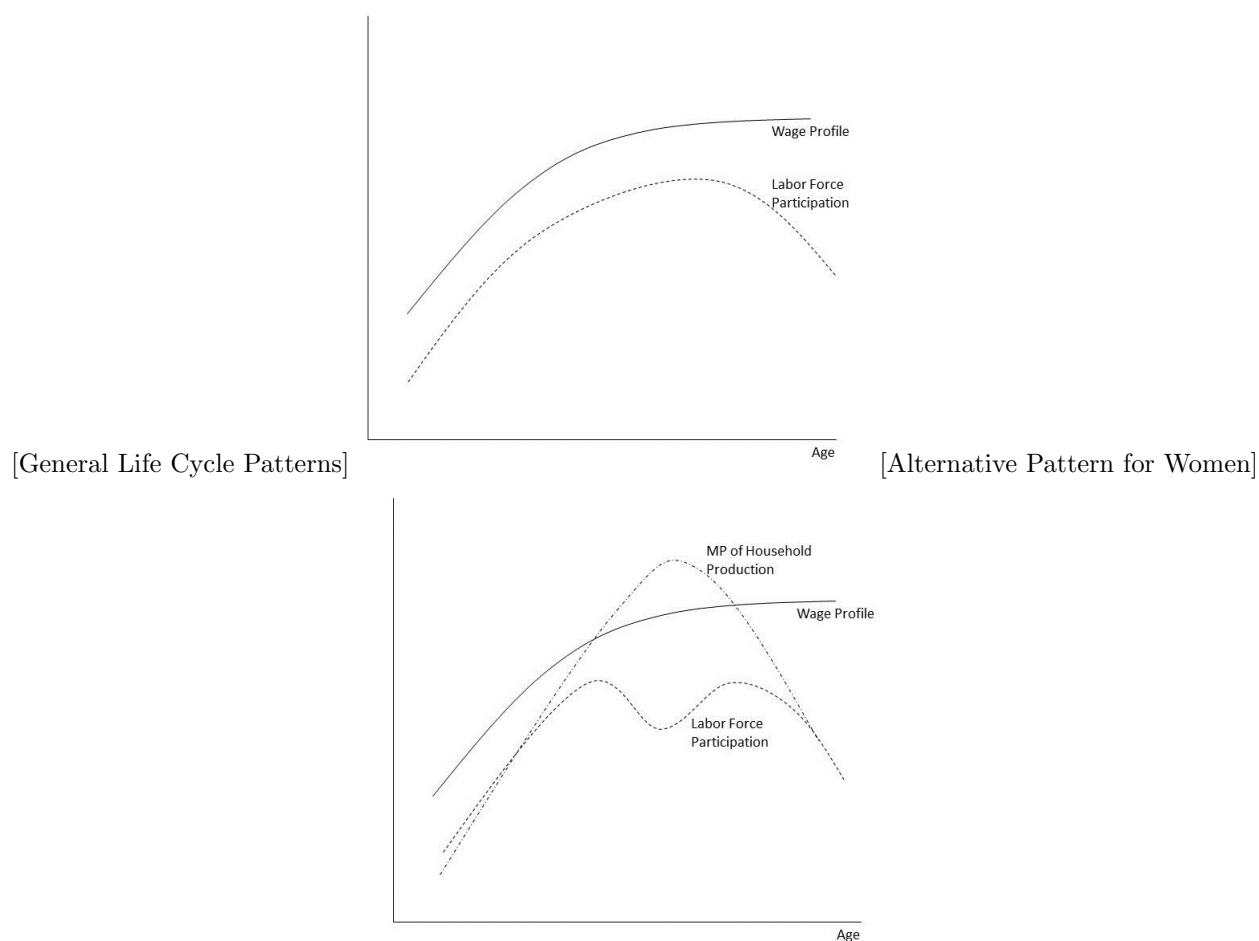
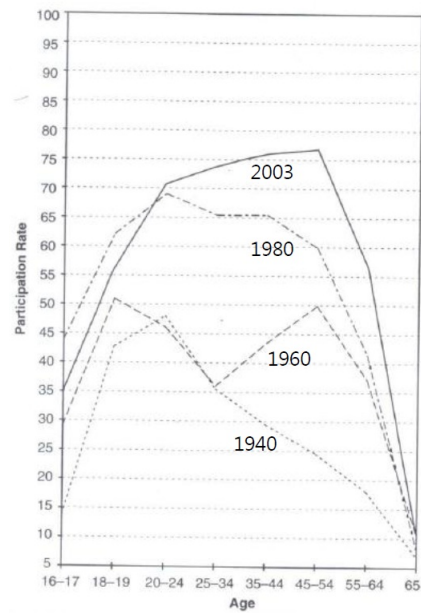


Figure 4.5: Life Cycle Profile of Wages and Labor Force Participation.

at 62 (early retirement) and receive full pension payments if they wait to collect until 67 (full retirement age) with increases if they differ to 70. In Korea, firms may set a mandatory retirement age, though it must not be below 60 (this was announced in 2013 and firms have until 2016 to comply). Individuals may collect national pension payments at age 60 (this age is being phased up to 65 by 2033). It has been common for companies to mandate retirement around 55 and push for early severance packages for individuals earlier - this can be seen in the fast declining participation rate for those 50+ (Figure 4.7).

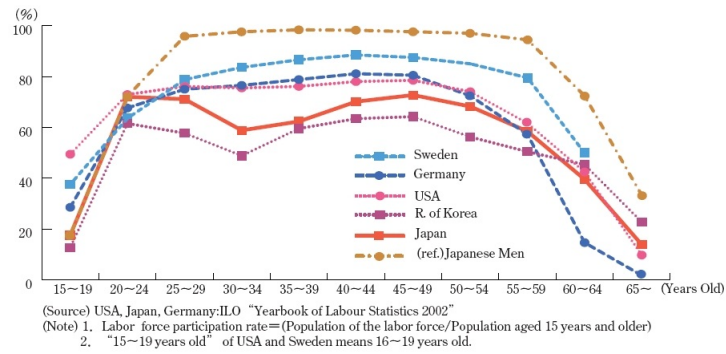
Despite this practice, the 'effective retirement age' - the age at which the individual withdrew completely from the labor force - was 71 for Korean men and 70 for Korean women in 2012, the second highest of OECD countries (thus different from official data). This is partly a reflection that Korea's pension system is relatively new - set up in 1988 - and the average payout for a retiree is about 300,000 won a month while many do not qualify because they did not enter the system. This is growing concern in Korea, especially in light of a fast aging population (as reflected in the governments new minimum retirement age). Though



[Changes in Women's Profile in the U.S.]

[Women's Profile in

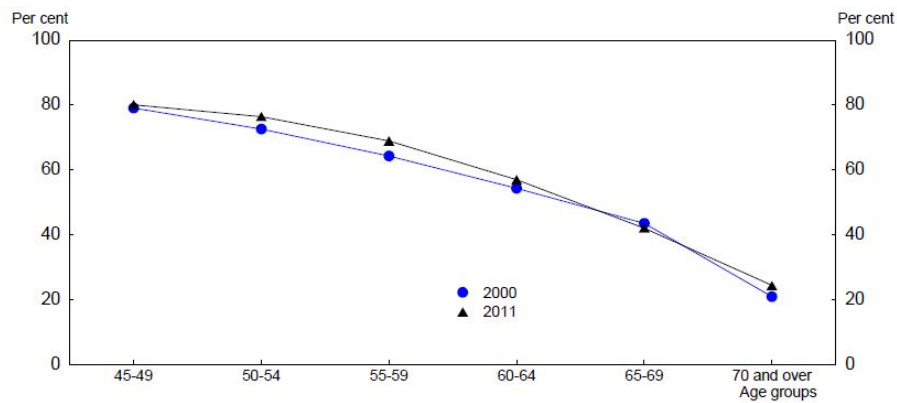
(Figure 9) International Comparison : Women's Labor Force Participation Rate by Age Bracket



Various Countries]

Figure 4.6: Women's Labor Force Participation Profile.

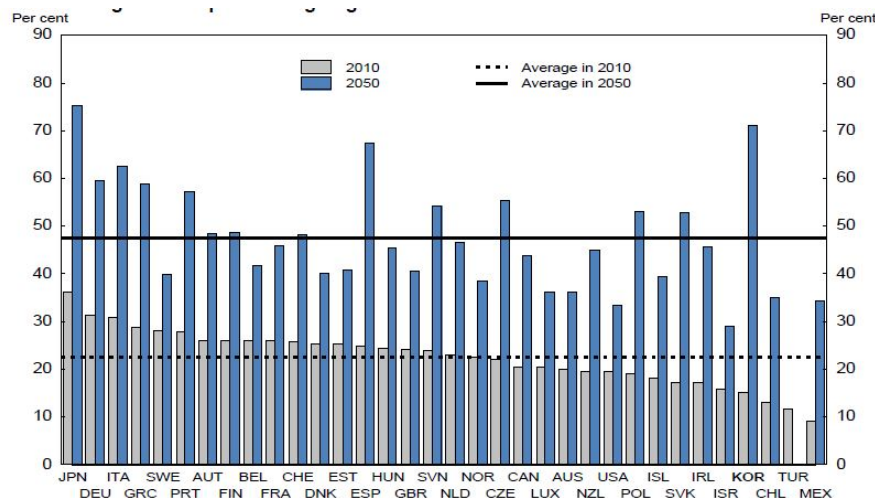
Figure 4.7: Older Workers Labor Force Participation.



Source: OECD Employment Outlook Database.

relatively young now, by 2050 it is expected to have one of the oldest largely due to a low birthrate (Figure 4.8).

Figure 4.8: Elderly Dependency Ratio.



1. The elderly dependency ratio shown in this figure is defined as the over-65 population as a share of the 20-to-64 population.

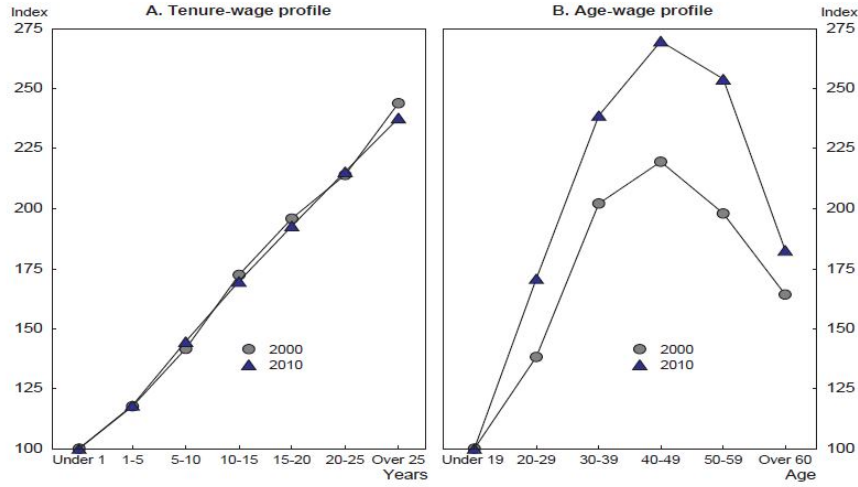
Source: Statistics Korea, *Population Projection for Korea* (2011 version) and *OECD Demography and Population Database*.

Does forced retirement make sense in Korea? Some argue this ‘frees up work’ for younger generations (a term commonly called by economists the “lump sum of labor fallacy” - the idea there is only ‘so much work to do’). Others point to seniority based pay schemes in Korea - meaning the oldest workers earn the highest wages despite possible declining productivity late in life (firms surveyed indicate they do not have the expertise to evaluate workers any other way, though there are likely cultural aspects to this). Note in Figure 4.9 (a) the steep tenure-wage profile in Korea and contrast that with Panel (b) which sees a peak around 50 and subsequent steep decline reflecting mandatory retirement and movement of older workers to secondary markets.

Moreover, given strict labor laws, mandatory retirement gives Korean firms one of the few ways to legally reduce their labor force. Though they would likely retain older workers if they could lower their wages, a mix of efficiency wages/culture may be blocking this from happening and maintaining older workers. A survey of firms indicates just such issues are high on the reason for not wanting older workers (see Figure 4.10).

Furthermore, a simple model by Edward Lazear (“Why is there mandatory retirement?” *Journal of Political Economy*) argues it could be efficient. Firms might structure contracts with wages under marginal product early on that then rise above marginal product later. The idea is that this will induce workers to provide their full effort to the firm. They do this because they know if they do not they may be terminated for

Figure 4.9: Korean Wage Profile.



1. Wages for 19-year-olds and younger and for less than a year are set at 100 in each year.

Source: Ministry of Employment and Labour, *Wage Structure Survey*.

Figure 4.10: Reasons Firms Reluctant to Hire/Maintain Older Workers.

Reasons	Per cent
Low adaptability to change	57.3
Lower work ability and capacity	44.8
High wages relative to productivity	43.1
Difficulty in assigning to posts	39.7
Unable to perform difficult tasks	32.9
Little motivation or enthusiasm for new work	25.8
Difficulty in accepting instructions	19.9
Frequent accidents	8.2
Lack of ability to co-operate with other workers	6.3

1. The survey included 648 firms. Firms were allowed to give three answers.

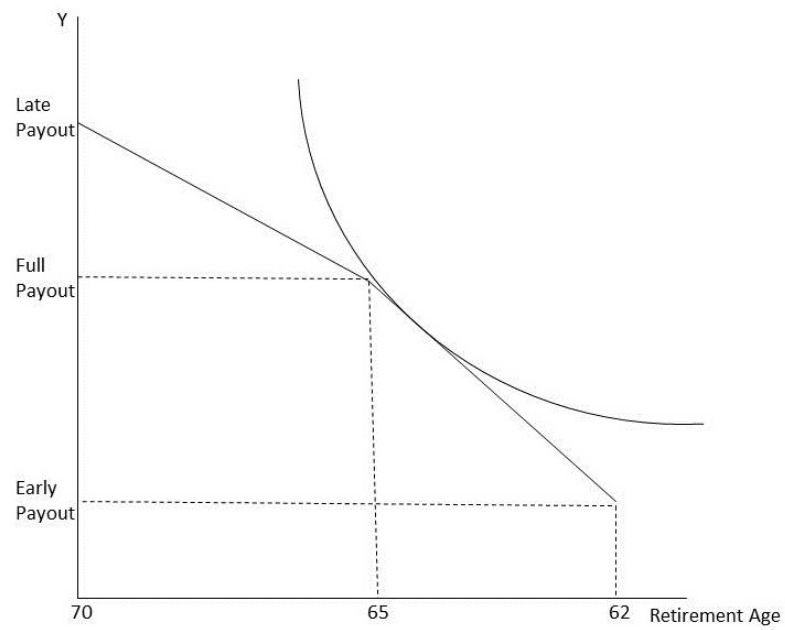
Source: Korea Labor Institute, *Survey on Firms Implementing the Wage Peak Compensation Scheme 2008*.

“cheating” and they must work the full contract to receive their full “value payment”. Thus wages at the end of the contract will be above marginal product. After this point the firm will terminate the contract. However, a full treatment of retirement and its issues is beyond this text as are the issues of pensions.

But we can still think about one decisions regarding age to retire if they face a U.S. style pension structure. Lets simplify this greatly (ignoring part-time work in retirement and probability of death for example) and just assume the agent has a utility function based on lifetime earnings and lifetime leisure (this is in essence retirement age): $u(y, l)$. Furthermore the pension structure creates an opportunity cost to retiring before 70. We can use our simple model to think about what might happen to changes in the pension system. Lets look at this graphically as it entails the exact same issues in our standard labor-leisure model. Furthermore the same income and substitution effects should be considered when thinking about expected changes in labor

supply in response to changes in the payment system (see Figure 4.11).

Figure 4.11: Retirement Decision.



Chapter 5

Education and Human Capital

Some Facts

Korea has some of the highest levels of education in the world. The percent of high school graduates entering college peaked around 2007 at over 80% (though has declined slightly since) as can be seen in Figure 5.1 (a). This is out of 95% of students that graduate High School and is a dramatic rise from the around 35% in 1990. Compare that to the 75-80% high school graduation rate for U.S. and that only about 65% of them attend college (with much higher college attrition rates). Part of this is due to Korea's dramatic jump into the international economy, though there is more to it than that - a complicated question we will come back to. Also the levels of spending on education in Korea have increased substantially. From about 8% of household spending to nearly 12% (see Figure 5.1 (b) - and note this is an average, it is estimated those with school-aged children spend as much as 25%). More dramatic is the rise in private spending which soared from about 20% to now around 80%. All of this shows - Korea consistently ranks near the top in international standardized tests. But there is another side - the amount of students attending vocational high schools has steadily dropped, and is expected to continue to do so, leaving a gap in the supply of certain skills. Furthermore the drive in costs has (in part) resulted in the low birth rate as nearly 60% of couples report educational costs as a major obstacle to having children

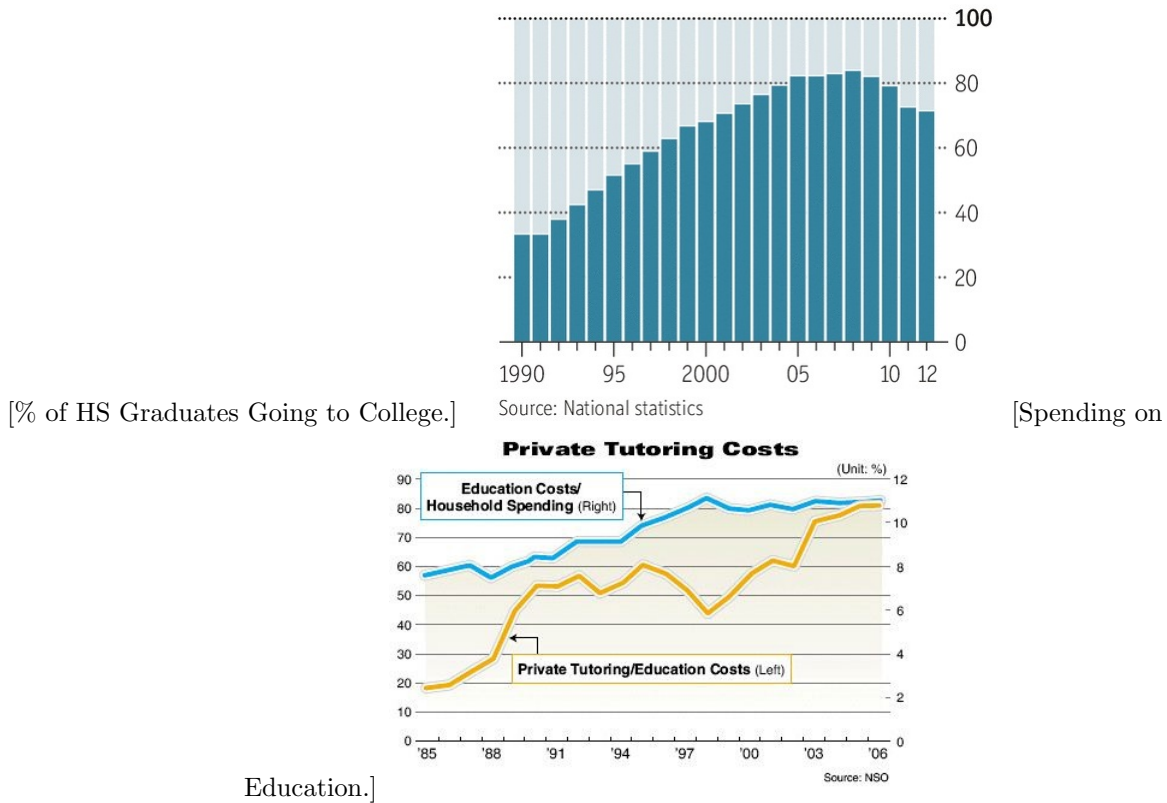


Figure 5.1: Education Trends in Korea.

Choosing Education and Selection

2.1 A Simple Model

Say there are two periods. A parent lives in the first, works (receiving y), consumes, saves and decides whether or not to send their child to school (so $e = 0,1$). The child lives and works in the second period and consumes. Furthermore the parent cares about the child's consumption, thus their utility is: $u = \ln(c_i) + \ln(\hat{c}_i)$ where $\hat{c}_i$ is the child's consumption. There is heterogeneity across children in such that the cost of education (θ) varies. Children that receive education earn w_s and those that don't earn w_u .

Assume there are no credit constraints and there is a single prevailing interest rate r . The choice of the parent is to maximize their (discounted)¹ utility by choosing education and consumption levels subject to

¹Discounted refers to the idea that a dollar tomorrow is worth less than a dollar today. For example if there is an interest rate of 5% then a dollar today is worth 1.05 dollars tomorrow ($\$1 \times 1.05$) and a dollar tomorrow is worth $0.952 = \$1/1.05$.

the constraint:

$$c_i + \frac{\hat{c}_i}{1+r} \leq \frac{w_u}{1+r} + e_i \frac{w_s - w_u}{1+r} + y_i - e_i \theta_i. \quad (5.1)$$

Note with no credit issues we can deal with this as a single discounted period problem, and it is assumed that the parents ‘control’ all of the money (they borrow on behalf of the children possibly). Note also that e is not in the utility function (it is not a directly consumed good) but only in the budget constrain. Thus the decision of e is solely to maximize the budget constraint (this is basically the ‘separation theorem’ for human capital investment). So education only occurs if:

$$\theta_i \leq \frac{w_s - w_u}{1+r} \quad (5.2)$$

ie. the cost of education for the child is less than the discounted gain in wages. It seems reasonable to assume this cost is inversely related to the child’s ability level, thus only high skilled individuals (low cost of education) go to school. Of course this model ignores likely credit constraints which would lead to more investment with higher income levels of the parents. These are interesting problems and ones addressed in the literature but beyond what we will do in this course.

2.2 Selection

Assume there are many individuals who have some unobserved ‘type’ z (uniformly distributed between 0 and 1) which possibly characterizes their innate ability. They decide wether to obtain education, at cost c , or not. Their wages are determined as:

$$\text{No Education : } w_0(z) = z \quad (5.3)$$

$$\text{With Education : } w_1(z) = \alpha_0 + \alpha_1 z. \quad (5.4)$$

Assume $\alpha_0 > 0$ and $\alpha_1 > 1$. The first means everyone gets higher wages for more education, the latter means more skilled individuals get a higher premium (note that with no education we have $w_0(z) = 1 \cdot z$). This is just a different formulation of the θ varying above - both ensure higher skilled are more likely to obtain education. Again one only obtains education if it is worth it, thus only those with a $z \geq z^*$ choose education where²:

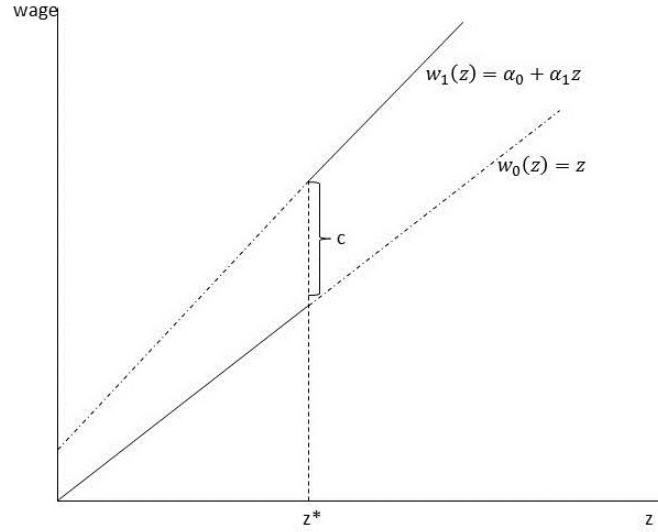
$$z^* = \frac{c - \alpha_0}{\alpha_1 - 1}. \quad (5.5)$$

²This is found by setting the benefit from education to be greater than c and solving for z .

Lets assume this is between 0 and 1 (otherwise either everyone or no one does and thats not very interesting).

The distribution of wages can be seen in Figure 5.2.

Figure 5.2: Selection in Education: Distribution of Wages.



This difference in the two payouts at z^* is c . Since individuals are uniformly distributed and everything is linear, the average wage of the groups is the wage of the ‘middle’ person. Thus:

$$\begin{aligned}\bar{w}_0 &= \frac{c - \alpha_0}{2(\alpha_1 - 1)} \\ \bar{w}_1 &= \alpha_0 + \alpha_1 \frac{\alpha_1 - 1 + c - \alpha_0}{2(\alpha_1 - 1)}\end{aligned}$$

People without education are those with $z \in [0, \frac{c - \alpha_0}{\alpha_1 - 1}]$ and the average person has $z = \frac{c - \alpha_0}{2(\alpha_1 - 1)}$. Repeat this for those with education to get the above payouts. Ok, now what type of question might we be interested in? Perhaps ‘how much does education increase wages for some group of people?’ Define $\bar{z}$ as the average ability of this group of interest, then the gain from education for the average person would be:

$$w_1(\bar{z}) - w_0(\bar{z}) = \alpha_0 + (\alpha_1 - 1)\bar{z}. \quad (5.6)$$

So what do we get by looking at the observed differences in wages? Well lets write the wage differential as:

$$\bar{w}_1 - \bar{w}_0 = \alpha_0 + (\alpha_1 - 1) \left[\frac{c - \alpha_0}{2(\alpha_1 - 1)} \right] + \frac{\alpha_1}{2}. \quad (5.7)$$

Note the first two terms give the potential wage gain for the low ability group, but then we have some other ‘remainder’. This remainder reflects the fact that those who chose to get education would have had higher wages anyway (ie. they have different ability levels), and just looking at the difference in wages overstates the effect of education. This is our ‘omitted variable bias’ spelled out with an economic model.

2.3 Why Higher Wages

We have not said *why* education leads to higher wages, just that those with education receive them. There are two theories when it comes to views on the role played by education in the labor market (though they are by no means mutually exclusive). One views the role of education as increasing human capital - that is investing in yourself increases productivity³ (this is the theory that became popular with Becker (1964)). The other views education as essentially a filter or signalling process - that is by observing your educational outcomes employers can learn about your innate ability (this is the theory put forward by Spence (1973)). Of course the likely reality is that both play a role. The question then becomes what role and how large of a role do each of these theories play, and do they differ for different groups.

Human Capital

We will distinguish between general and specific training/education. General can apply to any firm whereas specific can only be applied to a particular firm.

3.1 General Training/Education

Since general training can be used at any firm then a firm will not pay for general training. Say they did invest in you and now your MP is higher, well they cannot recoup their investment because you will demand your now higher MP for your wage because the next firm is willing to pay it (if markets are competitive). So the worker will have to pay for general education/training - in other words firms don’t pay your college tuition.

³There are several theories about what exactly we mean by ‘human capital increases productivity’, but they are not important for us.

Say you invest in yourself a level i of general training and your resulting marginal product is $y(i)$. Assume the training happens at the beginning of time at a cost of c_i per unit and the worker works two periods receiving their marginal product. Furthermore there is a discount rate β (β is less than one reflecting a dollar tomorrow is worth less than a dollar today, and in some instances (particular asset pricing) is defined as $\frac{1}{1+r}$ where r is interest rate). So the individual maximizes their payoff by choosing their amount of education:

$$\max_i y(i) + \beta y(i) - c_i i. \quad (5.8)$$

Maximizing investment yields optimal amount of investment implicitly defined by:

$$y'(i^*)(1 + \beta) = c_i. \quad (5.9)$$

So we see the individual will invest up to the point the discounted marginal return (which is the marginal change in their marginal product from investing) is equal to the marginal cost of investing. And the resulting lifetime payoff is: $y(i^*)(1 + \beta) - c_i i^*$. Moreover, individuals will differ in their investment in education dependent on their function $y(i)$ and their cost c_i which may differ reflecting innate ability as discussed in the previous models. So in this setting education raises wages because it directly increases marginal product - this is why education is many times referred to as ‘human capital’ and education is termed ‘human capital investment’.

3.2 Specific Training

Say now the firm hires worker at time zero and they work for two periods. The firm can invest in the worker thus raising their MP for them but no other firm (and assume there is some cost to the worker to move firms, thus since wages are competitive they will not move⁴). The firm's decision is the amount of investment that maximizes profits: $(y(i) - w)(1 + \beta) - c_i i$. First order condition yields our optimal level of investment: $y'(i^*)(1 + \beta) = c_i$. So they also invest up to point discounted marginal product equals marginal cost. What about wages? Well since competitive market profits are zero so: $(y(i^*) - w)(1 + \beta) = c_i i^* \rightarrow w = y(i^*) - c_i i^* (\frac{1}{1+\beta})$. And their discounted lifetime earnings are: $(y(i^*) - c_i i^* (\frac{1}{1+\beta}))(1 + \beta) = y(i^*)(1 + \beta) - c_i i^*$ (exactly what they previously earned).

⁴This is of course a strong assumption that makes things simple. There are interesting issues when this is more realistically modeled. This then requires some optimal contract. We will leave this interesting question for later.

And it should be noted that both of these (which are the same) are the social optimal level of investment. If we were to extend these simple two period models (to include multiple periods and ones in which earnings must be forgone to educate) we would also find the agent have reasons to attain all education in the beginning of life - that way they have their higher MP for longer - and at some point they will stop investing. This would imply a steep earnings profile which levels off - which is what we see. So it seems the human capital theory of education/training seems to match what we see in the world.

Education as a Signal

But the human capital theory may be wrong. Just because we see individuals with higher education levels receiving higher wages does not mean that the education increased their marginal product. Perhaps good standing in education settings merely signal to employers potential workers' innate ability which is directly related to their marginal product. So education does nothing to increase the students marginal productivity, it merely acts as a signal to firms as to the students potential marginal product if hired. Though perhaps the difference seems merely academic, they come to completely opposite conclusions on efficiency. We will see in the signalling model workers overeducate themselves.

Say there are two types of workers based on their innate ability/human capital/MP level - h^h and h^l . Furthermore education can be obtained (s) at a cost ($c_s = \frac{s}{h}$) but it does not affect MP. Workers obtain utility of:

$$u = w - c_s. \quad (5.10)$$

There is no leisure or disutility of work. There are three stages:

1. workers invest in education or not (this is instantaneous here so no loss of income etc.)
2. wages are offered
3. workers accept (or not) and work

If everything is competitive and firms can observe workers' MPs, then workers simply receive their MPs ($w^h = h^h > h^l = w^l$) and no one invests in education.

What if ability is not observed? Suppose there were no signal available. Furthermore say the proportion of the workers with high ability are δ and the proportion with low ability is $1 - \delta$. In this case the firms will simply randomly select workers and offer a wage equal to the expected (average) MP (recall firms are competitive): $w = \delta h^h + (1 - \delta)h^l$ and this is the wage accepted by both types.

However if ability is not observed then education becomes a way for high ability people to signal who they are to the firms. We can arrive at a *separating equilibrium* if there is some signal such that only high ability workers send a signal (s^*) and low ability type do not. Lets see what such an equilibrium might look like.

Assume low ability workers did signal they are high ability, then they would receive $u = h^h - \frac{s}{h^l}$, but if they do not signal they receive h^l . So for an equilibrium it must be that:

$$h^h - \frac{s}{h^l} \leq h^l \rightarrow s^* \geq (h^h - h^l)h^l. \quad (5.11)$$

Now high ability workers will want to send the lowest signal possible $s = (h^h - h^l)h^l$. And since low ability workers were indifferent between sending s^* and not, then high ability workers must strictly prefer sending s^* (because their cost of education is strictly less). Thus high ability workers send the signal and receive the high wage while low ability workers refrain and receive the low wage. Note that the higher wages are *not* based (from the firms perspective) on higher ability, but rather on the education signal.

Are we better off with signalling? Well from a social viewpoint certainly not. With and without the signal everyone works so output is the same, but with the signal there is energy wasted obtaining the signal. What about for each type? Well the low ability types strictly prefer the pooling or no-signalling equilibrium because they receive a wage of $w = \delta h^h + (1 - \delta)h^l$ rather than h^l . So what about high ability types? Do they prefer the signalling? Well it depends on whether the following inequality holds:

$$h^h - \frac{(h^h - h^l)h^l}{h^h} \geq \delta h^h + (1 - \delta)h^l$$

Which only holds if⁵

$$\delta \leq \frac{h^h - h^l}{h^h}$$

which means their proportion of the population must be smaller than their relative ability gap. So we see the signal model is inefficient, low ability workers strictly do not prefer it, and high ability workers only possibly

⁵ $h^h - \frac{(h^h - h^l)h^l}{h^h} \geq \delta h^h + (1 - \delta)h^l$ implies $(h^h)^2 - h^l h^h + (h^l)^2 \geq \delta (h^h)^2 + h^l h^h - \delta h^l h^h$ which implies $(h^h - h^l)^2 \geq \delta (h^h - h^l)h^h$

prefer it (though they must signal because with a potential signal the no-signal equilibrium is unstable).

4.1 Extensions

It is possible that a signalling model could be efficient. Though not a very interesting model (in my opinion) I will lay out the basics nonetheless. Assume again there are high and low types and some education signal which is costly and again costs are higher for the low types. However, now assume there is some disutility of labor (or possibly leisure in the utility model). If there are enough low types the wage offered in a pooling equilibrium may be too low to induce anyone to work. However if there is a signal available then the high workers will attain it, and be offered a high enough wage to overcome the disutility of work and offer their labor.

Alternatively we can construct a model where education both increases human capital *and* acts as a signal because type is unobserved by the firms. Consider again two types - high and low - who have the option to attain education: $e \in [0, \infty)$. The education is costly and increasing so. Specifically $C(0) = 0$ and $C'(e) > 0$, furthermore $C'_l(e) > C'_h(e)$ so at every point the slope of the cost function is higher for the low types (and thus always above). Also there are gains to MP ($= Y$) from education. For simplicity assume this is linear, with a zero intercept, and higher for high types. These are depicted below in the left figure.

Now first lets look at the optimal amount of schooling. Well we have seen this before it is just where $C' = Y'$. Denote these values of education e_l^* and e_h^* . We can see where these are on our figures by shifting up the cost functions until they are tangent with our Y functions. Now the question is what if type is not observed (thus neither is $MP = Y$) but only education (though firms do know the functions $Y(e)$). Can we find two levels of education ($\hat{e}_l$ and $\hat{e}_h$) such that only low types choose the lower and only high types choose the higher (note there may be many equilibrium). This way the high types can be distinguished and receive higher wages based on their higher human capital function Y (if low types might also choose the high's level of e then again firms cannot distinguish).

Assume that $\hat{e}_l = e_l^*$ and thus $w_l = Y_l(e_l^*)$. Now what is the lowest level of e for the highs such that the lows have no incentive to choose? Note if they do choose the high's level of e they will be thought a high and paid $Y_h(e)$ - this is what we don't want to happen, the signal must work. Whatever this level is, denote

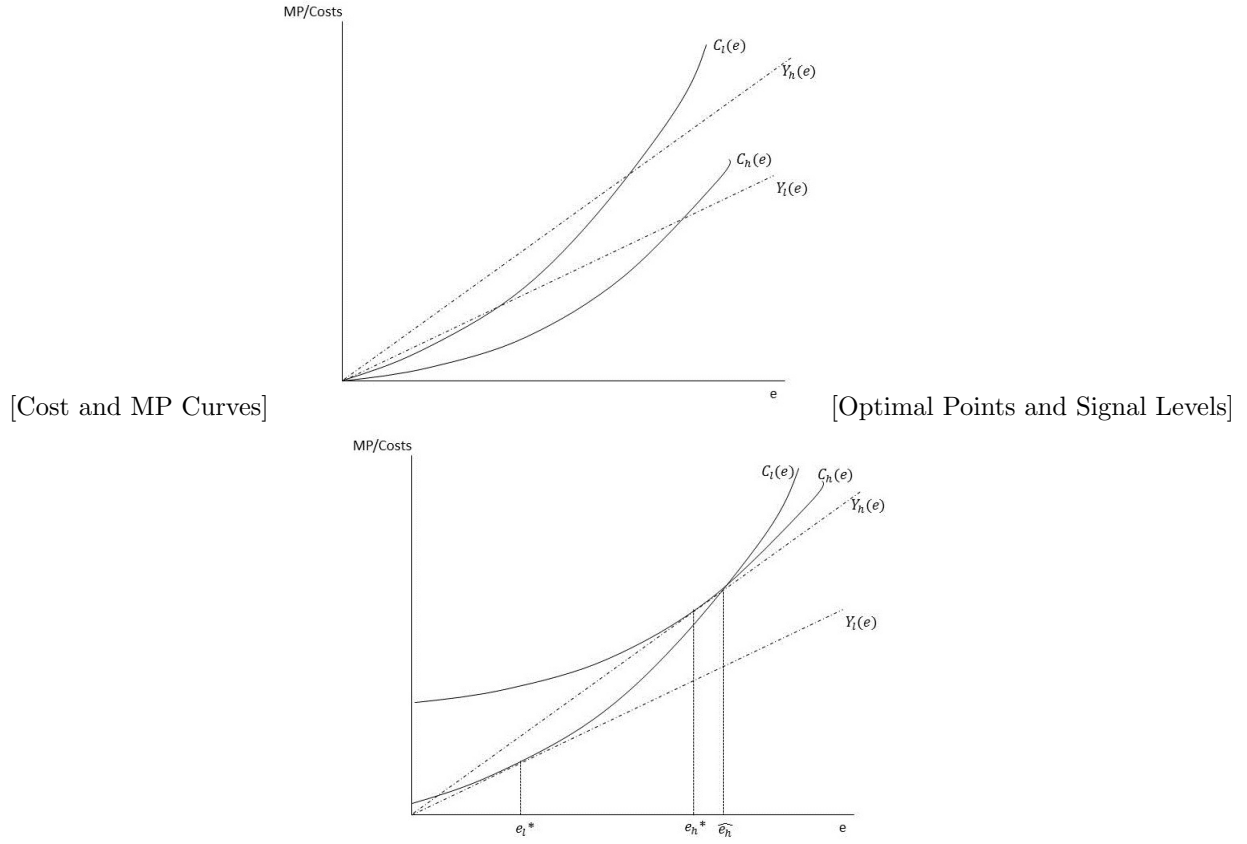


Figure 5.3: Education as a Signal and as an Increase in Human Capital.

it $\hat{e}_h$, we have the restriction:

$$Y_l(e_l^*) - C_l(e_l^*) = Y_h(\hat{e}_h) - C_l(\hat{e}_h). \quad (5.12)$$

The left side is what the lows get if they choose their designated level e_l^* . The right hand side is what they get if they did choose the highs level $\hat{e}_h$ - they get the high's payoff but pay thier cost. Now since they are equal and thus there is no benefit to deviate from e_l^* we will assume they don't. This level is depicted on the graph. Note if this level $\hat{e}_h$ were set lower the lows would gladly deviate because they would get a higher payoff if they could be thought a high at that level. And they strictly would be worse off at a higher level. So we know this is the lowest level $\hat{e}_h$ such that lows will never choose to pick it even if they were thought a high. But the question remains whether the highs want this level or would be fine picking the lows level of eduction and being thought a low, ie. do the highs have an incentive to attain this signal?

Note the high's payoff (below) can be written such that we know they will choose this level:

$$\begin{aligned}
Y_h(\hat{e}_h) - C_h(\hat{e}_h) &= Y_h(\hat{e}_h) - C_L(\hat{e}_h) - (C_h(\hat{e}_h) - C_l(\hat{e}_h)) \\
&> Y_h(\hat{e}_h) - C_L(\hat{e}_h) - (C_h(e_l^*) - C_l(e_l^*)) \\
&= Y_l(e_l^*) - C_l(e_l^*) - (C_h(e_l^*) - C_l(e_l^*)) \\
&= Y_l(e_l^*) - C_h(e_l^*)
\end{aligned}$$

The first equality adds and subtracts $C_L(\hat{e}_h)$. The second line switches out the second term with e_l^* s instead of $\hat{e}_l$ s and we know that this difference in the cost curves increases thus we know we go from an equality to an inequality (look at the graph of the cost curves to see this). The third line switches out the first term which we know are equal (based on how we defined $\hat{e}_h$). And the last line simply cancels terms.

So we see the lows have no incentive to attain the high's level of education and the high's do have an incentive to attain this higher level of education. But again note that type is not observed, so even though education does increase human capital and MP (through the $Y(e)$ function) it also acts as a signal. So what about efficiency? Well the way we have drawn this the signal is inefficient - the level of education for the high's is higher than the efficient level derived at the outset. The other possibility is that the cost curve of the lows increases so fast that this level is below e_h^* in which case we would get our efficient levels - but what fun would that be? What we see is that it is possible that education plays both roles, but we could still get inefficiencies.

Other Issues

Returns to education? Regardless of HC or signalling - is it worth it? Empirical question. Issues involve measuring benefits - better working life, teasing out education and induced training effects, longer working life, and reduced unemployment risk (theory of education, training and layoff - $MP > w$). Also issue with costs - what are real costs? Lost time, disutility. And of course there is the issue of measuring the causal effect.

Is there signalling? See paper on GED.

Furthermore maybe we are interested not just in private returns. Consider measuring social returns to

education. And maybe consider ‘spill-over’ effects (ie. peer effects) of education.

[Maybe more here later]

Chapter 6

Labor Demand

Labor is demanded in order to produce something else. Labor is not (in general, of course there are exceptions) hired because it directly increases the firms utility. Labor is hired as an input in production. In this sense labor is a factor of production and the labor market can be thought of as a “factor market”. Note that the firms we are discussing are involved in three markets: the labor market, the capital market, and the goods (output) market. Lets begin with the assumption that all markets are competitive - thus firms are ‘price takers’.

Basic Profit Maximization

Since all markets are competitive we will assume firms can hire an unlimited amount of labor (h) at a market set wage rate (w), they can rent an unlimited amount of capital (k) at rental rate (r), and can produce a single good (q) and sell an unlimited amount of it at the market set price (p). Each firm combines labor and capital to produce the good with the production function:

$$q = f(h, k)$$

where production is increasing in h and k ($df/dh, df/dk > 0$) at a decreasing rate ($d^2f/dh^2, d^2f/dk^2 < 0$). Also assume $d^2f/dhdk > 0$, this implies the derivative of one input increases when the other input increases -

if you are not familiar with second derivatives don't worry, we will discuss their interpretation in the context of our problem. A generic shape of a production function can be seen below to the left. Note also that if we 'slice' a curve out of the function across the z-axis we get a 'level curve' - every point on this curve gives the same amount of output. And not further the relation between h and k on this curve - this is plotted on the right. This is what is an isoquant curve.

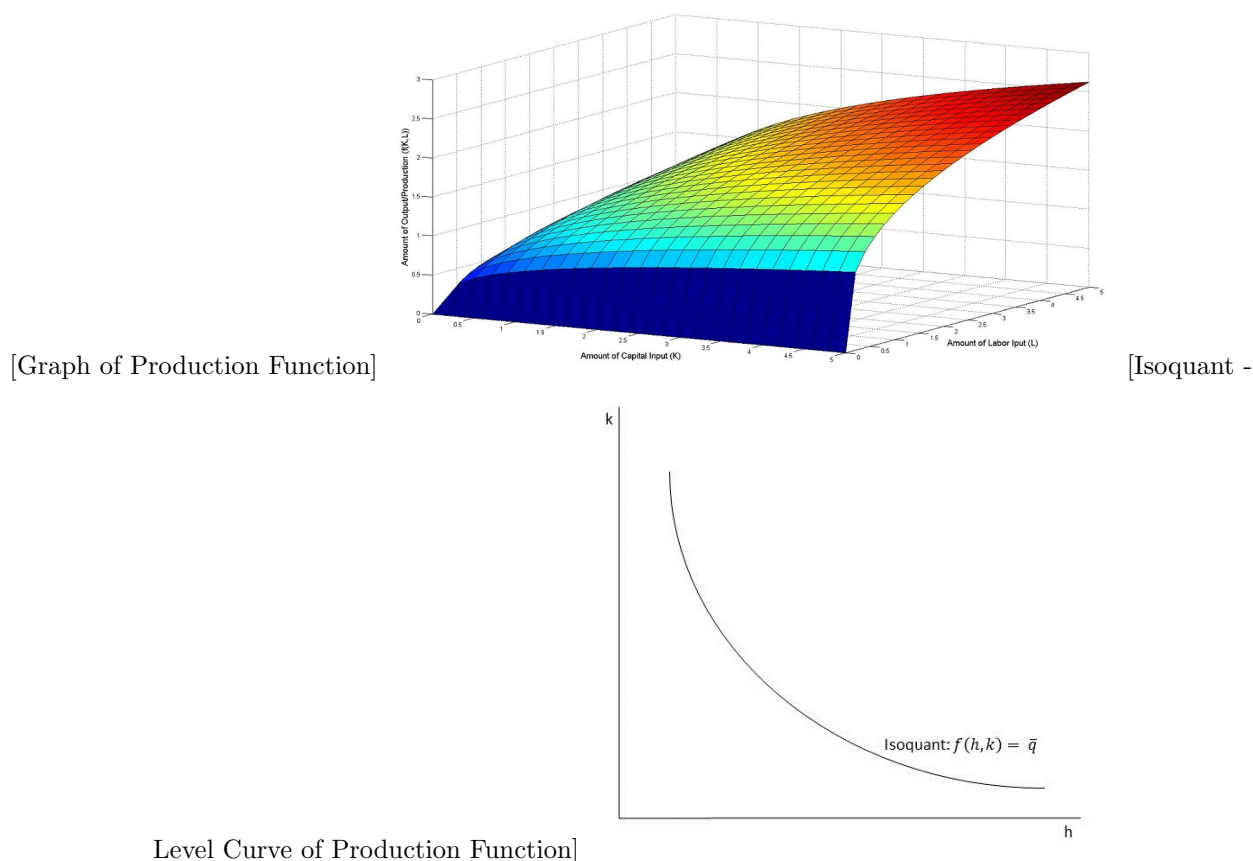


Figure 6.1: A Standard Production Function.

Profit (which the firms maximize) is defined as revenue minus costs:

$$\pi = Rev - Cost = pq - wh - rk$$

We assume firms are profit maximizers and thus solve the problem (noting $q = f(h, k)$)¹:

$$\max_{h, k} pf(h, k) - wh - rk$$

¹Note there might not always be a solution to this problem. For example if there are constant returns to scale then any amount of output maximizes profits (and profits are zero), as long as the FOCs are satisfied. In such a setting we would think about cost minimization instead given some output $\hat{q}$ as we did in the household production topic.

This yields two first order conditions (optimizing with respect to choice of h and k):

$$p \frac{\partial f(h, k)}{\partial h} = w$$

$$p \frac{\partial f(h, k)}{\partial k} = r$$

These results (along with the production function) implicitly define the optimal choices of h^* and k^* and also the optimal amount of production q^* . Before looking at these results graphically lets get some definitions out of the way so we can discuss the interpretation of these results.

$$MP_h = \frac{\partial f(h, k)}{\partial h} \text{ Marginal Product of Labor}$$

$$MP_k = \frac{\partial f(h, k)}{\partial k} \text{ Marginal Product of Capital}$$

$$MRTS = \frac{MP_h}{MP_k} \text{ Marginal Rate of Technical Substitution}$$

$$MRP_h = pMP_h \text{ Marginal Revenue Product of Labor}$$

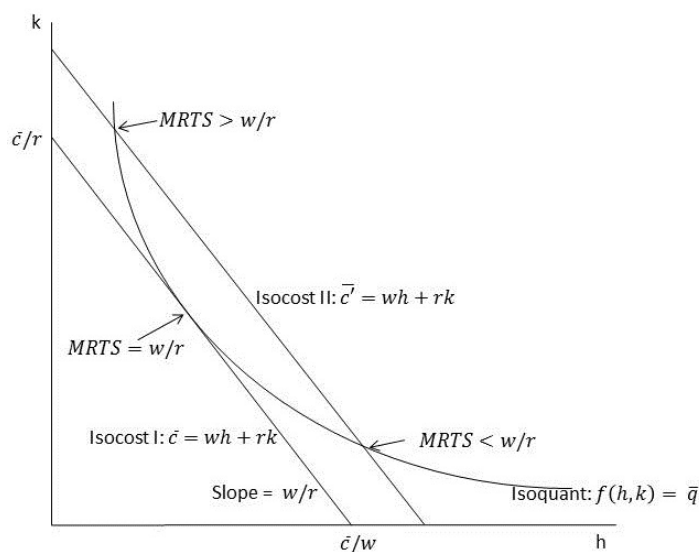
$$MRP_k = pMP_k \text{ Marginal Revenue Product of Capital}$$

So what do our first order conditions mean? Well note that our first condition can be rewritten as $MRP_h = w$ and the second as $MRP_k = r$. So at our optimal points the gains to the firm from using one more unit of labor (capital) - which is the interpretation of MRP_h (MRP_k) - should be equal to the cost of using one more unit - w (r). Note if the $MRP_h > w$ then the value of hiring more workers (or more hours) is greater than the wage rate. Thus it makes sense to hire more. Of course if you exceed that point then you should reduce your labor demand.

Lets look at this graphically. To do this lets first introduce an isocost curve (similar to a budget constraint). Note that the prices are fixed by the market and there is a trade off determined by the prices. Say there is some cost level $\hat{c}$, there are multiple bundles of labor and capital that cost this same amount - the isocost curve. So the slope of the isoquant should equal the ratio of the prices in the market.

Also note that $MRTS = \frac{w}{r}$. What can we make of this? Recall, similar to the MRS in the consumer case, the MRTS is the trade off between capital and labor that allows the firm to make the same output - it is

Figure 6.2: Profit Max/Cost Min.



the slope of the isoquant - $MRTS = \frac{MP_h}{MP_k} = \frac{dk}{dh}$. Think about why this makes sense. So the firm give up one unit of labor, then they need $\frac{dk}{dh}$ units of capital to be able to make the same output. Now if they give up that unit of labor they have w extra cash which can buy $\frac{w}{r}$ units of capital. So as long as $\frac{w}{r} = \frac{dk}{dh}$ you can buy just enough to produce what you were. If you are were using more capital than this point your MRTS would be greater than the ratio of prices and you would be at a higher cost level if producing the same amount. At this point the market is willing to give you more labor for capital (the price ratio) than you need - why not take that deal and lower your costs. Similar thought process should explain why you would not want to be at a point where MRTS is lower than the price ration.

Scale vs. Substitution Effect

Lets briefly consider what happens to the optimal amount of h demanded by the firm when its cost w changes (assuming all other costs remain constant). We can decompose this into two effects: a scale effect and a substitution effect.

Scale effect: when the price of labor goes down the average costs of production go down as well. This induces firms to produce more and in the process employ more labor. To think about this just go back to your intro micro course and recall what shifts the supply curve of a firm - costs of production. And what happens when

price is fixed but supply curve shifts out - quantity increases.

Substitution effect: if wage rate decreases it is now cheaper relative to capital then it was. Hence the firm has an incentive to utilize more labor and less capital - this brings the equality of $\frac{w}{r} = MRTS$ back into line.

Labor Demand Function

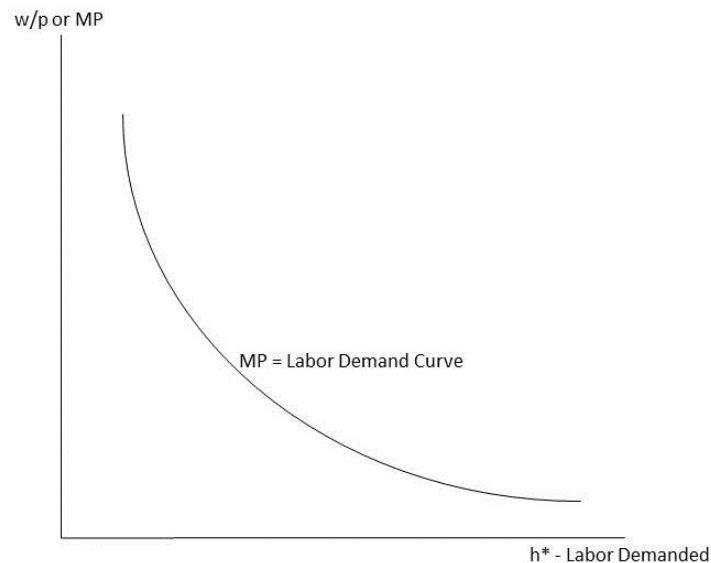
Lets first assume we are in the short run - the short run is defined by the time frame in which capital is fixed (of course this is a strong assumption but an ok starting point). Thus all that can vary when prices vary is labor. So in a sense we have a one factor production function $f(h, \bar{k}) = g(h)$ because now k is just a constant. Now using our same understanding of prices and marginal products we arrive at the condition:

$$MP_h(h) = \frac{w}{p} \quad (6.1)$$

The above notation makes explicit that the marginal product now only depends on the value of h . So if we also take p as fixed we see we have an implicit function linking wages and labor demanded - and this of course is the definition of a demand curve. Moreover we see the demand curve in terms of ‘real wages’ is simply the marginal product curve. And since we have assumed d^2f/dh^2 is negative, then this is a downward sloping curve (recall the second derivative describes how the first derivative changes with changes in the variable). So we see labor demanded is at the intersection of the real wage and the marginal product of labor. Furthermore, rather than graphing the demand as MP in relation to the real wage, we could graph it as MRP in relation to the nominal wage. If we graph it this way then the area under the MRP (demand) curve is equal to the firms revenue.

In the long run, even if p is fixed, the demand for labor depends on the wage rate and on the price of capital (thus we should rethink the scale and substitution effects discussed above but now in light of changes in the price of capital - I will leave you to this). Nonetheless it will be decreasing in the wage rate. This brings up the point that labor demand depends not just on the price of labor (and of course here we have no room in our model for regulation which certainly plays a role). This relation between price of an input and quantity demanded is directly related to the idea of elasticity. We will not cover the details here (see any introductory microeconomic text for the basics) but in short elasticity of demand is a measure of the responsiveness of

Figure 6.3: Short Run Labor Demand.



the quantity demanded to a price - and this price can be the price of that good (own price elasticity) or the price of other goods (cross price elasticity).

Another important caveat is that not all labor is the same, though this should be obvious. Moreover any given firm likely employs various ‘types’ of labor. For simplicity let's divide them between skilled and unskilled labor. The main interest in this divide is how the two types of labor interact differently with capital - specifically since capital more and more is in the form of high end technology. There are no more typewriters and storage cabinets but rather word processors and SQL databases - this evolution in the type of capital used for ‘similar’ work has changed the nature of the relation between capital and skilled/unskilled labor. Even though we will not cover in detail elasticity there are some good rules of thumb. These are called the *Hicks-Marshall Laws of derived Demand*. They tell us that, on average, own-wage elasticity will be high under these conditions:

- When Price elasticity of product being produced is high: a pure scale effect
- When factors can be easily substituted in production: a substitution effect
- When the cost of labor is a large share of the total cost: this leads to big scale effect
- When supply of other factors of production is highly elastic: thus they are easy to substitute

Much has been written - journalistic and academic in nature - about the interplay between technological gains and the demand for skilled vs. unskilled labor. All we will note here is some generalities. It seems to be that skilled labor is more likely to be a complement to capital (or if they are substitutes skilled labor is not so easily substituted) whereas unskilled labor is likely a substitute for capital. Of course we have not exactly said what skilled or unskilled labor is, nor what capital it is to which we are referring. Nonetheless I think the general idea is clear. The following are measures (cite needed) of elasticities between physical capital, and two types of labor - number of workers and human capital per worker (this second is a measure of 'skill'). Note that when price of number of workers increase firms demand less of them and more physical capital as well as more human capital. The key thing is physical capital and human capital move together whereas number of workers moves in the opposite direction.

Government Regulations and Labor Demand

There are several prominent types of government intervention that will likely affect labor demand and wages in a labor market and we will discuss three: minimum wages, employment taxes/mandated benefits, and employment protection laws.

4.1 Minimum Wage

Minimum wage is an example of a price floor - it is illegal to buy/sell a good/service below a certain price. Minimum wage came into effect in Korea in 1988 at a rate of 462 Won per hour. In 2014 it was set to 5,210 Won per hour. In 2012 it covered about 14% of workers (thus is likely to increase slightly with the new increase). Lets begin by looking at the effect of a minimum wage in the most basic supply-demand setting and then discuss some shortcomings to that model, some caveats, and the empirical literature.

The standard introductory microeconomic analysis of a minimum wage is depicted below. Note that a minimum wage is meaningless if it is below the market wage rate so it is only interesting if it is binding. So for who is it likely to bind? Most likely this will affect the market for unskilled and youth labor. Minimum wages tend not to be high enough to affect the market for nuclear engineers. The minimum wage covers relatively few workers, and has never affected more than (...)% of the population and currently affects about (...)%.

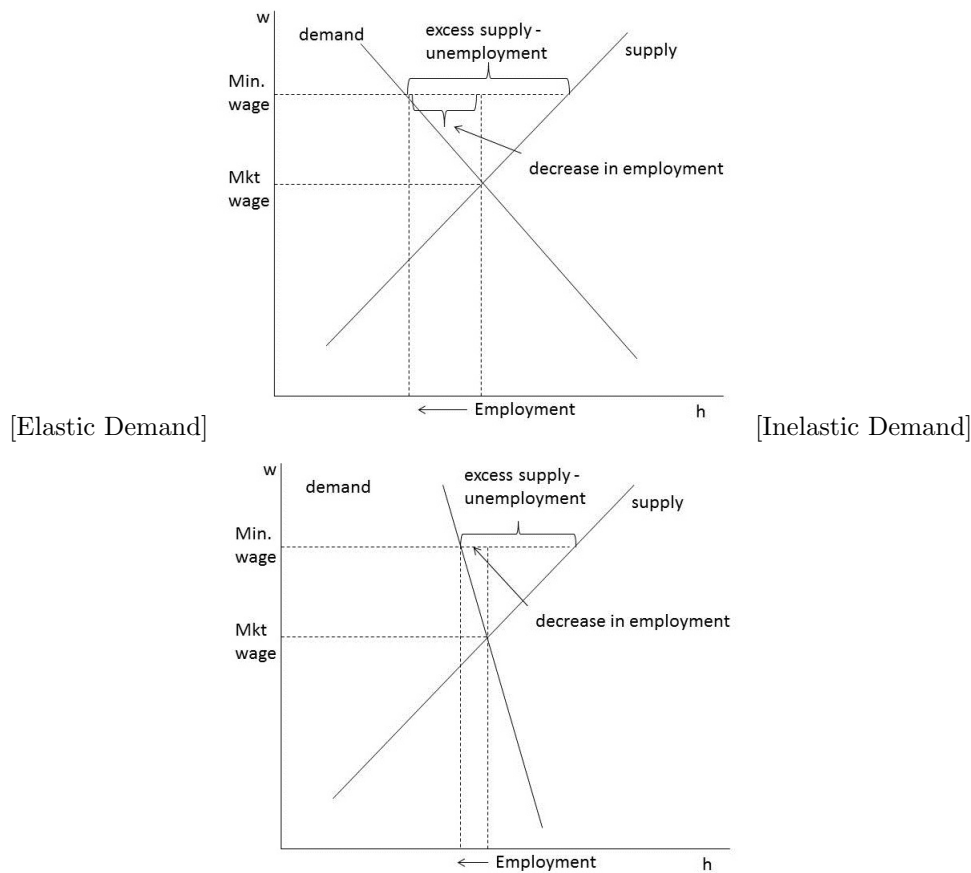


Figure 6.4: Effect of Minimum Wage.

One key thing to note is there is a difference in ‘unemployment’ - the difference between supply and demand - and the change in employment due to the minimum wage. The former includes the fact that labor supply will increase in response to a higher wage. However what is likely of most interest is the change in employment. Furthermore note that the unemployment amount changes depending on how elastic the labor demand is. If it is not very elastic (does not respond much to change in price as in Panel (b)) then unemployment is smaller. Moreover this difference is fully due to the change in employment rate. The elasticity of supply will also affect the unemployment rate, but only through the amount of labor offered to the market and so has no effect on the change in employment. So we see the effects of a minimum wage depend crucially on how this model is drawn - ie. what is the elasticity of demand? Well the empirical literature is not exactly settled. Though the measured effect (of *small* changes in minimum wage) seems to be small or negligible. However there is an important issue with the empirical literature on the effects of minimum wages - they are all short term measures. In a dynamic world it is very hard to measure what we would like to know - how does a minimum wage change employment levels in the long run?

In interesting finding by (Card and Kreuger 1996?)² was that the introduction of a minimum wage might actually increase employment. This, at first, seems perhaps a nonsensical finding. However this result would perhaps be expected if employers exert market force and act as ‘monopsonists’ - the only buyers in the market. Lets see what that might look like. To make this as simple as possible for us to understand, assume there is a market with some supply curve relating labor supply and wages: $LS(w) = w \rightarrow h = w$. Furthermore assume this market could be competitive or covered by a single firm with a $MP_l = 10 - h$ (thus a production function of $f(h) = 10h - 1/2h^2$), and the price of the good produced is ones: $p = 1$. It is fairly easy to solve for the competitive equilibrium of $w = h = 5$ from where the supply and demand curves cross. Furthermore if the wage were set at the equilibrium level just found and this firm took it as a given and maximized profits it too would choose to employ $h = 5$ from solving:

$$\max_h f(h) - w \cdot h \quad (6.2)$$

The key assumption here is that the firm cannot affect the wage - it is set at $w = 5$ which came from where supply and demand crossed (thus its marginal cost of labor is $MC_l = 5$). Now say the firm realizes that it is the only firm in the market and does not need to accept w as given as equal to 5. That is say it knows that $w(h) = h$ - its chosen level of employment determines the required wage rate (thus its cost of labor is $w(h) \cdot h = h^2$ and its marginal cost of labor is $MC_l = 2h$. Now the picture changes.

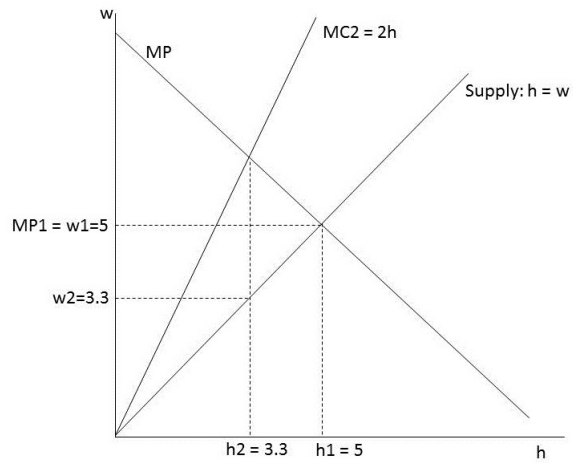
The firm now solves:

$$\max_h f(h) - w(h) \cdot h \quad (6.3)$$

which given our assumptions leads to $h = 3.33$ and $w = 3.33$. Both of these results are shown graphically below in the left panel. The key thing here is that employment is lower than would be in a competitive market, and wages are lower (in particular lower than the MP of labor). The firm employs less now because they affect the wages. The supply curve at 4 units of labor is $w = 4$. But the marginal cost of labor to the firm of employing that 4th person is not 4. To employ the 4th person they must pay them $w = 4$ in order to induce them to work, but must now also pay the 1st, 2nd, and 3rd workers an extra dollar/won assuming they cannot price discriminate (the supply curve at 3 workers is 3). Thus the marginal cost of the 4th worker is actually higher³.

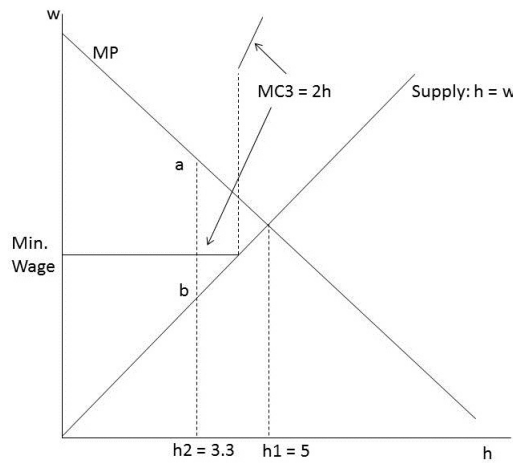
²Brief note on the paper maybe, if not done in class.

³If you add this up you get 7 which is not $2h$. This difference is from treating the supply curve as continuous (as in the equation) verses 1, 2, 3 ... as in the text explanation. Check this yourself and make sure you understand it is just a mathematical issue and not an economic one.



[Equilibrium and Monopsony Outcome]

[Monopsony with



Minimum Wage]

Figure 6.5: Effect of Minimum Wage in Monopsony Case.

Now in the right panel we see how a minimum wage might change the marginal cost of labor to the firm. Point 'b' is the level of wages the monopsony chose and 'a' is the wage at which the original marginal cost curve passed through the marginal product curve. If we set a minimum wage in between these two values we will see employment increase. So increasing the wage rate in the market through a minimum wage increases employment. This is because of the minimum wages effect on the firms MC_l curve. The line between a-b is the amount of labor the monopsony employed (h_2), setting a minimum wage between a-b reduces the marginal cost of the firm for some area above h_2 . Specifically the area between h_2 and where the minimum wage hits the supply curve - then it jumps to its original level. Of course maximum employment is reached if the minimum wage is set to the natural equilibrium wage rate.

Is this likely - the monopsony setting? It could be argued that since youths and the least skilled are the ones likely affected by a minimum wage are the same ones least likely to move for jobs and be limited in their

options regarding employers. This would likely lead to possibly enough job search frictions to cause such an outcome. On the other hand it should be noted that for the most part the results of the Card and Krueger paper were shown to be due to data issues, and once they are corrected their results disappear (cite).

Even if a minimum wage reduces employment perhaps a country wants it anyway though. As stated in the introduction of this text the results of positive economic analysis is not the end of policy debates. Even if there is a small reduction in employment maybe many feel it is *worth it* to increase the minimum wage. This is beyond this text, and in a sense beyond economics. Our duties, as economists, are to spell out the trade-offs, not dictate to voters which trade-offs they should make. With this in mind though some have advocated simply greatly increasing the Earned Income Tax Credit as a more efficient way to help the poor working class. This has some merit as it helps those with low wages while avoiding the possible employment effects of a minimum wage. Others might retort, however, that this puts the burden on taxpayers rather than the producers/consumers of goods that employ those workers. Again, this gets a bit beyond what I am willing to discuss here, but it is nonetheless worth noting.

4.2 Taxes and Mandated Benefits

Many governments have employment taxes - a tax a firm pays based on their payments to workers. In addition many countries have some mandated benefits such as insurance paid by the firm on behalf of the worker. But from the firm's point of view the distinction between tax and mandated benefit is largely arbitrary - either one is an additional payment that must be made in addition to wages. The only difference comes from the worker's side. Benefits are likely valued, and perhaps as much as wages. This distinction will be important.

Let's focus on mandated benefits as most of you have likely already seen the basic derivation of taxes and their effects on markets⁴. The difference is how much employees value the benefits that are being mandated. Below in the left panel we see a case where employees do not value the benefit at all. In this case it is exactly like a tax on labor - the cost is split and employment drops. In the right panel employees 'fully value' the benefit. In other words they would have purchased it themselves anyway. In this setting the supply curve shifts out because now in addition to the wage rate there is also some benefit. In this case the employees

⁴In brief, it does not matter who 'pays' the tax, rather the *tax incidence* is split between the supply (worker) and demand (firm) sides of the market, and the split depends on the elasticity of supply/demand with the least elastic paying more (they are less responsive so more cost can be pushed onto them). And of course the market shrinks causing a deadweight loss.

pay the full amount and there is no change in employment. Of course the likely outcome is somewhere in the middle where employees ‘partially value’ the benefit - possibly because of uncertainty or incomplete information.

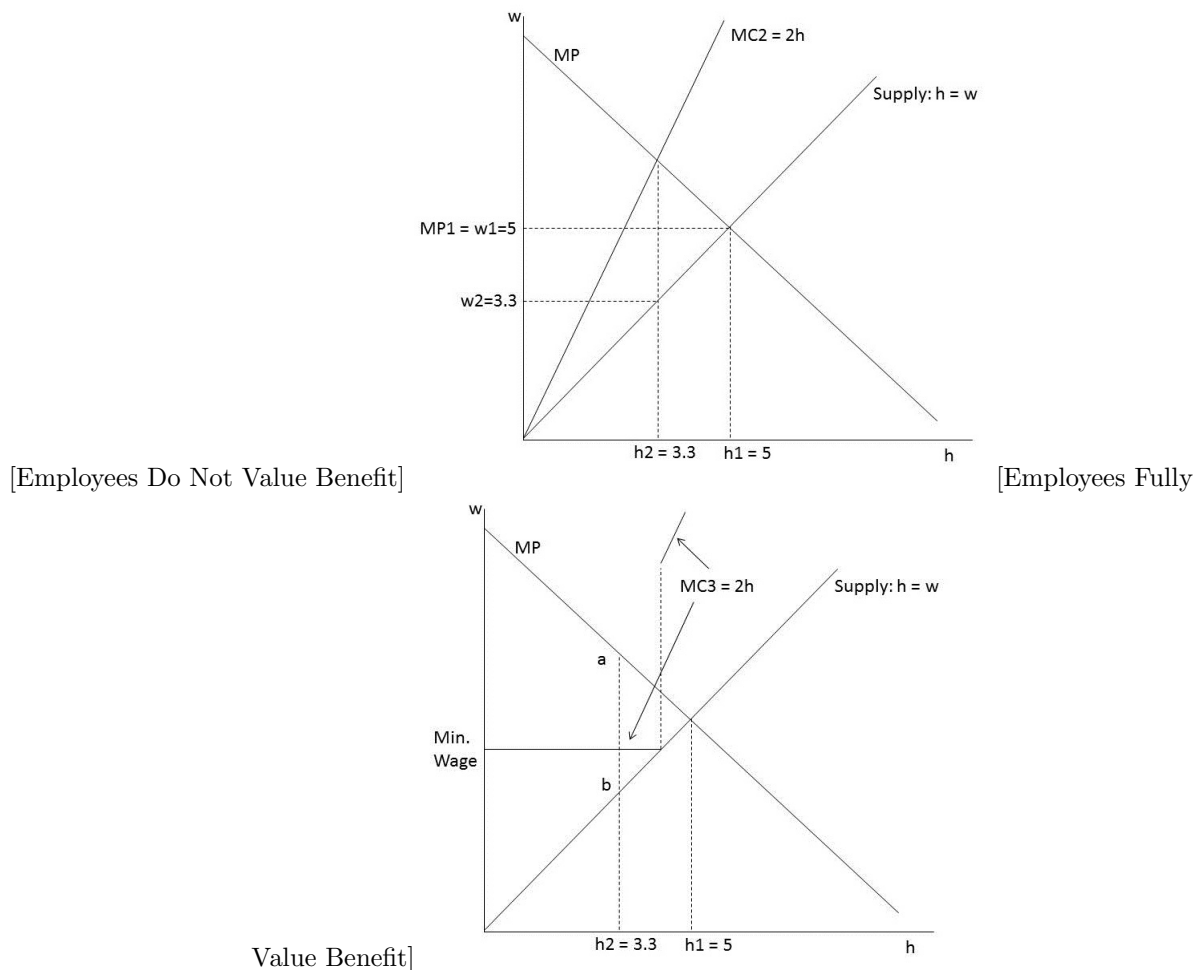


Figure 6.6: Effect of Mandated Benefit.

4.3 Employment Protection Laws

Many governments have some set of laws known as ‘Employment Protection’ laws. These laws make it difficult for firms to fire employees (at least once an employee is determined to be a ‘permanent’ worker - they do nothing in general for non permanent workers). They range from outright restrictions to specified level of severance payments for termination or possibly that the firm ‘prove’ to the government that some amount of terminations are needed in the face of current macroeconomic realities. In a perfect model of competition

with no frictions any form of intervention will lead to either lower employment or wages⁵. But of course we do not live in the textbook model economy and so discussion of labor policies must be more nuanced. The literature on the effects of such EPLs is a bit mixed and depends on what is being measured.

Nonetheless the existence, and gaps in such laws, are thought to be large drivers of the dual labor market split between ‘regular’ and ‘non-regular’ workers. These numbers have steadily risen in last decades and Korea stands as one of the highest in terms of employment which is ‘non-permanent’ (note the exact definitions differ slightly in the two figures below).

	2001	2002	2003 ^a	2004 ^a	2005 ^a	2006 ^a
Workers with a fixed-term contract	Less than or equal to 1 month	5.6	5.2	6.7	5.6	5.2
	More than 1 month to less than 1 year	2.8	2.7	4.9	4.7	4.9
	Exactly 1 year	1.5	1.9	3.3	4.4	5.3
	More than 1 year to less than 3 years	0.6	0.6	1.3	1.7	1.7
	3 years or more	0.5	0.6	0.7	0.8	0.9
	Subtotal	11.0	10.9	17.1	18.2	17.7
Workers without a fixed-term contract, where employment is not expected to continue due to involuntary reasons	2.9	3.8	4.3	7.6	5.9	5.9
Temporary agency workers	1.0	0.7	0.7	0.8	0.8	0.9
On-call workers	2.2	2.9	4.2	4.6	4.8	4.3
Total^c	16.6	18.1	25.9	29.7	29.4	28.8

Memorandum item
Total dependent employment (000s) 13 659 14 181 14 402 14 894 15 185 15 531

... Data not available.

a) There have been several changes in the survey from 2003 on requiring caution in interpreting trends. The entire sample of households for the Economically Active Population Survey was replaced in January 2003, and the question structure in the supplementary survey was partially modified.

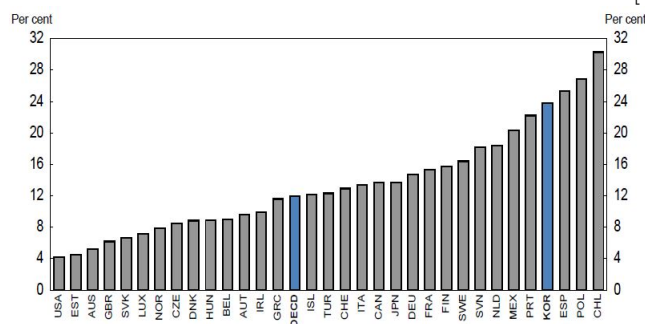
b) Since 2002 the category *Workers without a fixed-term contract, where employment is not expected to continue due to involuntary reasons* includes a subcategory of workers whose contract is being renewed on a regular basis, but who assess the durability of their employment to be limited. The figures in parenthesis exclude this subcategory.

c) Total is adjusted for overlapping categories.

Source: EAPS Supplement, various years, as reported by the Korean Ministry of Labor; and OECD Labour Force Statistics database.

[Rise in Temporary Workers]

[International]



¹ Temporary employees are defined as wage and salary workers whose job has a pre-determined termination date. For Korea, it includes only firms with a fixed-term contract, temporary agency workers and on-call workers (excluding double-counting).

Comparison - 2011] Source: OECD Employment Outlook Database.

Figure 6.7: Non-Regular Workers.

These laws that protect regular workers do not cover certain contract/temporary workers and thus many firms have turned to them for flexibility in their labor decisions. Furthermore they are less represented in unions (5% vs 17% of regular workers) and more importantly due to their non-permanent nature they have shorter tenures with employers and thus receive less training - both of which (union status, tenure, received training) lead to higher wages.

⁵Many point to the U.S.'s typical fast recovery to its flexible labor market (ie. low labor protection) - though this notion has been questioned in light of the latest atypical employment recovery post 2008.

Figure 6.8: Reasons for Using Non-Regular Workers.

	Reduce labour costs	Increase employment flexibility	Perform peripheral tasks	Perform short-term tasks	Other reasons	Total
All industries	32.1	30.3	18.5	13.9	5.2	100.0
Manufacturing	28.7	34.5	17.9	14.7	4.1	100.0
Non-manufacturing	35.4	26.1	19.1	13.2	6.2	100.0
By firm size						
Less than 30	35.5	28.9	15.8	13.2	6.6	100.0
30-99 workers	28.5	27.6	18.7	18.2	7.0	100.0
100-299 workers	37.7	26.2	15.5	14.3	6.3	100.0
300-499 workers	34.3	29.4	19.6	12.7	3.9	100.0
More than 500	28.1	49.9	22.9	9.6	1.6	100.0

Source: OECD (2007b).

Figure 6.9: Pay Differential.

	Total	Firm size ^c		Union membership	
		Small firms	Large firms	No	Yes
Regular workers	9 263 (100.0%)	7 994 (86.3%)	14 046 (151.6%)	8 679 (93.7%)	12 351 (133.3%)
Non-regular workers	6 256 (67.5%)	6 032 (65.1%)	10 572 (114.1%)	6 382 (68.9%)	9 513 (102.7%)

a) Hourly wages are calculated as average monthly total wage / (average weekly working hours x 4.3 weeks), using data from the EAPS Supplement, August 2005.

b) Figures in brackets denote the ratio of wages in a particular category to average wages of regular workers.

c) *Small firms* are firms with less than 300 workers, and *large firms* are those with 300 workers and over.

Source: Kang (2005).

Chapter 7

Compensated Wage Differentials

In all previous chapters we basically only discussed wages, ignoring the fact that things other than wages matter when it comes to labor supply/demand decisions. Here we will attempt a first try at incorporating these important aspects. First let's get out in the open a key aspect that is problematic with incorporating 'non-pecuniary' (ie. not money) aspects to labor decisions. That problem is that different people value different things. Some people love the outdoors while others love air-conditioning and central heating. In such settings it becomes very difficult to build simple models where people value other aspects of their job. However if we can find some aspect that everyone can agree on we might be able to get some traction. A good starting place is risk of death or injury. I believe you would be hard pressed to find people who would like higher risk of injury at their job for the same wage. So we will base our models on just that - wages and risk of death.

Basics

The basic model assumes that people don't want to die and so must be compensated for a higher risk of death. Furthermore reducing this risk in the workplace is costly to the firm. Moreover we assume the marginal costs of reducing risk are increasing. This should make sense. The idea is that it should be easy to reduce the most obvious hazards but reducing risk from 0.01% to 0.005% become very costly. Also assume there is competition and so firms are pushed to zero profits. In what follows we are implicitly assuming the

workers can freely move between firms.

Lets first think about the firm and their trade-off. Basically they are trading off risk-reduction and offered wages, and this trade off will reflect in essence some technology used by the firm. Lets start off at a point where the firm does nothing to mitigate risks - there is some offered wage which yields them zero profits. Now as they reduce risk they must reduce the wage offer or else they will have negative profits. Moreover, as they reduce the risk more, the wage reduction steadily becomes larger due to the increase in marginal cost of risk reduction. Thus we get an offer curve as depicted below on the left. Also there are likely many firms with different risk reduction technologies - the figure depicts multiple firms and their offer curves - these offer curves are *isoprofit* curves - every point has the same profit level (zero). Then in the right panel we have a single 'market offer curve' based on the outline the combination of firms' offer curves make.

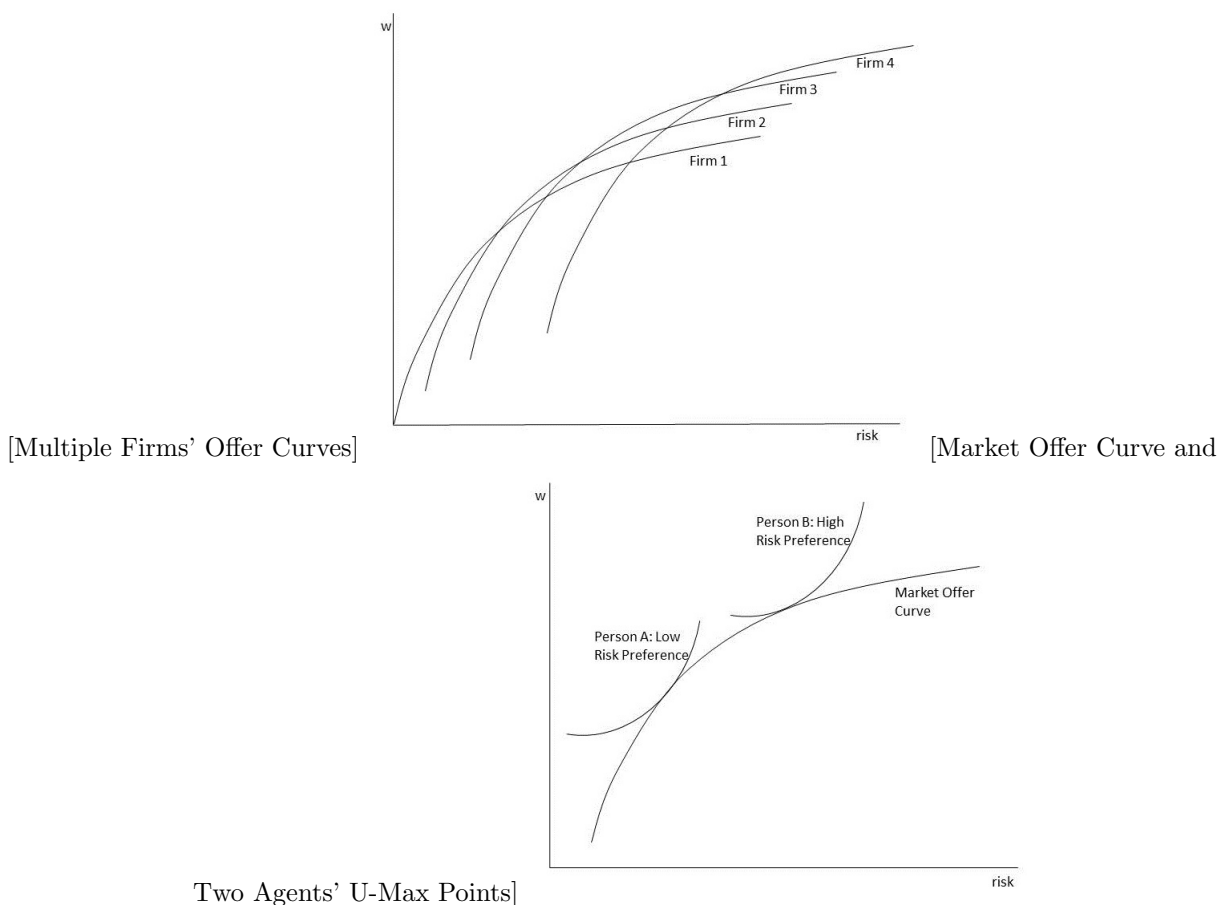


Figure 7.1: Compensated Wage Differentials.

Now since we have already seen workers' decisions regarding wages and leisure and we understand how they are shaped, it is not much to get the workers preferences into this picture. The workers have some utility

function based on wages and safety (the opposite of risk): $u(w, s)$. And note that higher wages and lower risk (higher safety) yields higher utility. In the panel on the left we have two different workers' indifference curves and their optimal choice of wage/risk - one places a lower weight on risk than the other. The same concept of tangency of indifference curve and trade-off presented by the market hold here. Say there is a function defining the market offer curve relating safety to wages: $s = g(w)$. Then the workers problem is:

$$\max_w u(w, g(w)) \quad (7.1)$$

Which yields the FOC:

$$\begin{aligned} \frac{\partial u}{\partial w} + \frac{\partial u}{\partial s} \frac{dg}{dw} &= 0 \\ \rightarrow MU_w / MU_s &= -\frac{g'}{w} \end{aligned}$$

Thus the worker matches their tradeoff between wages and safety to that offered by the market. And we see here that workers with different preferences sort into firms with different costs of risk reduction. The worker who doesn't mind the risk works for a firm that finds it more expensive to reduce risk - and in return for more risk he gets higher wages. The worker who does not want risk sorts to the firm who finds it relatively cheap to reduce risk - and in return for less risk gives up more on wages.

Application to Regulation

So we have seen that there is some explanation why wages might vary depending on certain risks and costs of reducing them. So how does potential government regulation come into play here? For example given some industry - maybe chemical manufacturing or nuclear power production - how does government regulation mandating some type of safety precautions affect the non-regulatory outcome? A key issue is whether workers truly understand their risks. So let's begin with the setting where workers completely understand the risks of certain chemicals and possible exposure to them.

The government might mandate some upper limit to the risk allowed in a workplace. This is likely implicitly done by mandating certain safety equipment etc. Call this limit r_L . In the left panel below we have a worker choosing a risk level r_1 associated with a wage w_1 . Once the regulation is implemented the workers utility is

maximized at the corner of the 'legal offer curve' and they receive wage w_2 and risk level r_L . Now it should be obvious this individual is 'hurt' by the regulation in the sense that their utility is now lower, but how much are they hurt? Can we perhaps put a number on how much they are hurt - maybe in terms of money? Note at their original utility level they would have been willing to accept risk level r_L and wage w' (though it is not offered by the market). After the regulation they receive r_L but only a wage rate of w_2 . So in a sense we can quantify their damages by $w' - w_2$. Note that their losses are not $w_1 - w_2$, because they do value the reduction in risk - just not as much as they had to give up.

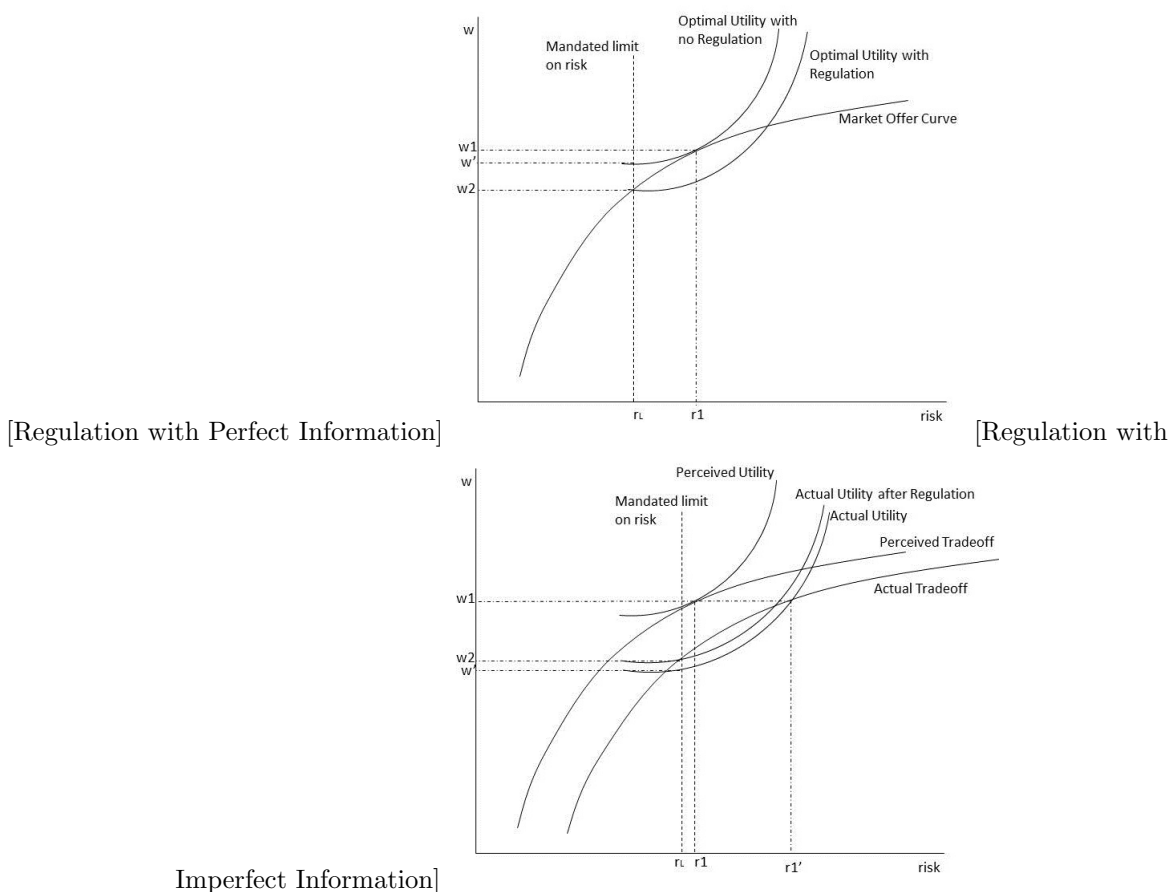


Figure 7.2: Effect of Regulation on Risk Level.

Now let's turn to the case where a worker does not have accurate information regarding the risk levels they face (though the government and firm are aware) - this is depicted in the right panel. Prior to regulation the worker chooses the firm with wage w_1 and perceived risk r_1 . However the actual risks are higher at r_1' , thus they are actually on a lower indifference curve. Now the regulation is again set to r_L . If the worker knew their true risk level they would be willing to take a pay cut to w' for this risk reduction (it is on their actual indifference curve). However they need only take a pay cut to w_2 . Thus in this setting the regulation actually made the worker better off (though of course they don't know this). Of course if the regulation goes

too far it can still hurt the worker.

Application to Pricing

This idea of compensating wage differentials can be used to price things that do not have actual observed prices in the market. For example how much do people value their lives? Now there might be a ‘gut reaction’ against such statements - a human life is priceless many would respond. Nonetheless we routinely make choices that put our lives at greater risk in return for other things we value - people drive fast in the rain for example to save time (or for that matter people drive at all rather than staying home in bed where it is safe). And moreover if we are to recommend policies that save lives but cost money - money that could be spent on other things people value like nice parks - then we need a number.

Recall the slope of peoples indifference curves is how much they are trading off risk and wages - thus it is how much they value safety (perhaps in terms of probability of death). Of course different people are optimizing with different slopes, but if in reality the firms are not that different, and the slope of the offer curve is not too much curved then comparing the wage/risk level between two groups of workers is a kind of ‘average value of safety’ (in the sense that the slope of the line connecting to optimal wage/risk points is a sort of average of slopes of indifference curves). Say the low risk workers receive 50,000,000 Won per year with risk of death of 0.1%, and the high risk workers receive 55,000,000 Won per year with risk of death at 0.2%. Thus we could say (on average) people are willing to take an additional 5,000,000 Won to have an increase of 0.1% risk of death. Well $5,000,000/0.001 = 5,000,000,000$. Thus we would say people value their lives at about 5 billion Won. So in making public policy decisions we would value projects that save a life a year at 5 billion Won a year.

Now in reality we do not have two points with nice data as the above example. In reality these are measured with ‘hedonic price models’ and the above number is called the ‘statistical value of life’. Furthermore this idea extends to settings outside of labor economics. For example maybe we want to find out how much people value good schools or nice parks when choosing housing. Well we would gather data on all the houses sold in the last year - their sale price, size, age, options, and local school quality. When we run the regression the number for the coefficient on school tells us how much people value having nice schools nearby because it tells us how much more they were willing to pay for it when they purchased the house.

Chapter 8

Discrimination

Women tend to earn less money than men in most countries (though the difference varies widely). Blacks tend to earn less money than whites in the U.S. What should we make of this? We have seen that things such as non-pecuniary benefits may be offset with lower wages (compensating wage differentials) and differing levels of human capital investment (schooling) can also lead to different wages. Perhaps these groups simply have chosen certain schooling levels and job characteristics that lead to lower wages or perhaps these groups are simply being discriminated against (though even if the former is happening it could be an optimal choice in the face of discrimination). Before we discuss how to possibly measure such cases let's first look at some theories of discrimination. By discrimination we mean when a group with identical characteristics is treated differently solely because they belong to that group.

Some Data for Korea

The main issue for Korea in terms of labor market discrimination is based on gender. Korea has some of the largest gender gaps in terms of labor market participation, wages/earnings, and the percentage of women in executive positions. In this vein the Korean government enacted an anti-discrimination law in 1988 and then in 2006 an affirmative action law (though whether these laws had any effects remains questionable). Figure 8.1 shows Korea's relative standing in regards to median earnings pay gap. We can see it is far above most nations at nearly 40%.

Figure 8.1: Gender Pay Gap in Median Earnings.

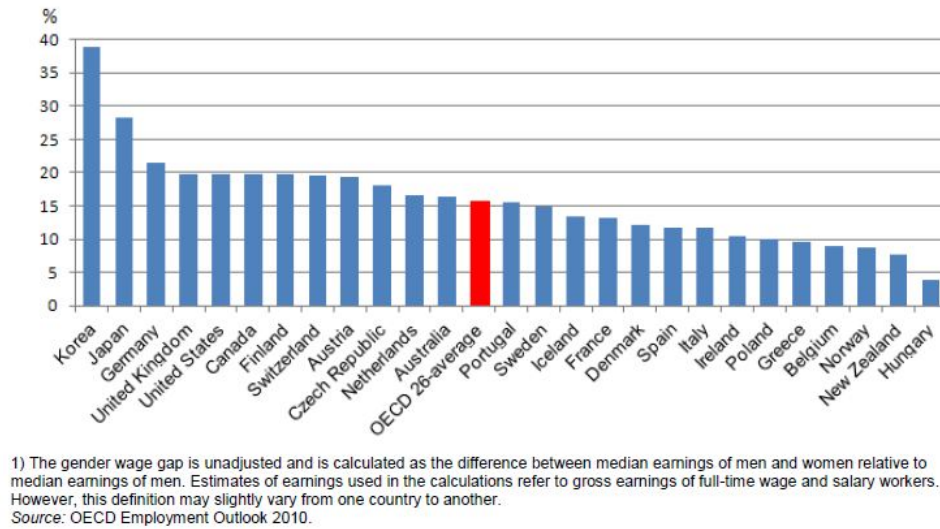
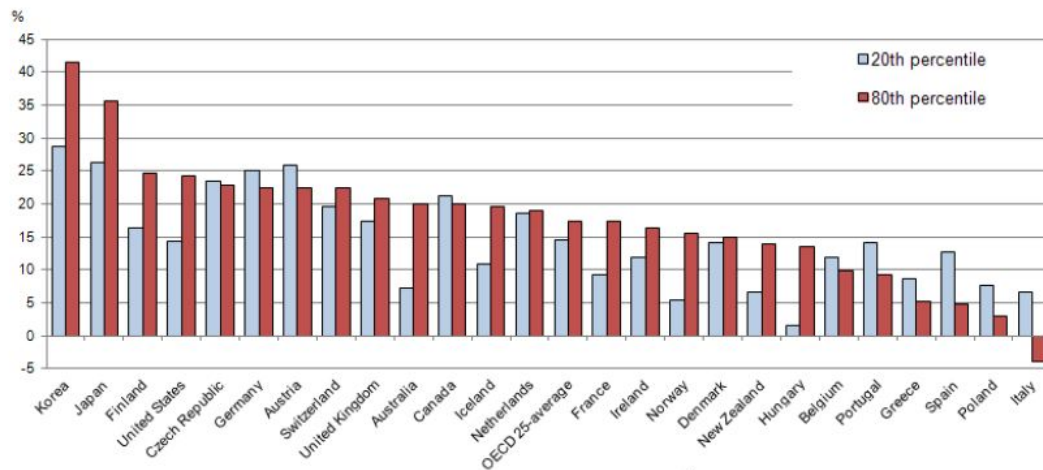


Figure 8.2 breaks this down to differences within the highest earners (80th percentile) and lowest earners (20th percentile). And here we can see Korea remains the highest in this breakdown with much larger gaps for the highest earners (this is the difference in the average earnings for the highest paid women and the highest paid men).

Figure 8.2: Gender Pay Gap by Percentile.



Countries are ranked in decreasing order of the gender wage gap for top earnings (80th percentile).
1) The gender wage gap is unadjusted and is calculated as the difference between top/bottom earnings of men and women relative to top/bottom earnings of men. 2) Data refer to 2005 for the Netherlands and to 2007 for Belgium and France.
Source: OECD Employment Database, March 2010.

Figure 8.3 breaks this down differently by education/age levels. These numbers are the percentage of male earnings received by women. We can see that the gap is smallest, with women aged 35-44 with a college

degree receiving 84% the earnings of men. Moreover this is the highest figure listed. This appears to contradict the previous figures. However it should be kept in mind that the labor force participation rate for educated women drops significantly in their 30's - from about 70% to about 57%.

Figure 8.3: Women to Men Earnings Percentage by Education.

		Below Upper Secondary		Upper secondary and post-secondary non- tertiary education		Tertiary		All levels of education	
		35-44	55-64	35-44	55-64	35-44	55-64	35-44	55-64
Australia	2005	88	99	87	77	80	76	88	84
Austria	2008	71	69	76	84	73	67	73	75
Belgium	2006	73	67	74	83	82	74	83	76
Canada	2007	68	63	70	73	76	59	75	63
Czech Republic	2008	72	81	73	85	67	78	66	75
Denmark	2006	68	57	83	87	72	76	78	78
Estonia	2008	63	66	61	72	64	74	68	76
Finland	2007	78	77	76	77	72	72	77	74
France	2006	76	63	78	82	81	55	84	65
Germany	2008	69	70	86	66	76	68	79	67
Greece	2006	61	45	78	67	68	89	77	60
Hungary	2008	83	86	86	105	57	75	79	86
Iceland	2006	67	90	67	69	58	70	68	74
Israel	2008	69	71	74	70	64	67	70	70
Italy	2006	71	83	81	84	52	45	77	76
Korea	2007	66	67	58	74	84	58	59	57
Luxembourg	2006	85	55	76	78	73	78	79	71
Netherlands	2006	76	77	83	74	79	65	85	74
New Zealand	2008	78	67	76	73	74	76	77	74
Norway	2007	74	78	72	74	68	69	74	73
Poland	2006	65	62	67	91	66	73	77	83
Portugal	2006	66	58	75	74	70	72	77	65
Slovak Republic	2008	71	74	72	83	61	79	68	80
Slovenia	2006	85	85	87	97	81	99	92	104
Spain	2007	72	74	85	86	82	75	86	79
Sweden	2006	94	82	77	80	72	77	78	83
United Kingdom	2008	82	78	69	72	77	78	76	77
United States	2008	67	65	69	75	68	62	71	65
OECD Average		74	72	76	79	71	71	77	73

Source: OECD Education at a Glance, 2010

What implications do these trends have? Many point to the low birth rate in Korea as a result of such findings - though this is speculative. Regardless of what the outcomes might be of this wage gap, simply having a wage gap does not itself mean there is discrimination. We will look at ways to measure (or attempt to measure) if wage gaps are truly a result of discrimination. But first let's look at some basic theories of discrimination.

Theory

There are two main theories regarding discrimination - one coined taste based discrimination which is due to Becker, and the other coined statistical discrimination which is due to Arrow. Let's begin with the former.

2.1 Taste Based Discrimination

Taste based discrimination is based on the idea the firms simply would rather not hire a specific group - they have some distaste for working with/employing them. Thus the worker must in a sense compensate the firm by accepting a lower wage. Call A the undiscriminated group and B the discriminated group. The firm maximizes its utility based on its profits and some monetary value it places on employing certain groups. Let d be the “coefficient of discrimination” - it is the monetary value of the distaste of employing the discriminated group ($d \geq 0$). Further more there is a production function that produces a good q , sold at price p by using both groups’ labor $N_b + N_a$. Thus the firm maximizes:

$$U = pf(N_B + N_a) - w_a N_a - w_b N_b - d N_b \quad (8.1)$$

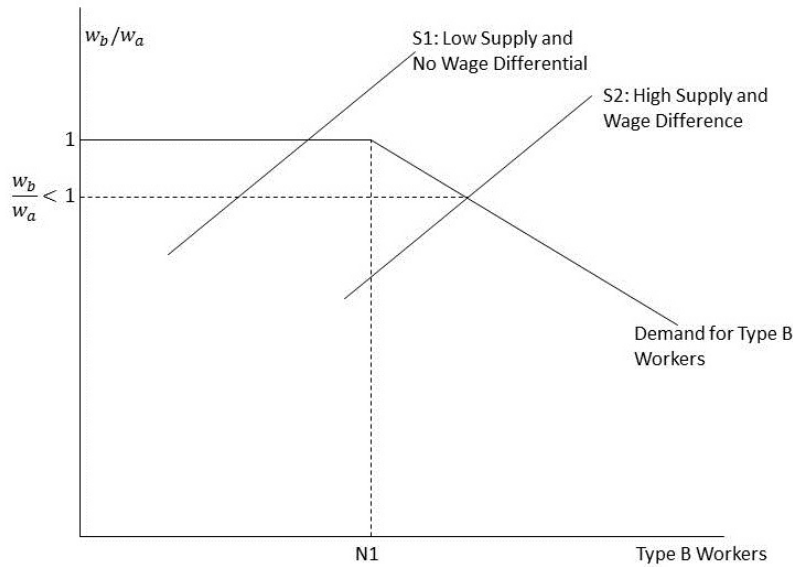
So we can see the firm acts as if the wage for group B is $w_b + d$. So the firm will only hire group B if $w_a - w_b \geq d$. Now we can easily solve for the firms maximized choices of N_a and N_b and note they are defined by the FOCs:

$$pf'(N_a) = w_a \quad (8.2)$$

$$pf'(N_b) = w_b + d \quad (8.3)$$

So if workers are the same (ie. have the same marginal product) this firm will only hire group B at a lower wage. However, the question we are asking though involves an entire market and what really matters the the relative size of the supply of type B workers and discriminating firms and how much they discriminate. So assume there is some distribution of d ’s throughout the firms in the market (with some possibly $d = 0$). Furthermore say that there is some set of firms that do not discriminate and will hire up to N_1 amount of type B workers at equal wages ($w_b/w_a = 1$). After that point we can graph the demand for type B workers by introducing the discriminating firms (with the low d firms first). If the supply of type B workers is small enough that these firms can employ them all then there will be no wage differential in the market. This is seen in the figure below. However if the supply of type B is large enough to surpass this amount then some will have to work for discriminating firms. And since there is free movement of workers this will drive the overage market wage of these workers down.

Figure 8.4: Taste Based Discrimination and Wage Differential.



Note that this implies that as type B workers increase the wage gap also increases, whereas as the number of nondiscriminating firms increase the wage gap decreases. The main critique to this explanation is that if there is perfect competition and free entry of firms then the discriminating firms will be competed away. To see this note that discriminating firms are maximizing utility and not profits - they are not hiring type B workers until wages equal MP. Thus their profits must be less because in a sense they are buying or paying for their discrimination by not using type B workers who's MP is higher than their wage. And if the market is competitive and profits are driven to zero then these firms must be making losses. Of course this is a simple static model of the labor market, and search based models of based discrimination easily show wage gaps can persist in competitive markets if there are search costs.

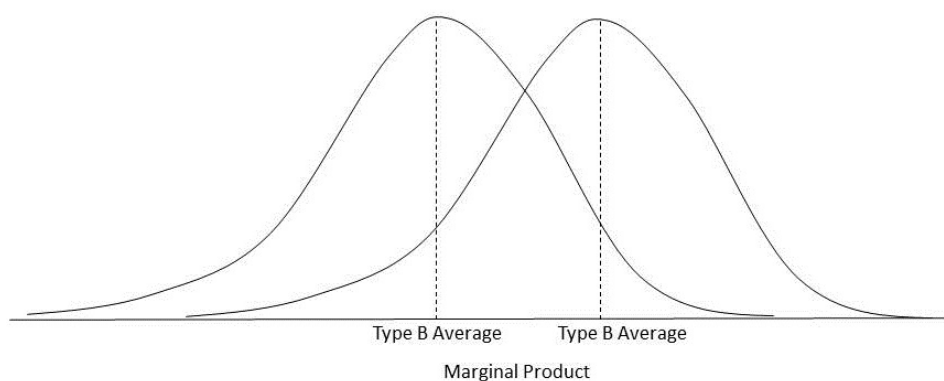
There might also be co-worker taste based discrimination. In this case one type of worker does not want to work with the other. This is normally modeled as $u = w(1 - d)$ where d is their discrimination coefficient. In such a setting firms will completely segregate, since to match the utility of working in a firm without the discriminated type of workers a firm would have to pay a premium to the discriminating worker. So firms completely segregate and there is no wage gap. If for some reason this is not possible then some workers receive a compensating wage differential to make up for the disutility of worker with the other type. On the other hand discrimination might come directly from the customer. Since the customer is the source they simply pay for their discrimination and the discriminated workers MP is lower and thus also their wages. There is no reason why this should be driven out of the market.

2.2 Statistical Discrimination

There are two types of statistical discrimination models we will look at. In one there are distributions of two types of workers' MP that may differ, but firms know the groups average and only see the individuals MP with 'noise'. In the second the two groups are the same. However firms observe 'noisy' signals of individuals MP based on some test score (related to investment in human capital) and that level of noise is higher for one group. We will see this latter one implies some interesting outcomes that make it complicated to measure 'discrimination' in a market.

The first model is fairly simple. There are two types of workers again A and B . But not all workers are the same in their MP for any group, rather there is a distribution of MPs. However type A workers have a higher average MP (AMP) than type B workers. This is depicted below.

Figure 8.5: Distribution of MPs.



Each firm sees some individuals MP but with 'noise' and know their groups average MP. Now we will not formalize this model, but will sketch the idea. This is the intuition: if a firm can not see your MP for sure, but only noisily see it, they will put some weight on your observed MP and some on your groups average¹. So for example the firm pays wages equal to expected MP and the expected marginal product is an average of your MP and your groups average: $w = E(MP) = 1/2MP + 1/2AMP$. So we can easily see that if you look at two people with identical MPs, one type B and one type A , that the type A worker will receive a higher wage. The simplest version of this is if there is no signal from the individual but only from the group average. In such a case everyone receives their group average MP.

¹The actual choice of the firm is formalized differently, but the math is a bit beyond this course and the idea sketched here captures the main idea nonetheless.

Now in the traditional sense of discrimination this does not really fit. The firm is not paying type B with an identical MP less because they do not like them, but because of incomplete information. So though any particular type B worker is discriminated against in the sense that they receive lower wages than an identical type A worker, on average there is no discrimination because average wages equals group average MP. Moreover, ‘low’ type workers are helped and ‘high’ types are hurt for both groups. Why might this happen? Perhaps firms know that, on average, schools in predominantly black neighborhoods are worse than those in white neighborhoods (why this is is a different, though related question). Then if they see two individuals, one white and one black, both with a high school diploma and some perceived level of ability, they have a reason to believe there is a likelihood the white worker may be better anyway.

A second model is slightly different, and in a sense more interesting. Say the two groups are identical in the distribution of their MP. And again firms see ‘noisy’ signals of their MP - perhaps a test score or GPA or just schooling S (which is a result of costly investment in human capital) - and so again firms put some weight on the signal and some on the known (identical) group average. However here the noise in the signal is smaller for type A workers. So type A workers receive wage (for example): $w_a = E(MP) = 1/2S + 1/2AMP$ and type B workers receive: $w_b = E(MP) = 1/4S + 3/4AMP^2$. Note here that type B workers with above average MP receive lower wages than identical type A workers (though the opposite is true for below average workers).

But, more importantly, what is the marginal gain to investing in schooling? Well it is lower for type B workers because the signal is worth less to them. So if we recall the chapter on human capital, if they are rational and equate marginal costs and benefits then type B workers will invest less in human capital. What is the result? They receive lower wages (on average) because they invest less in human capital and above average workers receive lower rewards for schooling. So if we look at the data it may seem the wage gap is yes partly discrimination (same S receive lower wages) but partly a result of their choice to invest less. However, they invest less precisely because they are responding to the market - so the lower S is actually a response to the ‘discrimination’. Of course measuring this becomes tricky.

²Now again we will not formalize this but rather explain the intuition, to formalize it we would need to actually establish an equilibrium. Thus in what we do here really ‘glosses over’ the inner workings of the model, but again the main idea holds.

Decompositions

So how do we go about measuring and putting an number on discrimination in labor markets? That is, how do we put a value on wage differences? One approach is through a wage ‘decomposition’. The goal is to decompose the observed difference in average wages into a part that can be ‘explained’ and a part that cannot - the part that cannot is loosely described as the differential associated with discrimination. The following is termed the ‘Oaxaca-Blinder’ decomposition after the economists who it is attributed to. Note that the average wage can be described as a function of the characteristics of the population - this is just a regression of wages on things like age, education, IQ, tenure, job type etc. So think about a type ‘A’ and type ‘B’ worker. They may have different average characteristics which we denote $\bar{X}_A$ and $\bar{X}_B$ and possibly different ‘returns’ to those characteristics which we denote β_a and β_b . Then we can write out their group average wages as:

$$\begin{aligned}\bar{w}_a &= \beta_a \bar{X}_a \\ \bar{w}_b &= \beta_b \bar{X}_b\end{aligned}$$

Then the difference in average wages is:

$$\bar{w}_a - \bar{w}_b = \beta_a \bar{X}_a - \beta_b \bar{X}_b \quad (8.4)$$

Now subtract and add $\beta_a \bar{X}_b$ and arrange and you get:

$$\bar{w}_a - \bar{w}_b = \beta_a (\bar{X}_a - \bar{X}_b) + (\beta_a - \beta_b) \bar{X}_b \quad (8.5)$$

The first part is due to differences in characteristics while the second is due to different returns to those characteristics. Two problems with this is that, one (as we just saw) choices of human capital are endogenous and chosen in light of the returns. Thus the choice of possibly lower returns due to discrimination will mean we underestimate it. On the other hand, we can only do this with things we have data on. If other things matter and are not picked up in the data we may wrongly assign them to discrimination and so overestimate it. Furthermore there are always two decompositions (do you see this?) though most papers report both.

Chapter 9

Search Theory and Unemployment

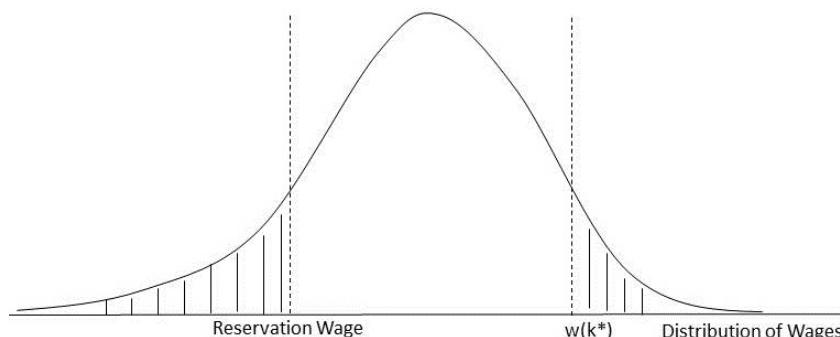
Up until now we have talked about labor demand and supply and the market as if it magically clears. In reality there is a process where one goes from unemployed to employed. This process is modeled with what are known as ‘search’ or ‘matching’ models. The basic idea is that a worker must actually look for a job and that take time and resources. Again here, as in some of our previous models, the math is a bit beyond this course so we will not formalize the model but rather explain the main idea, its intuition, and some of its applications.

Basics of Search

The basic idea regarding search is that there is some distribution of possible jobs. They have associated wages, some skill requirements (k), and possibly other characteristics. But let's start with the basic model where all jobs are the same except the wage offered and skill needed (actually the simplest does not have a skill set but this leads to interesting predictions). Furthermore the agent ‘searches’ and gets one wage offer (which has a skill requirement - skills needed increase with wages) at a time (from some known distribution) and must decide to accept and stay employed for the rest of their lives or search some more. The workers utility is based solely on consumption: $u = u(c)$ and they care about their discounted lifetime utility and have some skill level k^* (thus they will not be employable at jobs with skills higher than k^*). Assume if they are unemployed they receive some consumption b possibly from some government program and there

is possibly some cost to searching c . What is the choice this worker must make? Well they need to assume when to accept a wage/job offer, and since nothing changes except the wage this comes down to choosing a *reservation wage* $w_r(b, c)$. This is a wage, above which they will accept any offer. This is depicted below.

Figure 9.1: Distribution of Wages and Reservation Wage.



Once they accept a wage they earn that the rest of their lives, so waiting for a higher wage means once they get it they have higher wages for life. But while they hold out for that higher wage they must wait longer and lose out on the declined wages. The wage that just balances these two effects is the reservation wage. The probability they will accept a job is equal to the area not shaded above (the shaded area either they reject or are not qualified for). The expected eventual wage is the average wage in the unshaded area.

There are several key implications of this model that, though intuitive, could not be captured in previous models easily. One is that there will be unemployment - the probability of accepting (or even getting an offer they are qualified for) is not 1 for any period. And the expected duration of unemployment is (inversely) related to the probability of accepting a wage. Second, individuals will not use their full skill set - since they will not set $w_r = w(K^*)$ they will be 'underemployed'. Third, identical individuals, though will have identical reservation wages, will have different wages accepted.

Also we see there is room for some government policy here. For example we see that if the benefits they receive while unemployed are larger the individual will set a higher reservation wage (because the cost associated with rejecting any individual wage offer would be smaller). This will have two effects. One - they will be unemployed for longer. Two - they will finally accept a 'better match'. That is they are more likely to use their full skill set. And since we have seen human capital investment is a valuable thing this is a positive. Thus we see unemployment insurance is essentially a give and take between a bad incentive and a good one. Also, though it is not explicitly in the model, what if we could somehow make the individual

receive more offers? For example maybe a jobs program could be thought of as increasing the number of offers received. This would imply both better matches and shorter unemployment lengths.

Types of Unemployment

We generally divide unemployment up into specific types: frictional, structural, and cyclical (and to a lesser extent seasonal unemployment).

Frictional unemployment is the type of unemployment that naturally arises because it actually takes time to find jobs etc. This is the type of unemployment that the above search models essentially model - and this is the type that is considered the 'natural rate' of unemployment. Somewhat related is seasonal unemployment. This is the natural changes in employment levels due to the seasons. For example during big holidays retail sales pick up and more customer service workers are hired, or notably there is much less construction in the winter.

Cyclical unemployment is normally the type of unemployment that policy makers are concerned about. This comes about because of demand-deficiencies - fluctuations in the 'business cycle'. This is essentially unemployment due to recessions and the type that concerns macroeconomics. The big question then is, in the face of a recession why do wages simply not fall instead of employment? The leading theory is 'wage rigidity'. This boils down to the fact that firms generally do not slash wages. Why is this? Some point to union power (though this cannot explain but perhaps a small percent), or cultural norms, or perhaps most likely that firms know that slashing workers would significantly hurt moral in the firm and laying off a few workers is more profitable. Or a reason could be related to human capital - if they just cut wages then the best workers are likely to leave or the ones with the most experience. Perhaps best to simply let go of the least productive.

Structural unemployment might be a new concept to some. This is the unemployment that is due to inherit mismatches between the labor demand and supply. Of course if wages were completely flexible and occupational and geographic mobility had very low costs then this type of unemployment would not persist - but these are not likely assumptions. One example of the geographic aspect can be seen when looking at labor mobility in the South-West of the U.S. after the Cold War died down. There were many firms invested in military contracts in Southern California. When these contracts disappeared after the

war died down unemployment shot up. Then over time people migrated to nearby states and the cross-state statistics became more homogenous. Related, some people point to the housing slump in the U.S. post 2008 as a reason unemployment remained high. Because house prices dropped it made it harder for people to relocate to regions with lower unemployment levels. Though this seems likely the net effect on the unemployment level is unclear (one paper puts it at only about 0.5% points - “Geographic Reallocation and Unemployment During the Great Recession: The Role of the Housing Bust” by Karahan and Rhee). Perhaps more importantly many of the lost jobs were manufacturing/construction whereas the ‘jobs of the future’ are of a different kind and these workers do not have these skills. Another possible source is *efficiency wages*.

Efficiency Wages and The Wage Curve

The concept of efficiency wages is related to contracts and revolves around the idea that firms cannot at all times monitor their workers. So how do you get your workers to work efficiently (hard) for you? Well if you pay a premium, then it is worth it for them to work harder, because if they don’t and you fire them then they will have to work at the lower market wage rate. Of course this higher efficiency wage must mean that there will be equilibrium unemployment. So let’s look at a simplified model of efficiency wages called the “Shapiro-Stiglitz” Model (this version leaves a lot out but the full model is beyond this course).

Say a worker can choose to work hard or be lazy and ‘shirk’. Now if they shirk there is a chance they are caught and fired (probability = p) and then have to take the market wage rate w , but they only get this with some probability $1 - u$ - related to the unemployment rate. But putting in effort is hard and there is a cost to it. Now say a firm pays an efficiency wage w^* . If the worker works hard they receive:

$$u(e = 1) = w^* - c \tag{9.1}$$

whereas if they shirk they receive:

$$u(e = 0) = (1 - p)w^* + p(1 - u)w \tag{9.2}$$

So in order to ensure the workers don't shirk the firm must pay:

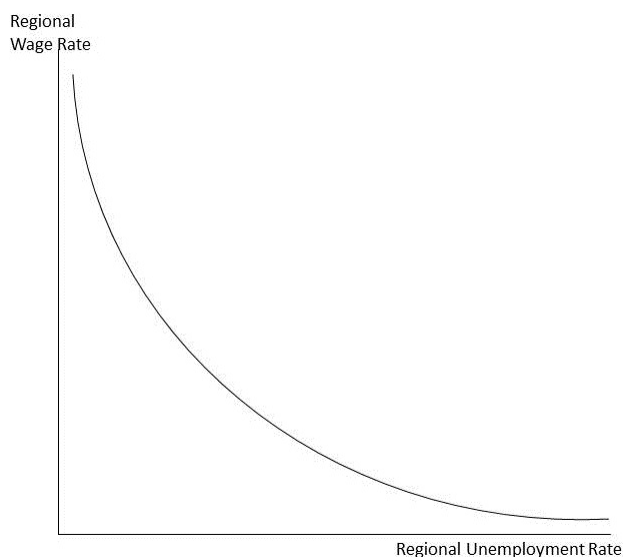
$$w^* > (1 - u)w + c/p \quad (9.3)$$

Now the firms will pay this if the benefit to production (the increase in output under non-shirking) outweighs the added cost. But at this higher wage level there will be lower employment rates.

Now we will not look at the full model, but it also equates economy wide unemployment rates and an equilibrium condition (note we just assumed an unemployment rate without forcing this to be endogenous to the model). Notably we find that the higher the unemployment rate the lower the efficiency wage rate need be. This is because being fired is in a sense 'punishment' for shirking, and the punishment is worse if unemployment is high. And if the punishment is high then the incentive to work hard does not need to be high.

This last point is a good explanation for a strange finding. In almost all countries surveyed when we look between regions (after controlling for important things like human capital) there is a strong negative association between the regional wage rate and the unemployment level. That is, regions within countries with low unemployment levels have high wages, while regions with high unemployment have low wages. This is called 'The Wage Curve' - and it is shaped very similarly for all countries.

Figure 9.2: The Wage Curve.



Why is this 'strange'? Well in a basic S&D model if wages are above equilibrium then there is more

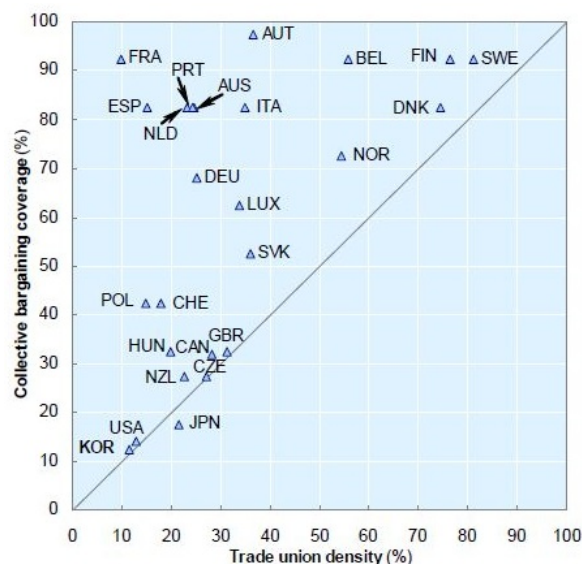
unemployment and the higher the wage above the equilibrium the higher the unemployment. So high wages would be associated with high unemployment. So this goes against the standard model. Furthermore, though the basic model predicts that high unemployment levels will lead to wages *falling*, the wage curve depicts wage levels not changes over time. The efficiency wage model gives a rational to this finding.

Chapter 10

Unions

About 10% of Korean workers are members of unions. This number peaked at about 20% in 1989. This 10% is low by international standards (see below) as is the percentage of workers in some other form of collective bargaining. Nonetheless the power of unions in certain key sectors is quite strong (see the 2011 Hyundai unions wage hike agreement) though their strength in general has been steadily eroding with both political push-back by successive governments (see for instance recent KORAIL strick issues) and changing perceptions.

Figure 10.1: Union Coverage.



Unions represent workers in negotiations with firms. Their demands may be about decision making etc, but we will focus on the monetary demands. In other words - a union works to increase wages for its members. Of course there is a trade off the union faces. Recall the firm has a downward sloping demand for labor reflecting the diminishing product of labor. And also recall a key issue is how easy is it to substitute the labor of a specific firm for other labor or capital - that is the elasticity of the demand is a key issue that determines the trade-off the union faces. Of course in a growing industry/country this trade-off could be mitigated (in a sense) by growth in the demand of the product leading to growth (shifting out) of the demand for labor. Though of course the trade-off is still there, just sort of 'covered up'. We will look at two simple models of union effects on employment/wages in the firm. One is a simple 'monopoly'-union model, the other is a more efficient 'bargaining' model.

Monopoly-Union Model

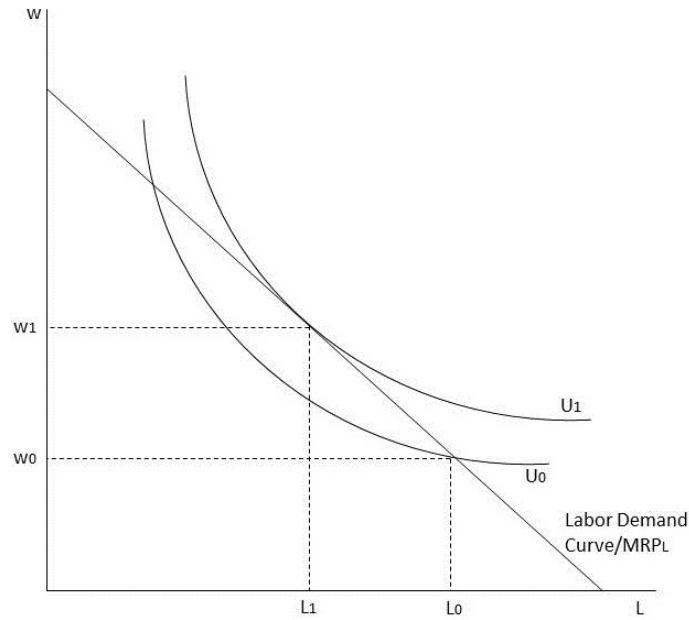
First lets assume the union has some utility function that it attempts to maximize. Specifically it cares about the employment level and the wage: $u = u(l, w)$. And assume this utility function has standard properties. Also assume there is some 'pre-union' market wage rate w_0 . Assume the union can simply set the wage level and then the firm takes this union wage as given and chooses the employment level as determined by the labor demand curve (marginal product of labor curve). It is fairly easy to see that the union will choose wage w_1 and then the firm will employ level L_1 in Figure 10.2.

So the union maximizes its utility (however that is determined) subject to a constraint which is the marginal revenue product curve (demand curve) of labor for the firm. Now one thing we will see is that this is not efficient (in terms of union/firm). Lets look more closely at this with a bargaining-contract model of unions.

Efficient Contracts

First note that by efficient we only mean for the union and firm, as opposed to the general definition which includes all of society. First lets think about what the labor demand curve is - it is the choice of employment levels that maximize profits when the firm faces a given wage rate. So say at w_0 the labor demand curve gives us L_0 , thus $MPL_0 = w_0$. So what happens if at that rate the firm employed less than L_0 ? Well

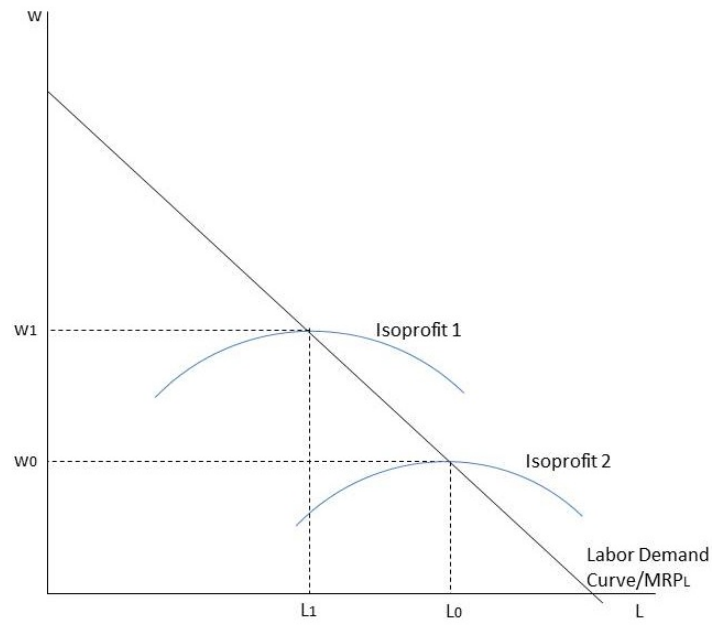
Figure 10.2: Monopoly-Union Wage Setting.



it would, in a sense, be ‘leaving some revenue on the table’. In other words, it could be employing more workers whose marginal product is higher than the wage they must pay (recall we are assuming decreasing marginal product). So their revenue and profits must be lower, thus to maintain the same *profit level* they must decrease wages. So what if it employed more workers than L_0 when the wage rate was w_0 ? Well now it is paying workers wages with lower marginal products than their wages, thus hiring them is in a sense ‘eating away at’ their revenue/profits - so again their profits must be lower - and again it must lower offered wages to maintain the same *profit level*. Also note that if all else is fixed then moving down the demand curve represents higher profit levels. Now let's put this together and add some ‘isoprofit’ curves to the firm's labor demand curve in Figure 10.3.

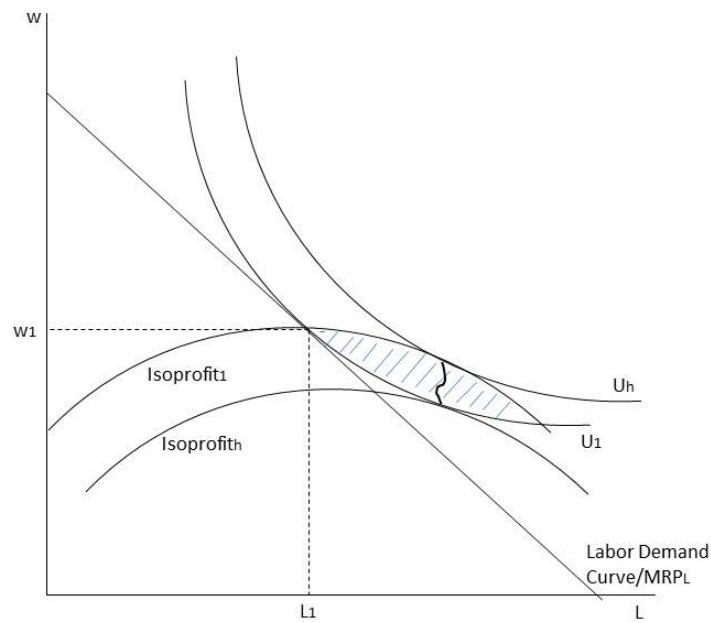
So rather than the union simply setting the wage and then the firm setting the employment level, what if they negotiate and determine both together. This can lead to a better outcome for both firm and union - the shaded area in Figure 10.4. Let's take as our starting point the outcome from the monopoly-union model - (w_1, L_1) . Note at this point the firm is on $Isoprofit_1$ and would be indifferent between any points on it, specifically the union could push its utility out to U_h with the firm maintaining the same profit level as (w_1, L_1) . Conversely note the union could maintain its same utility U_1 as the monopoly-union model and the firm could increase profits by moving to $Isoprofit_h$. So we see that, relative to the monopoly-union outcome, there are gains from trade to be made. In the two extreme examples graphed the union takes all

Figure 10.3: Isoprofits.



the gains in the first and the firm takes all the gain in the second. Furthermore there is a whole range of ‘efficient’ contracts graphed by the line in the shaded area, where the line represents the tangent points between the indifference curves and isoprofits as we move from one extreme to the other.

Figure 10.4: Efficient Contracts.



So while any point in the shaded area is ‘better’ than the point (w_1, L_1) , not all are efficient - only on the

contract curve are they efficient. The shape of this line is determined by how the isoprofits and indifference curves move.

A key thing to note is that regardless of the shape of the contract curve is beyond the labor demand curve. So does this seem realistic? The evidence is not clear cut, but there is some evidence for this when you look at teacher contracts (student/teacher ratio limits), or mandatory crew size for construction unions, or limits on jobs specifications (which means a firm must hire different workers for specific jobs leading to more employment than necessary).

Issues and Evidence

3.1 The Demand Curve

We have seen that the elasticity of the demand curve inherently impinges on a union's power - if it is too elastic then even slight increases in the wage will lead to large reductions in the employment level. This gives reasons unions may try to effect policies or firm actions to lower this.

One example would be directly trying to limit competing goods. If there is less competition for final goods then there is lower price elasticity of those goods and thus lower wage elasticity of labor. Thus, for example, in the U.S. unions have in general opposed free trade agreements and pushed for 'buy American' legislation/campaigns. Similarly unions have supported minimum wage laws and laws requiring those receiving government contracts to pay the 'prevailing wage' - ie. union wage. This acts to increase the price of non union goods thus enhancing the unions pricing power. Similarly as noted above many times union contracts have staffing requirements that dictate specific tasks which can and cannot be done. So perhaps a firm needs to hire a welder and a machine operator when realistically one person could do both jobs.

So how do unions press their power? In the private sector it is generally through a strike or simply a threat of a strike which costs the firm money not only in lost revenue but potentially publicity and legal costs as well. In the public sector, generally, strikes are illegal for many sectors. However there are usually mandated arbitration. Arbitration occurs when the two sides cannot agree. In such a case arbitrators come in and make the decision.

3.2 Effect of Unions

In addition to the employment protection laws in Korea, some site the strength of unions as a reason for the dual labor market. Though evidence of direct causation is hard to find we do know that non-regular workers in general do not have unions working for them so they would fall outside of the union sector. Unions in general lead to higher wages for their members. But what about uncovered firms/sectors - what happens to the workers in these places? There are a few possibilities. One is that the workers who do not get employed in the union sector spill over into the non union sector increasing the labor supply in that sector. This serves to suppress the wages in the non union sector. So the difference we see in wages between the two sectors is a mix of the union power and the spill over.

Alternatively there might be a threat effect. In this case the uncovered firms/sectors are worried the workers might attempt to form/join a union and so increase wages in return for remaining union free. In such a case the wage difference between the two sectors would be smaller than the actual union effect. And there may also be some degree of 'waiting-unemployment' where workers do not take non union jobs because it is more optimal for them to stay unemployed and hope for a union job to open up.