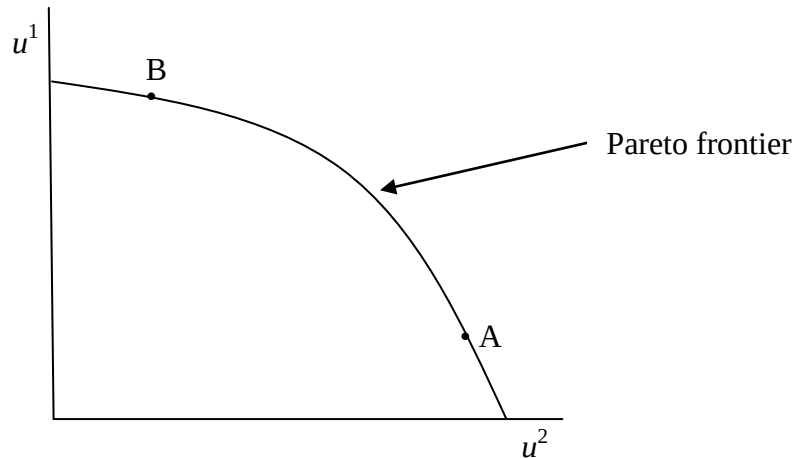


Economics 230a, Fall 2016

Lecture Note 1: Welfare Economics and the Role of Government

Public finance is the positive and normative analysis of government's role in the economy. To understand this role, let us start with the two fundamental theorems of welfare economics.



The first fundamental theorem says that, under certain assumptions, all competitive equilibria are Pareto optimal. That is, they lie on the Pareto frontier that defines the set of possible allocations among individuals; on the frontier it is not possible to make someone better off without making someone else worse off. But Pareto optimality defines optimality in only a limited sense; it does not allow us, for example, to rank outcomes A and B in the figure. To do this, we need some mechanism for ranking allocations.

Measuring Social Welfare

We typically use a social welfare function, $W(u^1, u^2, \dots, u^H)$. We typically assume that $W_h \geq 0$, i.e., that the social welfare function is non-decreasing in individual well-being and therefore achieves a maximum at some Pareto optimum. Note that, because the scale of the utility function representing an individual's preferences is arbitrary, so is the social welfare function. Only the combined effect of an increase in individual h 's income on utility and an increase in individual h 's utility on social welfare, the product $W_h \partial u^h / \partial y^h$, matters. Indeed, all that matters is the relative values of $W_h \partial u^h / \partial y^h$ for different individuals. We typically refer to this product as a welfare weight given to individual h .

Note that in using a social welfare function we must make interpersonal comparisons. This may be straightforward for individuals with identical preferences, as we can normalize welfare weights so that two individuals with same income, who choose the same bundle, have the same welfare weight. But among individuals with different preferences, there is no obvious unique normalization. For example, we can assign the same welfare weights to individuals with the same level of income for a given price vector, but this normalization implies different welfare weights at the same level of income if relative prices change. For example, suppose that person 1 has a stronger preference and hence a larger budget share for good i than person 2. Then an increase in the price of good i makes person 1 worse off relative to person 2. With a larger decline in real income, person 1 should now receive a higher welfare weight than person 2.

Further restrictions on the standard social welfare function include it being based only on individual well-being. This limit is not as restrictive as is sometimes thought. For example, it does not rule out the existence of concerns for the equal or similar treatment of equals, a criterion generally referred to as horizontal equity, if one interprets this criterion, following Auerbach and Hassett, as there being a greater aversion within the social welfare function to differences in outcomes among individuals with observationally similar exogenous characteristics such as age. However, the social welfare function approach does not allow separate weight to be given to specific criteria, such as inequality *per se*, as this could lead to the choice of a Pareto dominated outcome (e.g., by making the most well-off individual worse off to reduce inequality). And, of course, we assume that individuals have stable preferences. Some applications from the perspective of behavioral economics (for example, hyperbolic discounting) reject this assumption, leaving us having to decide which version of an individual's preferences to use in our measure of social welfare.

Finally, social welfare functions relate to *outcomes*, not initial conditions or the process by which the outcomes are reached. This means, in principal, that we should be equally happy with a given outcome regardless of whether it was delivered by a democratic process or a dictator, and regardless of how far the final distribution of well-being is from that in some initial state. There is a long-standing philosophical debate on this point, going back at least to 19th century principles like one calling for "equal sacrifice," and some evidence, as Weinzierl discusses, that individual perceptions of fairness take this form.

With our measure of social welfare specified, we can make use of the second fundamental theorem, which says that each Pareto optimum can be achieved via a competitive equilibrium, if lump-sum taxes and transfers are available to shift individual endowments. For example, if initial endowments yield point A and our social welfare function prefers point B, we can impose a lump-sum tax on individual 2 and give it to individual 1 to induce this shift in the resulting equilibrium.

Based on the fundamental theorems and our measure of social welfare, we have established a role for government, but it is a very limited one: imposing lump-sum taxes and transfers to choose the socially optimal point among Pareto optima. This is quite removed from the government activity we observe. What's missing?

First, government's ability to use lump-sum taxes to improve the distribution of resources is limited. In practice, we observe few if any taxes that are independent of individual choices. This leads to the use of more realistic taxes and transfer payments.

Second, the analysis so far does not take account of market failures, which can result for many reasons. If market failures exist, then a competitive equilibrium will generally not be Pareto optimal, so government intervention in the form of government spending, non-lump-sum taxes, and regulations, may improve outcomes, even in the hypothetical case that lump-sum taxes are available.

Important Market Failures

The two classic types of market failures are public goods and externalities.

“Pure” public goods are defined as having two key characteristics:

1. Nonrival in consumption: $x^1 = x^2 = \dots = x^H = x$.
2. Nonexcludable: no individual can be kept from consuming all of x if it is produced.

Characteristic 1 means that we want everyone to consume the good, because it is costless for them to do so once the good is produced. Characteristic 2 means that private provision, even inefficient provision in which individuals have to pay to access the commodity, is not feasible, since individuals cannot be excluded from consuming and therefore can choose to pay 0.

If both conditions are satisfied, only public provision (or publicly funded private provision) is possible. If only condition 1 is satisfied, then purely private provision with non-negative profits is possible (examples: software, pharmaceuticals) but will not be efficient if a single price is charged, since average cost greatly exceeds marginal cost (approximately zero).

Optimal provision: $\max W(u^1, u^2, \dots, u^H) \text{ s.t. } F(\mathbf{X}, G) \leq 0$, where $u^h = u^h(\mathbf{x}^h, G)$, $\sum_h \mathbf{x}^h = \mathbf{X}$. $F(\cdot)$ is a very general production function that is convex and obeys constant returns to scale (i.e., homogeneous of degree zero), where inputs are negative arguments and outputs are positive; $\mathbf{X}$ is the vector of private inputs and outputs and G is the output of the public good.

Form a Lagrangian $L = W(u^1, u^2, \dots, u^H) - \mu F(\mathbf{X}, G)$; first order conditions are:

$$x_i^h: W^h u_i^h = \mu F_i \quad \forall h, i \qquad G: \sum_h W^h u_G^h = \mu F_G$$

Combining the first condition for different i and h yields the standard result that $MRS = MRT$ for all private goods and all individuals. Dividing the second condition by the first (ranging over h) yields:

$$\sum_h \frac{u_G^h}{u_i^h} = \frac{F_G}{F_i}$$

This says that we should sum MRS^h and set equal to MRT , because everyone consumes the public good. (This is sometimes referred intuitively as vertical summation of demand curves, although there is no market – and no demand curves – in this case.) This classic result is due to Samuelson (*REStat* 1954) and is commonly referred to as the Samuelson condition.

Problem: if we don't have a market, how do we know individual valuations? This lack of information explains why we might settle for private provision (in the case of excludability), even if it falls short of Pareto optimality.

Externalities represent a market failure or market absence that is associated with a functioning market. For example, pollution may result from production in a market that is competitive, but there is no market for the pollution itself. There are many ways to represent externalities, but consider an “atmospheric externality” to which all individuals contribute and which affects all. That is, individual utility is $u^h(x^h, X_N)$, where X_N is aggregate output of the N^{th} consumption good. X_N can have a positive or negative effect on utility, corresponding to positive and negative externalities.

Assuming a CRS production function $F(X)$ and forming a Lagrangian, we get the first order conditions:

$$x_i^h: W^h u_i^h = \mu F_i \quad \forall h, i \neq N$$

$$x_N^h: W^h u_N^h + \sum_h W^h u_{N+1}^h = \mu F_N \quad \forall h$$

The second condition includes an extra term to account for the impact that individual h 's consumption has on all others. Dividing the second condition by the first (ranging over h) yields:

$$\frac{u_N^h}{u_i^h} = \frac{F_N}{F_i} - \sum_h \frac{u_{N+1}^h}{u_i^h} \quad \forall h$$

How can we achieve this outcome? In theory, we can do so by imposing a Pigouvian tax (subsidy) on each individual, equal to the damage (benefit) that individual's consumption of good N causes others. Again, though, we must know the damage or benefit in order to do so.

Other sources of market failure include imperfect competition and imperfect information. One may also include in this category so-called merit goods – cases where we may wish to override individual decisions for reasons of paternalism or because individual choices for some reason (other than imperfect information) fail to reflect the individuals' underlying preferences.