

Microeconomics for Smarter Students

by

John P. Conley
Vanderbilt University

August 2026

Draft
Not for General Distribution

Preface

Economics, as a field, has grown more and more technical in recent years. At the same time, the world has grown more and more complicated. Intermediate microeconomics, however, continues to be taught in much the same way it has always been. The technical gap between undergraduate and graduate microeconomics is enormous. Students come to graduate and professional school unprepared for the rigor of the courses they will need to take. Requiring students to take more math and statistics is only a partial solution. Economics is still a social science, although a very technical one. What gives it its power is the connection between its formal and rigorous methods and messiness of the real world.

One might argue that the fraction of intermediate students who eventually go on to advanced study is relatively small. Although it might be a disservice to leave these students so unprepared, our greatest responsibility has to be to the majority who just want to gain a basic understanding. A little knowledge, however, can be a dangerous thing. It is important that the type and level of training we give our students equips them to make good judgments about problems they are likely to face in the real world. We do not want students to accept a bad argument more readily because it comes wrapped in an economic model that they understand poorly or which may be misapplied.

Almost anything can be made to sound plausible, and showing a few graphs or statistics to back up a claim can give a wrong idea the veneer of scientific truth. People increasingly live in echo chambers that repeat and reinforce the ideas they believe to be true. Fake news, selective statistics, and simplistic solutions to complicated problems will only become more widespread in the future. Our students will need tools to take arguments apart and see what assumptions are being made, what linkages and feedback paths are being ignored or minimized, and to identify which parts are facts and logic, and which are guesses and opinion.

We see economics first and foremost as a powerful tool for critical thinking. One of the chief goals of this book is to teach students how to write down clear and unambiguous statements of ideas. This is the foundation of scientific communication and makes it possible to test and falsify arguments. You cannot debate an idea until you agree about what it is you are discussing. Imprecision leads to confusion, talking at cross purposes, and never leads to useful solutions.

Precision and clarity in economics are just a means to the end of understanding real world questions of political, social, and commercial importance. A second goal of this book is to develop what is called “economic intuition.” By this we mean the ability to look at a situation and get an immediate feel for what economic factors are at work. This requires fast and dirty applications of our analytical tools.

We will therefore spend a lot of time developing and exploring economics using pictures and graphs. These graphical tools are simple, but they often show us where the bodies are likely to be buried. Of course, as in any science, we must then take out our tools and rigorously verify that our intuitions are correct. Without economic intuition, however, we would have no idea where to put our shovels in the ground.

The strategy of this book is to take a three layered approach. We start with explanations and examples of the economic concepts at hand. Next, we develop graphical tools to help formalize these concepts and allow students to experiment with different scenarios and assumptions to gain a deeper understanding. Finally, we provide formal definitions of the key concepts and mathematical statements of the underlying models.

The objective is to show the essence of the factors at work and the assumptions that underlie the graphs and words which inform and limit how they can be applied. Students will see how assumptions lead to conclusions and when those conclusions require us to ask further empirical questions or to make value judgments beyond the scope of economics.

Some students are verbal and prefer learning through reading and lectures; others are spatial and prefer pictures and graphs; still others are analytic and prefer formal logic and mathematics. All of these are important.

Our hope is that this book will provide a bridge that allows all three types of students to venture into less familiar ways of thinking. We hope that students with verbal orientations will see the connections between informal language and the languages of logic and math. At the same time, we hope to show more analytical students how the formal statements and models they write down can be connected to real economic questions, and why it is so important to do so.

To conclude, our focus is to give students an exposure to the rigorous, formal methods that give economics its power. At the same time, we root our development firmly in the real world questions that give economics its relevance.

We have included an appendix that covers all the mathematical notation and concepts that we use in the book. A minimal familiarity with calculus is all a student should need. We have made a particular effort to ensure that the formal parts of the book run parallel to the text and graphs. Students who have difficulty with more mathematical approaches will not find themselves lost as they read through the chapters. Although we think that students will profit to the extent they understand the real foundations of economics, it profits no one if math becomes a barrier to entry.

Table of Contents

Chapter 1. Basics and Big Ideas.....	9
Section 1.1. What is Economics?.....	9
Section 1.2. Economics is Hard!.....	11
Section 1.3. Economics is Easy!.....	12
Section 1.4. Supply and Demand.....	14
Section 1.5. Prices and Cost.....	17
Section 1.6. Comparative Advantage and Gains from Trade.....	18
Chapter 2. Consumers: Objectives and Constraints.....	26
Section 2.1. Commodity Bundles and Consumption Sets.....	26
Section 2.2. Preferences.....	32
Section 2.3. Properties of Preferences.....	34
Section 2.4. Indifference Curves.....	52
Section 2.5. Needs, Wants, and Preferences.....	57
Section 2.6. Utility Functions.....	59
Section 2.7. Constraints.....	68
Chapter 3. Consumers: Optimization and Demand.....	80
Section 3.1. Optimization.....	80
Section 3.2. Income and Substitution Effects.....	89
Section 3.3. Giffen Goods.....	98
Section 3.4. The Primal Problem and Marshallian Demand.....	102
Section 3.5. The Dual Problem and Hicksian Demand.....	107
Section 3.6. Elasticity.....	109
Section 3.7. Price Indices.....	113
Section 3.8. Head Taxes and Income Taxes.....	121
Chapter 4. Producers.....	132
Section 4.1. Production Functions and Isoquants.....	132
Section 4.2. Returns to Scale.....	135
Subsection 4.2.1. Natural Monopoly.....	136
Section 4.3. Cost Minimization.....	139
Section 4.4. The Long Run, the Short Run, and Cost Functions.....	143
Chapter 5. Supply and Demand in Competitive Markets.....	154
Section 5.1. Aggregate Demand of Competitive Consumers.....	155
Section 5.2. The Supply Curve for Competitive Firms.....	158
Section 5.3. Short Run and Long Run Shutdown.....	159
Section 5.4. The Short Run Aggregate Supply Curve.....	162
Section 5.5. The Long Run Aggregate Supply Curve.....	163
Section 5.6. Example: Short and Long Run in a Long Run Constant Cost Industry.....	173
Subsection 5.6.1. Change in Demand.....	174
Subsection 5.6.2. Change in the Cost of a Fixed Factor.....	176
Subsection 5.6.3. Change in the Cost of a Variable Factor.....	178
Chapter 6. Equilibrium.....	184
Section 6.1. Partial Equilibrium.....	184

Section 6.2. Decentralization.....	190
Section 6.3. Competitive Markets and General Equilibrium.....	193
Section 6.4. Pure Exchange Economies and the Edgeworth Box.....	195
Subsection 6.4.1. The Model.....	196
Subsection 6.4.2. Prices and Competitive Equilibrium.....	197
Subsection 6.4.3. Welfare Theorems and the Core.....	206
Section 6.5. A Simple Economy with Production.....	211
Section 6.6. General Equilibrium, More Generally.....	214
Chapter 7. Welfare Economics and Social Optimum.....	226
Section 7.1. What is Welfare Economics?.....	226
Section 7.2. Consumer and Producer Surplus.....	227
Section 7.3. Social Optimum and Positive Economics.....	232
Section 7.4. Policy Analysis.....	233
Subsection 7.4.1. Sales Tax.....	233
Subsection 7.4.2. Subsidy.....	236
Subsection 7.4.3. Price Ceiling with Shortage.....	238
Subsection 7.4.4. Price Ceiling without Shortage.....	246
Subsection 7.4.5. Price Floor without Surplus.....	249
Subsection 7.4.6. Price Floor with Surplus.....	252
Subsection 7.4.7. Tariff in a Small Country with no Domestic Production.....	259
Subsection 7.4.8. Tariff in a Large Country with no Domestic Production.....	261
Subsection 7.4.9. Tariffs and Trade War Between Large Countries.....	263
Subsection 7.4.10. Tariffs and Quotas in a Small Country with Domestic Production....	265
Section 7.5. Bergson-Samuelson Social Welfare Functions.....	268
Section 7.6. Arrow Impossibility Theorem.....	271
Chapter 8. Market Power.....	278
Section 8.1. Price Taking and Making.....	278
Section 8.2. Monopoly Pricing.....	279
Section 8.3. Monopsony.....	282
Section 8.4. Why Do Monopolies Exist?.....	284
Subsection 8.4.1. Natural Monopoly.....	284
Subsection 8.4.2. Network Externalities.....	286
Subsection 8.4.3. Patents, Copyrights, and Trademarks.....	287
Subsection 8.4.4. Legal Barriers to Entry.....	289
Subsection 8.4.5. Natural Resource Monopoly.....	290
Section 8.5. Government Solutions to Monopoly.....	291
Subsection 8.5.1. Subsidizing Monopoly.....	291
Subsection 8.5.2. Price Ceilings.....	293
Subsection 8.5.3. Rate of Return Regulation.....	294
Subsection 8.5.4. Sophisticated Monopoly Pricing Strategies.....	297
Subsection 8.5.5. Price Discrimination.....	298
Subsection 8.5.6. First Degree Price Discrimination.....	298
Subsection 8.5.7. Second Degree Price Discrimination.....	299
Subsection 8.5.8. Third Degree Price Discrimination.....	301

Subsection 8.5.9. A Numerical Example of Third Degree Price Discrimination:.....	303
Subsection 8.5.10. Two-Part Tariffs.....	303
Subsection 8.5.11. Resale Price Maintenance.....	305
Subsection 8.5.12. Durable Goods Monopolists.....	306
Section 8.6. Acquiring Market Power.....	308
Subsection 8.6.1. Mergers.....	308
Subsection 8.6.2. Horizontal Integration.....	308
Subsection 8.6.3. Vertical Integration.....	308
Subsection 8.6.4. Predatory Pricing.....	309
Subsection 8.6.5. Duopoly and Oligopoly.....	309
Chapter 9. Noncooperative Game Theory.....	320
Section 9.1. What is Noncooperative Game Theory?.....	320
Section 9.2. Normal Form Games.....	323
Subsection 9.2.1. Nash Equilibrium.....	323
Subsection 9.2.2. Dominant Strategy Equilibrium.....	324
Section 9.3. Examples of One Shot Games.....	326
Subsection 9.3.1. Coordination Games.....	326
Subsection 9.3.2. The Hotelling Location Game.....	327
Subsection 9.3.3. Prisoners' Dilemma.....	329
Subsection 9.3.4. Zero-Sum Games.....	330
Section 9.4. Mixed Strategies and the Folk Theorem.....	332
Subsection 9.4.1. Mixed and Pure Strategies.....	332
Subsection 9.4.2. Repeated Games.....	333
Subsection 9.4.3. The Folk Theorem.....	334
Section 9.5. Extensive Form Games.....	338
Subsection 9.5.1. Subgame Perfect Equilibrium.....	343
Subsection 9.5.2. Games with Imperfect or Incomplete Information.....	345
Section 9.6. Examples of Extensive Form Games.....	349
Subsection 9.6.1. The Ultimatum Game.....	349
Subsection 9.6.2. The Cake Eating Game.....	350
Subsection 9.6.3. Centipede Game.....	351
Subsection 9.6.4. All Pay Auction: Dollar Bidding Game.....	352
Section 9.7. Applications of Noncooperative Games.....	353
Subsection 9.7.1. Standards.....	353
Subsection 9.7.2. Standards Wars.....	354
Subsection 9.7.3. Auctions.....	358
Subsection 9.7.4. Lock-in.....	362
Subsection 9.7.5. Hold-up.....	365
Subsection 9.7.6. Bank Runs and Bandwagons.....	366
Subsection 9.7.7. Experimental Approaches.....	366
Chapter 10. Cooperative Game Theory.....	377
Section 10.1. What is Cooperative Game Theory?.....	377
Section 10.2. The Core.....	378
Subsection 10.2.1. Pareto Optimality and Individual Rationality.....	379

Section 10.3. Examples of the Core.....	381
Subsection 10.3.1. Example 1.....	381
Subsection 10.3.2. Example 2.....	383
Subsection 10.3.3. Superadditivity.....	384
Section 10.4. The Shapley Value.....	385
Section 10.5. Cooperative Bargaining Theory.....	388
Subsection 10.5.1. The Nash Solution.....	391
Subsection 10.5.2. The Kalai-Smorodinsky Solution.....	395
Subsection 10.5.3. The Egalitarian Solution.....	396
Appendix A. Conventions.....	402
Section A.1. A List of Mathematical Symbols and Notation.....	402
Section A.2. Index and Vector Notation.....	403
Section A.3. Index Sets.....	404
Subsection A.3.1. Sets, Vectors, and Functions.....	405
Subsection A.3.2. Markets.....	405
Subsubsection A.3.2.1. Partial Equilibrium Analysis.....	405
Subsubsection A.3.2.2. General Equilibrium Economies.....	405
Subsection A.3.3. Games.....	406
Subsubsection A.3.3.1. Noncooperative Games.....	406
Subsubsection A.3.3.2. Cooperative Games and Coalitions.....	406
Appendix B. Mathematical Review.....	407
Section B.1. Vectors, Sets, and Order.....	407
Subsection B.1.1. Euclidean Spaces and Vector Order.....	407
Subsection B.1.2. Sequences.....	409
Subsection B.1.3. Open, Closed and Bounded Sets.....	409
Subsection B.1.4. Convex and Comprehensive Sets.....	411
Section B.2. Hyperplanes and Separation Theorems.....	412
Section B.3. Rules for Calculus.....	413
Subsection B.3.1. Derivatives.....	413
Subsection B.3.2. Algebra of Exponents.....	413
Subsection B.3.3. Algebra of Natural Logarithms.....	414
Section B.4. Functions.....	415
Subsection B.4.1. Continuous Functions.....	415
Subsection B.4.2. Monotonic Functions.....	415
Subsection B.4.3. Homogeneous, Linear and Affine Functions.....	416
Subsection B.4.4. Concave and Convex Functions.....	416
Section B.5. Optimization.....	418
Subsection B.5.1. Optimization without Constraints.....	418
Subsection B.5.2. Optimization with Constraints.....	420
Appendix C. Logic and Proofs.....	422
Section C.1. Logical Statements and Propositions.....	422
Section C.2. Basic Logical Connectives.....	422
Section C.3. Propositional Connectives.....	425
Section C.4. Quantifiers.....	429

Section C.5. Proofs.....	430
Appendix D. Game Theory.....	433
Section D.1. Noncooperative Static Games.....	433
Subsection D.1.1. Definitions.....	433
Subsection D.1.2. Equilibrium Concepts.....	434
Subsection D.1.3. Mixed Strategies.....	434
Section D.2. Extensive Form Games.....	435
Subsection D.2.1. Definitions.....	435
Subsection D.2.2. Subgames and Subgame Perfect Equilibrium.....	437
Subsection D.2.3. Information Partitions.....	438
Section D.3. Cooperative Games.....	440
Subsection D.3.1. Games in Characteristic Function Form.....	440
Subsection D.3.2. Bargaining Theory.....	440

Chapter 1. Basics and Big Ideas

Section 1.1. What is Economics?

The word “economics” is derived from the Greek: *oikonomikos* meaning home (*oikos*) management (*nomos*). From this humble beginning, economics has become something of a monster (*terástios*, in Greek) in the social sciences.

Economists now do field experiments like sociologists, experiments on human subjects like psychologists, and fMRI studies like medical doctors or neurologists. Economists use their tools to study questions from political science, geography, law, computer science, history, philosophy, and even theology.

Economics has been called the *imperial science*. There are few topics that economists will not comment upon, from how to run a country, to how to find a mate. In short, economics is the barbarian knocking at the gates of almost every other field of study. We are the Borg of the academic universe. Whether you choose to resist or assimilate, it is best to understand what you are facing.

So what is economics anyway? Many definitions have been offered. For example:

Economics is extremely useful as a form of employment for economists.

– John Kenneth Galbraith

Economics is what economists do.

– Jacob Viner

Economics sometimes gets a bad name for its seeming focus on greed and selfishness as the prime movers of human behavior. Adam Smith famously described this in the *Wealth of Nations* as follows:

It is not from the benevolence of the butcher, the brewer, or the baker that we expect our dinner, but from their regard to their own interest.

– Adam Smith

Put more succinctly:

There ain't no such thing as a free lunch.

– Robert A. Heinlein

Perhaps this is why Thomas Carlyle characterized economics as the “dismal science.” One which takes as given that you cannot have all that you want and that both individuals and societies must and do make difficult choices motivated to a great extent by self-interest.

This is a rather negative take on what is essentially a profound and positive observation. What Smith was pointing out is that it is self-interest that causes people to think very hard about the needs and wants of others because it profits them to do so. We can count on the help and support of others, and they, in turn, can count on ours, precisely because of self-interest.

Self-interest is a powerful, even irresistible, force that holds the social fabric together. Self-interest is the invisible hand that makes it all work.

Because of this, economists seek to understand the details of how this mechanism works and how it might be improved to better serve our interests. This notwithstanding, economists have a rather poor reputation:

If all the economists were laid end to end, they would not reach a conclusion.

– George Bernard Shaw

Give me a one-handed economist! All my economists say: On the one hand, on the other.

– Harry S. Truman

I guess I should warn you, if I turn out to be particularly clear, you've probably misunderstood what I've said.

– Alan Greenspan

Perhaps it does not matter too much in any event. After all, as John Maynard Keynes reminded us:

In the long run, we are all dead.

– John Maynard Keynes

Section 1.2. Economics is Hard!

It may not surprise you to learn that economics is hard. The subject matter is almost impossibly complex. Consider the following:

People are complicated: They act due to love and hate, fear and loathing, and duty and honor. They make mistakes, and they are sometimes irrational. Let's be sure not forget our old friends, ignorance, and stupidity. Certainly people are more than rational maximizers motivated by greed.

The world is unpredictable: We do not know who will be chosen president in the next election, whether there will be a drought or an oil shortage, if a massive terrorists attack will galvanize the nation, or if a new disruptive innovation like the internet will appear. Future events such as these have an enormous effect on the path of the economy, and no one, economists included, can credibly predict them.

The world is a complex system: Everything is related to everything else. Prices and quantities in all markets are determined simultaneously. For example: higher oil prices lead to a lower demand for cars, which decreases the demand for steel, which reduces demand for coal, which lowers coal prices, which leads to cheaper electricity, which lowers costs for most producers, which lowers most output prices, which causes consumers to demand more goods, which requires more freight and transportation services to get our purchases delivered from Amazon, which raises the demand for, and therefore the price of, oil ... This is only one of millions or billions of possible feedback channels. How can economists hope to make predictions or offer advice in such an environment?

Information is imperfect: How many people are actually employed in China? What is the true level of unfunded pension obligations of state and governments? Are a set of mortgage securities really as good as the rating agencies say they are? How do I figure out if people are reporting their full income to the IRS? Will this new educational reform have the effects that were promised? Does the executive branch have the expertise to put together a working health-care exchange? In other words, not only is it impossible to predict future events of economic importance, it is quite difficult to get a clear picture of current and past events as well.

How can economics possibly be useful in light of all this? How can economists give credible predictions of future inflation and unemployment, give financial and investment advice, or make policy recommendations about trade, welfare, and macroeconomics? Maybe we should just go have a beer and call it a day.

Section 1.3. Economics is Easy!

The questions economists are asked to address are very difficult indeed. Despite this, it may not surprise you to learn that what economists actually do is considerably easier.

Economics is, at heart, a **behavioral science**. We study people, but very stylized people. We do this by using **models**.

The world is complex, and so to gain traction, economists simplify by imposing assumptions and elementary mathematical structures. By analyzing these simple models, economists hope to gain insight into what happens in the real world. The art of making a model is to include the elements of the real world that are most important while getting rid of less important complicating details.

For example, we could model a tree as a rational maximizer of photosynthetic potential. Notice that (a) this is completely false, but (b) it gives fairly good predictions about how the leaves and branches are distributed.

Given that the model is based on the false assumption that a tree has a brain and can make decisions, should we reject it out of hand? It depends upon what our purpose is. If we want to predict how leaves are distributed over the branches of the tree, this is a fairly good model. If instead we were concerned with how the tree feels about losing its leaves in the fall, the model would lead us badly astray.

Suppose instead that we wanted to know how long it would take a brick dropped off a highway overpass to hit the ground. We might model the brick as a perfect sphere of uniform density falling in a perfect vacuum towards a planet of known mass and uniform density. We would therefore use the rule that objects on Earth fall at a rate of 32 feet per second squared if otherwise unimpeded.

All of these assumptions are false. The brick is rectangular, it is falling in an atmosphere that will slow it down, the Moon or the Rocky Mountains may exert some gravitational pull of their own and slightly deflect the brick from falling directly towards the center of the Earth, the brick might hit a butterfly on the way down, and so on. If we took these into account, our predictions would be more accurate, but not by much. Clearly, it would not be worth the effort to complicate the model with these details.

What if we dropped the brick out of an airplane? What if we dropped a sheet of paper off the overpass instead of a brick? In both cases, the objects would reach their terminal velocity long before they hit the ground due to wind resistance. Now our simplifying assumptions would not only be false, but would lead us to make very poor predictions about when the objects should hit the ground. Thus, what may be a good model in one situation, may be a bad one in another.

The point is that abstracting from reality through models can make extremely complicated problems, tractable. The fact that a model is too simple, unrealistic, or even outright false is, in itself, irrelevant. The proof of the pudding is in the eating. If the model gives systematically good predictions when applied correctly, it is a good model.

At the most basic level, economists model firms, consumers, and other economic actors as *agents who maximize objectives given constraints*. As we said before, people are much more complicated than this. Nevertheless, it turns out that in large groups, firms, and consumers *behave as if they were rational maximizers* in a statistical sense.

The deviations we see from what a hypothetical rational maximizer (sometimes referred to as **homo economicus**) would do tend to be randomly distributed around the “rational choice”. Thus, on average, groups of agents behave rationally even if no individual member of the group is rational as an economist would define it.

The models we develop in this book are aimed at addressing questions from a positive rather than a normative standpoint. By this we mean:

Positive economics: Statements about what is. No opinions are offered, just facts and analysis that follow from assumptions.

Normative economics: Statements about what should be. These involve value judgments based on political, religious, philosophical, and ethical beliefs.

This is why we talk about economics as a science. Just like a doctor or a physicist, our job is to give the best possible prediction of how cause leads to effect. It is not the doctor’s job to tell you that you should exercise, just that you are likely to die sooner if you do not. It is not the physicists job to tell you that you should not drop bricks off of highway overpasses, but to tell you the consequences if you choose to do so.

In the same way, it is not the job of economists to tell you to support lower taxes, free trade, or welfare reform. Our job is to use analytical tools, like the ones we will develop in this course, to understand the implications of such choices and then allow others (and ourselves) acting as citizens to decide which outcomes we prefer as individuals or a society.

Section 1.4. Supply and Demand

One of the most useful tools that economists have is supply and demand analysis. We will derive supply and demand curves from first principles later in this book, but it will be useful to have a little background and vocabulary before we do so.

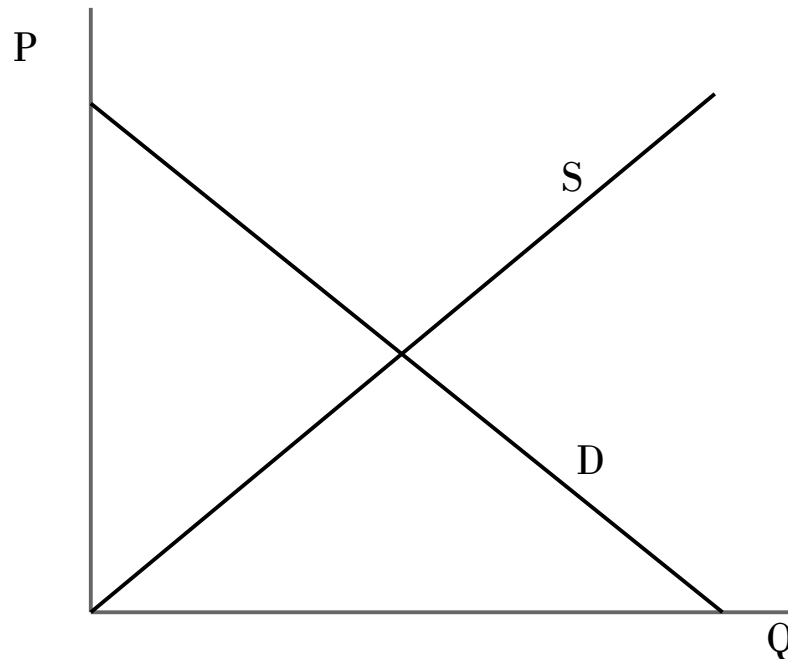


Figure 1: Supply and demand curves

Demand Curve: A statement of voluntary and optimal consumer behavior.

Supply Curve: A statement of voluntary and optimal firm behavior.

More specifically, supply and demand curves tell us how much of a commodity producers or consumers are willing to supply or demand at any given price. To understand what is being conveyed by a supply or demand curve we need to know the following:

- **What:** What physical commodity is being considered including where, when, and why it is delivered.
 - Where: Where is the good delivered (commodities markets)
 - When: When is the good delivered (futures and derivative contracts)
 - Why: Under what contingency is the good delivered (insurance and derivative contracts)
- **Who:** Who or what group of firms or consumers are represented by the curve.

- **When:** Over what unit of time is supply or demand expressed.
- **Unit of Consumption:** What is the unit of consumption on the quantity axis.
- **Unit of Exchange:** What currency or commodity is represented on the price axis.

Thus, a demand curve might be for apples, consumed by residents of Idaho, each month, as measured in pounds, and with price given in dollars. Alternatively, a supply curve might be for #2 red winter wheat, to be delivered on March 1, 2024, to Chicago, by farmers participating in the Chicago Board of Trade March futures market, in dollars per bushel. Given all this, we can now state two fundamental economic laws:

THE LAW OF DEMAND: DEMAND CURVES ARE DOWNWARD SLOPING.

THE LAW OF SUPPLY: SUPPLY CURVES ARE UPWARD SLOPING.

Our direct observation and examination of the evidence (that is, our **empirical** investigation) suggests that both of these laws are generally true in the real world. This is especially so for the law of demand. However, we can generate theoretical conditions under which demand curves would be upward sloping and supply curves downward sloping. Thus, both of these are empirical rather than theoretical laws. We will discuss this at more length below.

Finally, when markets are competitive, and all agents are price takers, we can state a third fundamental economic law:

THE LAW OF ONE PRICE: ANY GIVEN COMMODITY OR SECURITY MUST TRADE AT ONE AND ONLY ONE PRICE IF MARKETS ARE COMPETITIVE.

The basic idea here is that in any given competitive market, one equilibrium price will prevail. If your local bar offered frosty pints of the same locally sourced, organic, GMO-free, artisanal IPA for both \$6 or \$8, you would always choose pay \$6 No \$8 beers would ever be sold. Similarly, you cannot have two identical houses side by side selling for different prices.

To extend the idea, suppose that the price of steel in Shanghai is different than it is in Cleveland once exchange rates, taxes, duties, transactions costs, and so on are taken into account. All it takes to turn a ton of steel in Shanghai into a ton of steel in Cleveland (or the reverse) is adding some transportation services to it. Thus, the price gap can be no more than this.

For example if it costs \$8, to move a ton of steel between these two places, steel can be at most \$8 more expensive or cheaper in Shanghai as compared to Cleveland. Otherwise, there would be an arbitrage opportunity and one could profit from buying in China and selling in Ohio, or the reverse.

On the other hand, two beers of different kinds, or at different times of day, or at different bars, might have different prices, but then these are not the same good. Recall that the same physical good delivered at a different time and place has different value and is therefore a different good. In other words, the law of one price also implies that if one good (a ton of steel in Tangshan) can be converted to another good (a ton of steel in Gary, Indiana) by moving it, storing it, painting it, shaping it, teaching it to read, blessing it, or through any other process, the price difference can be no more than the cost of applying the process.

Of course, if markets are less than competitive because people do not know the price of beer, houses, or steel in every location, or because firms have monopoly power, the law of one price may fail.

Section 1.5. Prices and Cost

Supply and demand curves describe the relationship between price and quantity in a market. What do we mean by price here? There are two important variations:

Nominal Price: The number of dollars that can be exchanged for a unit of a given good.

Relative Price: The number of units of another good that can be exchanged for a unit of a given good.

Most of the time, we see nominal prices (sometimes also call “absolute prices”) in the markets in which we participate. This is because dollars are a convenient **unit of exchange**. People accept dollars in exchange for goods because they expect that other people will also accept dollars when they wish to buy other goods.

Without currency to mediate our exchanges we would have to barter with one another. This means that if I have goats and I wish to eat bread, I have to find a baker who is in need of goats. In other words, there has to be a **mutual coincidence of wants** in order to make beneficial exchanges. You can imagine how difficult it would be to buy a week’s groceries if we had to find individuals who wanted goats, but happened to have extra eggs, milk, beer, toilet paper, etc. with whom we could trade.

On the other hand, a relative price exists for each good with respect to each other good. The relative price of apples compared to oranges is simply the ratio of their nominal prices. If apples cost \$1 and oranges \$8 each, then the relative price of an apple is two oranges. Although money is very convenient, only relative prices matter in the end. If all nominal prices doubled (including your wage rate), you could (and would) buy exactly the same bundle of goods you did before.

When an agent considers buying or selling something, he compares costs and benefits. What do we mean by cost exactly? Part of the cost is the money cost, since if you buy a good, you have to give up some money that might have been used to buy something else. There may be other costs, however.

Suppose you are hungry for a Snickers bar at 11:00 pm. A Snickers bar costs \$1.25 at the local convenience store, and this is its **direct cost**. You also have to get out of bed, miss half an hour of Dr. Who, walk through the rain, add wear and tear to your car, use up gas, and so on. These are **indirect costs**. The sum of all these direct and indirect costs is the “opportunity cost” of getting a Snickers bar.

Opportunity cost: The total value of all aspects of the best foregone opportunity necessary to obtain an object.

Economists assume that agents weigh the opportunity costs of an action against its benefits when making decisions.

Section 1.6. Comparative Advantage and Gains from Trade

Economics is mainly focused on exchanges or trades that people make with one another to improve their own welfare. But why do people trade? Consider the following example:

Chore	My Wife	Me
Mow the grass	3 hours	4 hours
Buy the groceries	2 hours	5 hours

Table 1: Gains from trade and comparative advantage

Not surprisingly, my wife is better at everything than I am. Why in the world would she stay married to me in this case?

Suppose that we considered two possible ways to share the chores. First, we could *alternate* doing the weekly chores. I would mow the grass and buy the groceries one week, and she would do so the next. Second, we could *split* the duties with me mowing the grass and her buying the groceries each week.

If we alternate weeks, we each do each chore twice over the course of a month. I spend a total of 18 hours and my wife spends a total of 10 hours.

If we split chores, we each do our own chore four times over the course of a month. I spend a total of 16 hours and my wife spends a total of 8 hours.

It looks like this marriage can be saved! Even though my wife has an *absolute advantage* over me in all kinds of production, I have a *comparative advantage* over her in mowing the grass (and she has a comparative advantage over me in buying the groceries).

Absolute Advantage: One agent can do something using fewer resources than another.

Comparative Advantage: One agent can do something at a lower relative opportunity cost than another.

I have a comparative advantage in mowing since the opportunity cost for me is .8 of a trip to the grocery store while the opportunity cost for my wife is 1.5 of a trip to the store. In other words, I have to give up 80% of a shopping trip to mow, but my wife has to give up a shopping trip and a

half to mow. I give up less shopping (and thus have smaller opportunity cost) than my wife does to do the same mowing job. Symmetrically, my wife has a comparative advantage in shopping since the opportunity cost for her is .67 of a lawn mowing while the opportunity cost to me is 1.25 of a lawn mowing.

This is one reason that developing countries can trade with the US. Although the majority of their work force may be poorly educated and untrained, they are relatively good at assembling electronics and toys since the opportunity cost of a typical US worker to do the same job is something much higher.

Comparative advantage creates a potential for **gains from trade**. By trading mowing for shopping in the example above, my wife and I each gain two hours of time compared to when we alternate chores. Comparative advantages in production can come from several underlying factors:

Differences in Abilities: I teach economics, my barber cuts hair, my mechanic fixes cars. Each of us is comparatively better at doing our own jobs. We specialize since this gives us the largest income, which we can then trade for other goods and services.

Sometimes these differences are **exogenous**, that is, they are not due to choices made by the agents but instead are simply unchangeable facts or initial conditions. My barber is a nice guy, open, and honest, who deals easily with other people and is good at listening. He is naturally a great barber, but these same talents would make him a terrible economist. Let's just say I would make a terrible barber and leave it at that.

Other times, these differences are **endogenous**, that is, they are the result of choices made by agents either individually or collectively, perhaps in response to market conditions. It might be that my mechanic could have been a barber, and my barber could have been a mechanic. However, both of them have trained for and now practiced in their professions for years. At this point, they have both acquired a comparative advantage at doing their specific jobs even though may have been born with the similar abilities.

An interesting example of endogenous differences can be seen in the large passenger aircraft industry. The US was an early entrant into this industry. The Second World War led to a huge expansion in production.

It is estimated that the US produced something like 300,000 aircraft during the War. As Boeing, McDonnell-Douglas, and others ramped up production, they learned what worked and what did not. They modified their production processes, experimented with cheaper materials, discovered efficient ways of doing things, and so on.

All of this practical experience served to keep US companies further along the technology curve than companies in other countries who were building fewer aircraft and who started later.

This is a process called **learning by doing**. The more you do something, the better you get at it. Because of this, the US has had superior abilities to build aircraft for a long time. Note, however, this is not because there is something about the US that gives it an innate advantage. There is every

reason to believe that if another country had jumped in early and gained experience more quickly than we did, they would now have the comparative advantage.

The information technology sector is similar. There is nothing remarkable about Silicon Valley *per se*. In the 1950s and 1960s it was something of a backwater in the swampy part of San Francisco Bay and was given over to almond orchards and other agricultural uses for the most part. Stanford University, the University of California at Berkeley, several national laboratories, NASA, the military, and several private companies such as IBM shared an interest in electronics and other emerging technologies.

Part of this was driven by the Cold War and the space race. Partnerships and spin-offs were founded to supply the military with their requirements and to exploit these new technologies. It was easy to begin these businesses near Stanford in the 1950s. Land was cheap, graduates of the local universities liked the area, and continuing collaboration with people in university research labs was easier.

You can imagine what happened. This began to take on a life of its own. Soon, being in Silicon Valley gave you access to all kinds of talent and expertise. It was the place to find the right connections, hook up with the right team, find a firm that needed your technical skills, and see the bleeding edge of technological advance.

This is an example of a **neighborhood effect (which** are closely related to **network externalities**). The greater the number of people who work in the IT industry in Silicon Valley, the better it is for everyone. There is an advantage in sharing a neighborhood with others in the same industry.

This is similar to a **network good** in which the value of the good to any given consumer is determined by how many other people consume the good (perhaps by joining the network). Think about Facebook, for example. If only three people have profiles, it is a pretty useless site. On the other hand, if fifty million people have profiles, it is still pretty useless, but you are having too much fun wasting time to notice.

More seriously, it would be difficult to start an IT firm in Alaska. You would forego the benefits of seeing what other entrepreneurs were doing, what is selling, what technologies are in the offing, and who would be a good hire in any given area. Again, we find that the difference in ability that leads to a comparative advantage was an endogenous consequence of how history unfolded, and not to some specific exogenous feature of the West Bay of San Francisco.

Differences in Endowment: Saudi Arabia sends the US oil and the US sends them wheat in return. We both like each of these commodities, but our national endowments are specialized in different ways. If we did not trade, we would have too much bread, and no heat in the winter. The Saudis, in turn, would have all the energy they could possibly want, but would be hungry.

Another example might be France and California. Both places are very good and skilled at producing wine, and both have populations that enjoy wine. However, both are endowed with unique combinations of soil and climate and this causes the wines they produce to be different. Since variety is the spice of life, wines are traded back and forth.

Differences in Taste: Suppose every nation in the world had the same abilities and the same endowments. Then all nations have identical absolute costs of production and so no nation has a comparative advantage at producing anything. Despite this, trade can still be beneficial if there are differences in tastes between nations. For example, matsutake mushrooms grow in Japan, the Pacific Northwest of the US and Canada, and several other places. They are naturally occurring in evergreen forests and cannot be cultivated. Harvesting them is not especially hard, so no country has a particularly better or worse ability to produce them. Both Japan and the US have similar endowments of forests suited to their growth. Both Japanese and Americans like to eat matsutakes, however, the Japanese like them a great deal more. Because of this, the price that the Japanese are willing to pay far exceeds what they fetch in most other countries. Both countries are made better off when they trade American matsutake mushrooms for Sony flat screen TVs.

At the individual level, suppose my wife and I are equally good at doing everything. Since we are married, what is mine is hers, and what is hers is mine. Thus, our endowments are the same. However, I hate killing spiders, while my wife has a certain blood thirstiness in her nature. On the other hand, I find ironing the sheets to be relaxing after a hard day at the office proving theorems. She finds it boring. You can see that even though both of us are completely capable of doing the job we dislike, we are both better off if we trade sheet flattening for spider squishing services due to differences in preferences.

We conclude that any one of these three factors by itself is enough to provide a foundation for mutually beneficial trade.

Glossary

Absolute Advantage The ability of one economic actor to complete some action with fewer inputs, or at lower costs than another. For example, the ability of one country to produce a product more cheaply than another, or of one worker to complete a task more quickly than another.

Behavioral Science: A branch of study that examines the behavior of humans (and animals) using theoretical, experimental, empirical and other scientific methods in order to understand decision-making, communication, and interaction.

Comparative Advantage: The ability of one economic actor to complete some action at a lower relative opportunity cost than another. For example, suppose China can produce an airplane using the resources it would have taken it to produce ten buses, while France could produce airplane using the resources it would take it to produce five buses. Then France has a comparative advantage in producing airplanes (since the opportunity cost is giving up five buses rather than ten) while China has a comparative advantage producing buses (since the opportunity cost is only one tenth of an airplane instead of one fifth).

Demand Curve: A statement of voluntary and optimal consumer behavior.

Direct Cost: The money costs of labor, material and other inputs required to make a product, or more generally, to undertake any activity.

Economics: The science of allocating scarce resources over competing ends.

Endogenous: Originating, developing, or proceeding from within. In an economic context, something is endogenous if it is the result of choices made by agents either individually or collectively, perhaps in response to market conditions. For example, when we observe that a person is a vegetarian, is majoring in physiology or owns a new car, we are noting facts that are endogenous since they are the result of choices made by that person.

Exogenous: Originating, developing, or proceeding from outside. In an economic context, something is exogenous if it is the result conditions out of the control of an agent. For example, the fact that an agent is allergic to animal protein, is not smart enough to major in economics, or has wealthy parents are exogenous conditions. They may affect the choices made by an agent, but they are not themselves choices.

Gains from Trade: The amount by which agents benefit from buying, selling, bartering, and engaging in other types of voluntary exchange. This might be measured in money or consumer and producer surplus, for example.

Homo Economicus: A fictional idealized agent who narrowly and rationally maximizes his own self-interests.

Indirect Cost: The value of things that are required to make a product (or more generally, to undertake any activity) which are not directly paid for with money in the production process. For example, the value of the labor I devote to running a business for which I receive no salary or the time I spend standing in line to buy something.

Invisible Hand: The unseen force which seems to coordinate the actions of firms, consumers and other agents to act in ways that benefit one another and generate stable economic outcomes, even though all agents pursue only their own personal interests and neither know, nor care, about the choices made by other agents in the economy. This metaphor comes from the writings of Adam Smith, especially *The Wealth of Nations* (1776) :

*Every individual necessarily labours to render the annual revenue of the society as great as he can. He generally neither intends to promote the public interest, nor knows how much he is promoting it . . . and by directing that industry in such a manner as its produce may be of the greatest value, he intends only his own gain, and he is in this, as in many other cases, led by an **invisible hand** to promote an end which was no part of his intention. Nor is it always the worse for the society that it was not part of it. By pursuing his own interest he frequently promotes that of the society more effectually than when he really intends to promote it. I have never known much good done by those who affected to trade for the public good. It is an affectation, indeed, not very common among merchants, and very few words need be employed in dissuading them from it.*

Law of Demand: Demand curves are downward sloping.

Law of One Price: Any given commodity or security must trade at one and only one price if markets are competitive.

Law of Supply: Supply curves are upward sloping.

Learning by Doing: Becoming faster, better, or more efficient at an activity as you do it more often. In particular, learning by doing causes the average cost of output produced in the future to be less than it is in the present.

Model: A simplified abstraction used to illustrate or understand complex situations or objects. Economists often describe an economy or market using a rigorous mathematics framework that makes it possible to see what different assumptions regarding behavior, objectives, or structure imply about economic outcomes. These implications may then be tested empirically to see if they agree with real world observations.

Neighborhood Effects: A beneficial crowding externality. For example, shops of a certain kind often choose to locate close together. Although they must compete with each other and this may lower prices, customers shopping for rugs, cars, furniture, and so on, will be more willing to travel to a place with many alternatives than to a single store. Thus, co-locating brings in more customers and may actually raise demand, prices, and sales volume. Neighborhood effects are

also seen on the production side with similar firms locating in the same place to benefit from larger and deeper labor and supply networks. Silicon Valley and Detroit's auto and steel manufacturing centers are good examples of this.

Network Externality: A situation in which the benefit to an agent of consuming a product, joining a group, or undertaking some other action is positively correlated to the number of other agents who undertake the same actions. For example, the more people who have a telephone, the more useful it is for me to have a telephone since I can reach more people. Facebook and other social networks are more modern examples of the same phenomenon.

Nominal Price: The number of dollars (or other currency) that can be exchanged for a unit of a given good. In other words, the money price of a good.

Normative Economics: Statements about what should be. This involves value judgments based on political, religious, philosophical and ethical beliefs.

Opportunity cost: The total value of all aspects of the best foregone opportunity necessary to obtain an object. This includes both the direct costs (since the money paid for inputs cannot be spent on other things) and all indirect costs.

Positive Economics: Statements about what is. No opinions are offered, just facts and analysis that follow from assumptions.

Relative Price: The number of units of one good that can be exchanged for a unit of another good. Note that there is a relative price for each good in terms of each other good (and so $N \times (N-1)$ relative prices in total if there are N goods in the economy).

Supply Curve: A statement of voluntary and optimal firm behavior.

Problems

1. What is the difference between positive and normative economics (or more generally, science)? Are there any circumstances you can think of where an economist, acting as an economist, would be able to say that one policy or outcome is better than another? Could an economist or a scientist ever be justified in stating which policy or outcome is the best regardless of other circumstances?
2. Explain how the following things would change the equilibrium price of wheat, and why.
 - a. The price of corn goes up.
 - b. The price of fertilizer goes down.
 - c. The price of butter goes down.
3. Oranges grow in Florida and California. Yet, it is often observed that the average quality of oranges that are shipped to the Midwest for sale is much better than those that remain in Florida and California to be consumed locally. Give an economic explanation for the strange fact that more bad oranges are sold where they are grown, and more good oranges are sold everywhere else. (Hint: Think about the relative prices of these two commodities in the Midwest compared to California.)
4. It is often said that you can save money by cutting out the middleman. What role do the middlemen play in the economy? Are they parasites or do they contribute to economic efficiency? Try to discuss this question by making specific reference to the issues in this chapter.
5. On the small Pacific island of Tiniwini, the people are not very well off. One day, a United Nations official docks at Tiniwini's excellent and busy harbor. He is shocked to discover that despite the extreme difficulty that the islanders experience in collecting enough food to eat, most of the people are engaged in the manufacture of pukka-shell necklaces. He strongly advises the UN to start a program to encourage subsistence agriculture on Tiniwini. Do you think the people of Tiniwini owe this farsighted official a debt of gratitude, or can you imagine a reason why they should persist in making junk jewelry when it is almost impossible to find enough to eat on the island?

Chapter 2. Consumers: Objectives and Constraints

Section 2.1. Commodity Bundles and Consumption Sets

In microeconomics, consumers are modeled as agents who maximize objectives given constraints. Consumers are assumed to be “small” relative to the market as a whole, and so are treated as nonstrategic **price takers**. In other words, consumers in competitive markets (discussed in more detail in Chapter 9) have no power to influence the prices of the goods they consume. Game theoretic approaches (discussed in more detail in Chapter 9), in contrast, allow agents to behave strategically. Consumers and firms may be able to affect prices in the markets in which they participate in they are “large.”

We begin by considering the *consumer’s problem*:

Consumer’s Problem: Choose the most preferred consumption bundle from a set of feasible alternatives.

A consumption bundle is made up of different amounts of the various commodities available in the economy. We use the term **commodity** because of its generality. A commodity may be:

- Desired (**goods**) and not desired (**bads**) by agents: For example, iPhones are goods unless you happen to really hate being part of Apple’s ecosystem (in which case they might be bads).
- Produced by firms, found in nature, or part of the endowment of agents: For example, iron, iron ore, and the time of ironworkers, respectively.
- Material or immaterial: For example, corn or T-shirts in contrast to legal services, or MP3s.
- Divisible or indivisible For example, seconds of cell phone usage or gallons of gasoline in contrast a marriage ceremony or a car.
- Unique or common For example, the Mona Lisa or the Empire State Building in contrast to apples or tires.
- Durable or non-durable: For example, A car or a house are durable, and we actually consume the stream of services they provide while an apple or a hamburger disappear as they are consumed.

- **Storable or perishable:** For example, a bottle of wine, or an MRE, are storable goods and can be consumed at any time after purchase, while a seat on a specific flight has to be consumed at a specific time or not at all. Storable goods may appreciate, like a good Bordeaux or an oak tree that increases in value if consumption is delayed. Both storable and perishable goods may depreciate as well such as a suit that slowly goes out of fashion as time goes on or a loaf of bread that gets stale, then moldy, and ultimately turns to dust.
- **Rival or non-rival:** For example, a radio broadcast is a non-rival or public good since when one person tunes in he does not reduce the ability of others to turn in and consume the same amount of the broadcast. A hamburger is a rival or private good since every bite I eat leaves exactly one less bit for everyone else.

It must be the case, however, that every unit of a given commodity is identical to every other unit. That is, a given commodity type defines a homogeneous class of items.

Technical Details

Text surrounded by a green box is meant to indicate that it contains technical details that may not be completely familiar. However, the material is not part of the main economic discussion. If the technical details are familiar to you, just ignore the box and continue reading where it ends. If not, have a look the mathematical appendices for a complete exposition.

Typically, we consider a finite space of commodities. For example, we might think of a world with only three goods: food, clothing and shelter. We would represent a typical consumption choice as a vector:

$$x \equiv (x_f, x_c, x_s) = (4, 0, 1) \in \mathbb{R}^3.$$

In words, this statement means that a consumption bundle, (x_f, x_c, x_s) is an element of a three-dimensional Euclidean space, $\mathbb{R}^3$, and lists a level of consumption of each good. $x_f = 4$, $x_c = 0$, and $x_s = 1$. In this example, our consumer seems to be well-fed, adequately housed, but naked since he has four units of food, one unit of housing, but no clothes at all in his consumption bundle. It is sometimes useful to go beyond this and consider an infinite dimensional goods space, or even a generalized metric space of choices (in finance, for example), but we will not do so here.

A List of Mathematical Symbols and Notation

We will provide precise definitions of the various concepts and tools we develop in this book. This requires the use of mathematical notation which may not all be immediately familiar to all students. This technical box provides a summary list of these symbols, with their English interpretation. More details definitions will be provided as we use them in the chapters to follow. You may also want to have a look at.

$\forall$: For all, “for every”

$\exists$: There exists, “for some”

$\nexists$: There does not exist, “for no”

$\in$: Element of, In

$\notin$: Not an element of, “not in”

$\subset$: Strict Subset of, contained in but not equal to

$\subseteq$: Subset of, contained in or equal to

$\cup$: Union of sets

$\cap$: Intersection of sets

$\wedge$: Logical “and”

$\vee$: Logical “or”

$\Rightarrow$: Implies, “if”

$\Leftarrow$: Implied by, “only if”

$\Leftrightarrow$: If and only if

$\emptyset$: Empty set

$\mathbb{R}$: Real numbers

$\mathbb{N}$: Natural numbers

Not all of this consumption space may be permissible for consumers. For example, how would we interpret negative consumption? If this does not make sense, we might want to restrict consumers' choices to the positive orthant: $\mathbb{R}_+^N$. Alternatively, suppose that some consumption bundles do not permit the consumer to survive. Too little food or too little clothing might result in death. We might take the view that we cannot have meaningful preferences over such bundles.

Index and Vector Notation

Indices keep track of attributes of vectors or other objects. For example, which good is being consumed or produced, by which agent or firm. We will use the following conventions:

Index Set: An ordered set natural numbers, $\mathbb{N}$, which may be finite or infinite. If an index set is finite, it runs from 1 to some upper bound. We will use a lower-case letter to denote an **element of an index set**, an upper-case letter to denote the **upper bound of an index set**, and a script letter to denote the **entire index set**.

$$(1, \dots, n, \dots, N) \equiv \mathcal{N}.$$

N-Dimensional Vector: An ordered set of N real numbers. We will denote vectors with lower-case letters, the n^{th} **component of a vector** with a subscripted index, and the **set that an element is a member of** with an upper-case letter:

$$(x_1, \dots, x_n, \dots, x_N) \equiv \mathbf{x} \in \mathbf{X} \subseteq \mathbb{R}^N.$$

When we consider multiple agents producing or consuming multiple goods, we need two index sets to distinguish the elements of a vector along these two separate sets of attributes. Thus, agent i 's consumption of good n is the scalar:

$$x_{i,n} \in \mathbb{R}.$$

Here, the first subscript indicates the agent while the second indicates the good.

Dropping a subscript will indicate a vector composed whole list over the entire index set of the attribute. Thus, we indicate the consumption vector of agent i , and the vector of consumption amounts of good n for each agent, respectively:

$$\mathbf{x}_i = (x_{i,1}, \dots, x_{i,N}) \in \mathbb{R}^N \text{ and } \mathbf{x}_n = (x_{1,n}, \dots, x_{I,n}) \in \mathbb{R}^I,$$

You may notice that this creates a potential ambiguity since x_2 could be either a consumption vector for agent 2, or the consumption level of good 2 for all agents. In practice, we will refer to generic agents or goods, as $\mathbf{x}_i$ and $\mathbf{x}_n$, and so the choice of index will provide context. We will not use more complex notational forms that would prevent such ambiguities in the interest of readability.

In general, we will restrict the domain of a consumer's preferences to a **consumption set** denoted:

$$X_i \subseteq \mathbb{R}^N,$$

where the subscript refers to a particular consumer $i \in \mathcal{I}$. In principle, each consumer might have a different consumption set. We will typically need to assume that the consumption set is (*weakly*) *convex* and *bounded from below*, in order to prove welfare and existence theorems discussed later.

To keep things simple in this text, we will follow the convention that $X_i \equiv \mathbb{R}_+^N$. In words, all agents have preferences over all bundles of goods that have weakly positive amounts of each of the commodities. Of course, the consumption bundle chosen by a given consumer will typically have positive quantities of a few commodities, but zero quantities of the great majority of the millions of goods in the N-dimensional commodity space. I, for example, consume no Tesla Model S's, Bottles of 1990 Louis Roederer Cristal Brut, or vintage gold lamé jumpsuits, although I would dearly love to do so.

The figure below shows a possible consumption set in two dimensions. Note the curve defines the lower bound of the consumption set, and anything above is part of the consumption set. We do not show the entire consumption set because we ran out of paper.

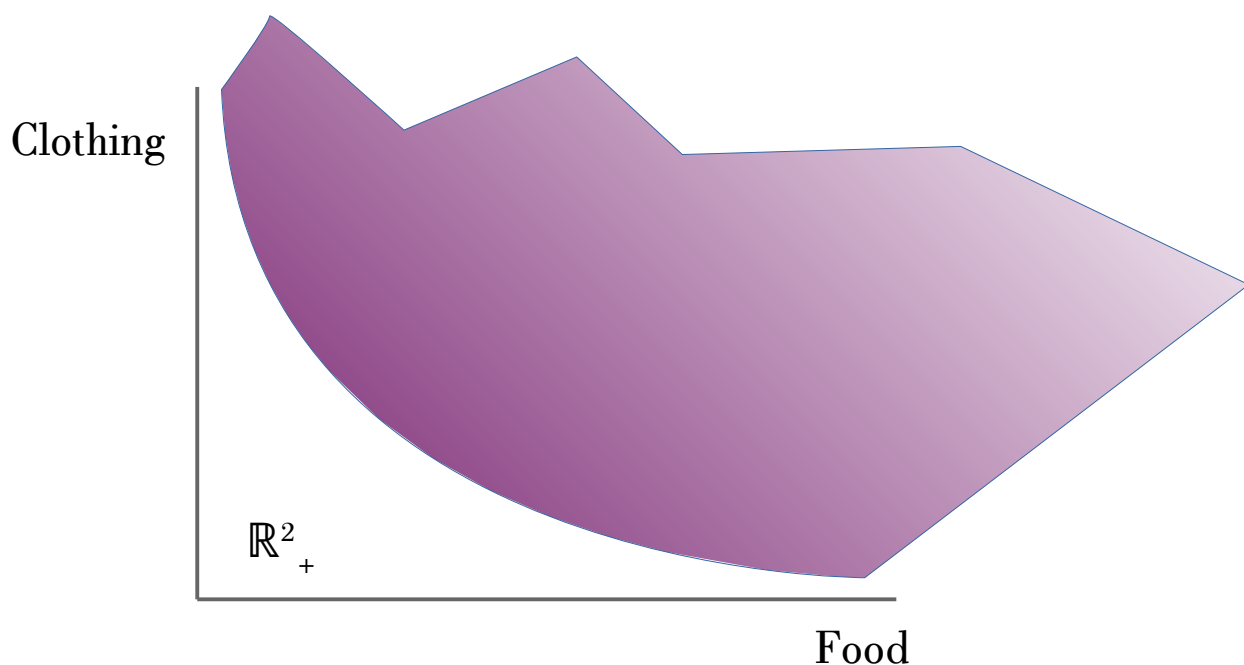


Figure 2: Example of a consumption set contained in the positive orthant

Ordered and Unordered Sets

It is conventional to distinguish ordered and unordered sets with round and curly brackets, respectively.

Ordered Set: A set for which the order of elements is part of its definition, and are which denoted by round brackets:

$$\begin{aligned}(1, \dots, N) &\subset \mathbb{N}^N \\ (x_1, \dots, x_N) &\in \mathbb{R}^N \\ (x_1, \dots, x_I) &\in \mathbb{R}^{I \times N} \\ (x^s)_{s \in \mathbb{N}} &\end{aligned}$$

Examples of ordered sets are index sets, the set of scalar components of vectors ordered by their competent index, sets of vectors ordered by agent or other index, and infinite sequences ordered in a defined, but otherwise arbitrary, way.

Unordered Sets: A set containing a collection of elements which is defined only by the elements in the collection, not by the order in which they are listed, and which are denoted by curly brackets:

$$\begin{aligned}\{a, b, c\} &= \{b, c, a\} \in S \\ \text{Spref}(x) &\equiv \{z \in X \mid z \succ x\}\end{aligned}$$

Examples of unordered sets are arbitrary collections elements, which may or may not be vectors, and abstractly defined sets. In the first case, sets that contain the same elements are equivalent, regardless of the order of the elements. In the latter case, there may be no meaningful way, or need, to order the elements.

Section 2.2. Preferences

Now consider any two commodity bundles in agent consumption set: $x, \bar{x} \in X_i$. We will use the symbols, $\succ_i, \succsim_i, \sim_i$, to denote an agent's **preferences**.

Weak Preference: If x is either strictly preferred to or exactly as good as $\bar{x}$ (equivalently, at least as good as $\bar{x}$), then we write: $x \succsim_i \bar{x}$.

From the weak preference relation, we define the strict preference relation as follows:

Strict Preference: $x \succsim_i \bar{x}$ and $\bar{x} \not\succsim_i x \Rightarrow x \succ_i \bar{x}$.

In words:

if x is at least as good as $\bar{x}$ and $\bar{x}$ is not at least as good as x
then
 x is **strictly (or strongly) preferred** to $\bar{x}$.

Also note that: $x \succ_i \bar{x} \Rightarrow x \succsim_i \bar{x}$ and $\bar{x} \not\succsim_i x$.

In **formal logic**, $x \succsim_i \bar{x}$ and $\bar{x} \not\succsim_i x$ are called **statements**, and have a **truth value** which must either be **true** or **false**. If we write $\{x \succsim_i \bar{x}\} \wedge \{\bar{x} \not\succsim_i x\}$ we make a **compound statement** using the **logical “and”**. For this compound statement to have a value of true, both of the **sub-statements** must be true. The complete set of statements gives a **definition**. That is, the compound statement on the left is logically identical to the statement on the right. In other words, the left-side statements are jointly true if and only if the last right-side statement is true. Equivalently the left-side compound statement is false if and only if the right-side statement is false. You may wish to have a look at [Appendix C. Logic and Proof](#) which provides a more detailed review.

We also define the indifference relation from the weak preference relation:

Indifference: $x \succsim_i \bar{x}$ and $\bar{x} \succsim_i x \Leftrightarrow x \sim_i \bar{x}$.

In words:

x is at least as good as $\bar{x}$ and $\bar{x}$ is at least as good as x
if and only if
 $\bar{x}$ is **exactly as good** as x .

We will also say that agent i is **indifferent** between x and $\bar{x}$. People often say that “ x is indifferent to $\bar{x}$.” This is not correct, of course. Bundles of goods like x and $\bar{x}$ don not have any feeling about one another, but this shorthand conveys the idea that an agent is indifferent between two alternatives.

Formally, these three symbols denote a **binary relation** over the consumption set, similar to $\geq$ which partially orders $\mathbb{R}^N$. Because of this, we will sometimes refer to these symbols as denoting an agent's **preference relation**. If you look carefully, the preference relation symbols are curved at the ends instead of straight like the symbols “greater than or equal to” or “greater than” relation.

Section 2.3. Properties of Preferences

In the abstract, a consumer's preferences could rank alternative consumption bundles in any way at all. In real life, however, consumers seem to follow certain regularities in their preferences over commodities (at least on the average). This is fortunate since what seems to be empirically true about preferences allows us to demonstrate several important things about markets and aggregate behavior. We summarize these properties in five formal assumptions and two variants, below.

These assumptions are sometimes referred to as “axioms of rational behavior”. This is not really accurate, however. Consumers who do not follow these axioms are not truly “irrational”. For the purposes of economic modeling, “rationality” simply means maximizing objectives within constraints. This says nothing about the nature of the objectives (as the axioms explicitly do).

Although it is quite possible to analyze consumers who do not satisfy the axioms we discuss below, it would not be particularly useful to do so. Our primary interest is to understand policy questions and consumer behavior in realistic situations and so not much would be gained by studying a more general setting that does not seem to have much real world relevance. Because we will impose these assumptions on all agents, we will drop the subscript denoting agents from the preference relation and the consumption set for the remainder of this section.

1. Completeness: $\forall x, \bar{x} \in X, x \succsim \bar{x} \text{ or } \bar{x} \succsim x \text{ (or both)}.$

In words, any pair of consumption bundles in the consumption set can be ranked against one another by the weak preference relation. Either the first bundle is at least as good as the second, OR the second is at least as good as the first (or both, in which case they would be equally good). This means there will never arise a situation in which a consumer cannot choose between two alternatives. For example, a consumer might prefer two slices of pizza and one beer, to one hamburger and two beers, or the reverse. He might even think they are equally good. What we don't allow for is an agent to throw up his hands and say that the two bundles are so different that he simply can't compare them, and neither one is at least as good as the other.

This is an example of a compound statement using the **logical “or”**. For the compound statement $\{x \succsim \bar{x}\} \vee \{\bar{x} \succsim x\}$ for all $x, \bar{x} \in X$, it must be true for any pair of consumption bundles we choose, at least one of these two statements is true. It need not be the same statement for each pair of bundles, and may also be the case that both statements are true. The complete compound statement defines a property of the weak preference relation. If the statement is true for any given preference relation, then the relation is said to have the property of “completeness”.

2. Transitivity: $\forall x, \bar{x}, \hat{x} \in X, \text{ if } x \succsim \bar{x} \text{ and } \bar{x} \succsim \hat{x} \text{ then } x \succsim \hat{x}.$

If you like BMWs at least as much as Hondas and Hondas at least as much as Fiats, it would be very strange indeed did not think that BMWs are at least as good as Fiats. This kind of circular preference would make it impossible for an agent to identify an optimal choice. Such preferences

are unlikely to be observed in the real world. (Still, there is the story of the king who had three daughters, each of whom was more beautiful than the other ...)

Transitivity can be shown to hold for the indifference and strong preference relation as well using the assumptions we make below.

Vector Inequalities

We distinguish three types of vector inequality or vector ordering.

Consider two vectors, $x, y \in \mathbb{R}^N$:

Strictly Greater than: $x \gg y \Leftrightarrow \{ \forall n \in \mathcal{N}, x_n > y_n \}$.

for example: $(2,3,7) \gg (1,0,6)$.

Greater than: $x > y \Leftrightarrow \{ \forall n \in \mathcal{N}, x_n \geq y_n \text{ and } \exists m \in \mathcal{N} \text{ such that } x_m > y_m \}$.

for example: $(2,3,7) > (1,0,6)$ and $(2,3,7) > (2,3,6)$.

Greater than or Equal to: $x \geq y \Leftrightarrow \{ \forall n \in \mathcal{N}, x_n \geq y_n \}$.

for example: $(2,3,7) \geq (1,0,6)$, $(2,3,7) \geq (2,3,6)$, and $(2,3,7) \geq (2,3,7)$.

Of course, it may be that none of these apply. For example: $(2,3,7)$ and $(1,8,6)$ are not vector ordered by any of the above inequalities.

If the commodities in the consumption set are goods, then larger consumption bundles should be preferred to smaller ones. If the commodities are bads, on the other hand, a bundle with less of each commodity would be preferred. A consumer would probably like another beer, but would not want the contents of the ashtrays in the bar.

It is also possible that agents might become **satiated** in some or all of the commodities. Giving agents more or less of such commodities leaves agents neither better nor worse off. For example, you are currently breathing all the air you want to. If I gave you the right to breathe twice as much as you do currently, you would not change your behavior and so would not experience any change in your well-being.

Since economics is often described as the study of:

ALLOCATING SCARCE RESOURCES OVER COMPETING ENDS

We typically assume that all commodities are goods, and that agents are not satiated in any of them. Formally:

3A. Strong Monotonicity: $\forall x, \bar{x} \in X$ such that $x > \bar{x}$, it holds that $x \succ \bar{x}$.

The idea of strong monotonicity is that agents are made strictly better off if you move their consumption bundle in a positive direction. More specifically, they are made strictly better off if you give them strictly more of at least one good while keeping them at least at the same level of consumption of all other goods. Note that we use the terms *strong* and *strict* interchangeably.

3B. Weak Monotonicity: $\forall x, \bar{x} \in X$ such that $x \geq \bar{x}$, it holds that $x \succcurlyeq \bar{x}$. In addition, $\forall x, \bar{x} \in X$ such that $x \gg \bar{x}$, it holds that $x \succ \bar{x}$.

This weaker version of monotonicity allows the possibility that agents might be satiated in some goods, but not all goods at once. This definition says that if you gave an agent strictly more of every good, they would be strictly better off, while if you gave him strictly more of only some goods and kept him at the same consumption level of the others, you would certainly not make him worse off, although you might or might not make him better off.

Note that if add the assumption of continuity (defined below), the second sentence in the definition of weak monotonicity $\{\forall x, \bar{x} \in X, x \gg \bar{x} \Rightarrow x \succ \bar{x}\}$ is implied by the first $\{\forall x, \bar{x} \in X, x \geq \bar{x} \Rightarrow x \succcurlyeq \bar{x}\}$ and so is sometimes omitted.

Clearly, both of these monotonicity assumptions only apply to goods. If we wanted to consider a consumption space of bads, we would need to reverse the preference relation in each statement. That is, more bads would make us worse off in the appropriate sense rather than better off. We could extend this to a mixed space of goods and bads as well by requiring that agents prefer consumption vectors with more of the goods and less of the bads.

To illustrate:

- Bundle B has more of everything than bundle A (that is $B \gg A$). Thus:
 - S. Mon $\Rightarrow B \succ A$.
 - W. Mon $\Rightarrow B \succ A$.
- Bundle C has more hot dogs, but the same amount of coke than bundle A (that is, $B \geq A$). Thus:
 - S. Mon $\Rightarrow C \succ A$.
 - W. Mon $\Rightarrow B \succcurlyeq A$.
- Bundle D has more cokes, but fewer hot dogs than bundle A ($D ? A$). Thus:
 - S. Mon $\Rightarrow D ? A$.
 - W. Mon $\Rightarrow D ? A$.

Neither strong nor weak monotonicity tells us anything about which bundle is preferred to another when the bundles are not vector ordered, as in the case of A and D.

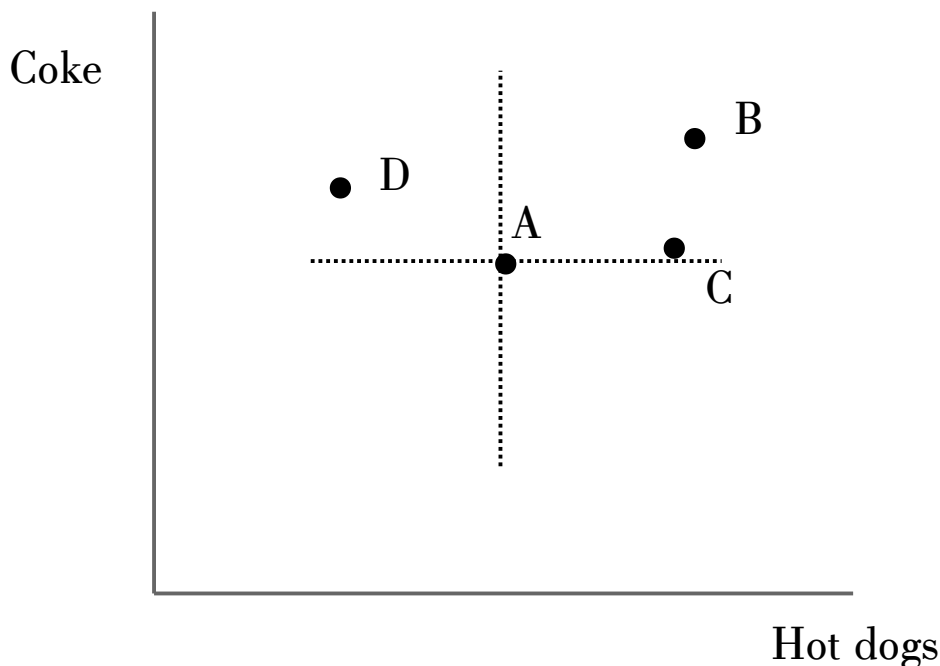


Figure 3: Strong and weak monotonicity

While monotonicity assumptions are both intuitive and convenient, they are really more than we need. Versions of essentially all of our results go through provided that agents are never satiated in all goods at the same time. That is, provided that no matter where an agent is in the consumption set, there is always another bundle in every neighborhood (no matter how small) that he strongly prefers. This preferred bundle does not need to be in either a positive or negative direction, and so in particular, this allows for the possibility that commodities may go from being goods to bads and back again at different levels of consumption. Formally:

3C. Local Nonsatiation: $\forall x \in X$, and $\forall \varepsilon \in \mathbb{R}_+$, $\exists \bar{x} \in X$ such that $\|x - \bar{x}\| < \varepsilon$ and $\bar{x} \succ x$.

The Euclidean Norm

Most of what we will do in this book takes place in $\mathbb{R}^N$. Informally, this is based on an **N-dimensional coordinate** space of real-valued vectors. If $n = 2$, this is the 2-dimensional coordinate plane that we all were introduced to sometime in middle-school. More formally, this is an **N-dimensional real vector space** endowed with the Euclidean distance metric, a linear structure and an inner product operation.

The **Euclidean metric** is the distance measure for Euclidean spaces and is defined as:

$$d(x, y) \equiv \|x - y\| \equiv \sqrt{\sum_{n \in \mathcal{N}} (x_n - y_n)^2}.$$

This is sometimes called the **Euclidean norm**. More generally, given an arbitrary space, S , a **metric** is a measure of distance, $d: S \times S \Rightarrow \mathbb{R}$, that satisfies four conditions $\forall x, y, z \in S$:

Identity Condition: $d(x, y) = 0 \Leftrightarrow x = y$

Symmetry: $d(x, y) = d(y, x)$

Nonnegativity: $d(x, y) \geq 0 \Leftarrow x \neq y$

Triangle Inequality: $d(x, y) + d(y, z) \geq d(x, z)$

The Euclidean metric measures the distance between two points in a real space, and you can easily see that walking from point A to point B on the straight line connecting them is the quickest path. If you take a detour to include a third point C, the trip will take longer unless C happens to be on the line between A and B, in which case, it will take the same amount of time (this is the triangle inequality).

Monotonicity assumptions tell us something about the preferences of agents for consumption bundles that are larger or smaller. What can we say about bundles that are larger in some, but smaller in other dimensions? In particular, the bundles above a bundle A are preferred while the bundles below it are inferior. What about the regions marked “?” in the next figure below?

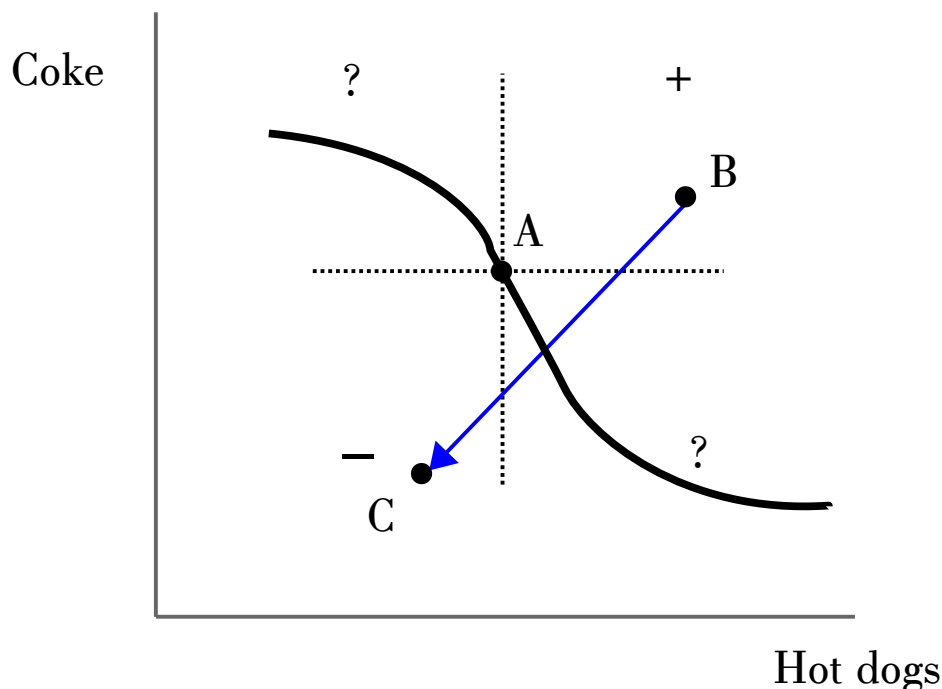


Figure 4: Continuity of Preferences

Consider a bundle $\hat{x}$ which is above, and therefore preferred to A, and a bundle $\bar{x}$ which is below, and therefore inferior to A. Connect these two bundles with a line and suppose we start at B and head in the direction of C. At first, we find bundles that are still preferred to $\hat{x}$ along this line, but as we get closer to C, we eventually find bundles that are inferior. By monotonicity, once we encounter a bundle that is inferior to $\bar{x}$ all the bundles below this as we continue to head toward C must also be inferior. Thus, there is a boundary between bundles that are strictly better and strictly worse than A.

It would be strange if we could go from a bundle that is strictly preferred to A to one that is strictly inferior without encountering a bundle in between that is exactly as good as $\hat{x}$. Such jumpy behavior would be discontinuous and violates the idea that similar consumption bundles are similarly valued. We will therefore assume that this never happens. Instead, we will assume that as we go from a preferred bundle to an inferior bundle, we must pass through an indifferent bundle. This assumption will also imply that in every neighborhood of a bundle, we will be able to find other bundles that are exactly as good.

Suppose we collected all the bundles that were indifferent to a given bundle $x \in X$ and connected them up. This would give a curve (or a surface, if the consumption set had more than two commodities). Informally, an **indifference curve** is a set of bundles in the consumption set that are exactly as good as one another under the preferences of a given agent. That is, two bundles x and $\bar{x}$, are on the same indifference curve if and only if $x \sim \bar{x}$.

Vector Mathematics and Linear Metric Spaces

Recall that $\mathbb{R}^N$ is a **linear metric space**. This means that certain addition and multiplication operations are well defined and satisfy certain properties.

Vector Addition: $x + y = (x_1 + y_1, x_2 + y_2, \dots, x_N + y_N)$

Scalar Multiplication: $\alpha x = (\alpha x_1, \alpha x_2, \dots, \alpha x_N)$

Inner Product or Dot Product: $x \cdot y \equiv \sum_{n \in \mathcal{N}} x_n y_n$

The **inner product** or **dot product operation** for vectors comes up a lot in economic contexts. Let $p \in \mathbb{R}^N$ be a price vector, that is, vector that gives the price of each of the N goods in the economy. Similarly, let $x \in \mathbb{R}^N$ be a consumption vector. Then the dot product of these two vectors is defined as follows:

$$p \cdot x = p_1 x_1 + p_2 x_2 + \dots + p_N x_N \equiv \sum_{n \in \mathcal{N}} p_n x_n$$

which is simply the total cost of bundle x under prices p . If a space is **linear**, the following axioms hold for these operations:

Commutativity of Vector Addition: $x + (y + z) = (x + y) + z$

Associativity of Vector Addition: $x + y = y + x$

Associativity of Scalar Multiplication: $\alpha(\beta x) = (\alpha\beta)x$

Distributivity of Scalar Addition: $(\alpha + \beta)x = \alpha x + \beta x$

Distributivity of Vector Sums: $\alpha(x + y) = \alpha x + \alpha y$

Scalar Multiplication Identity: $1(x) = x$

Existence of an Additive Inverse: $x + (-x) = 0 \equiv (0, \dots, 0)$

Additive Identity (existence of a zero): $x + 0 = 0 + x = x$

Under monotonicity, all the bundles above an indifference curve though x are strictly better than x and all bundles below this indifference curve are strictly worse than x . This gives us the notion of the **strongly** and **weakly preferred sets**:

$$Spref(x) \equiv \{z \in X \mid z \succ x\}.$$

and

$$Wpref(x) \equiv \{z \in X \mid z \succsim x\}.$$

Recall that a set in $\mathbb{R}^N$ is *closed* if it includes its boundary where it is bounded. Given this, a formal definition of an **indifference curve** is:

$$IC(x) \equiv \{z \in X \mid z \sim x\} \equiv \{z \in X \mid z \in Wpref(x) \text{ and } z \notin Spref(x)\}$$

That is, the set of bundles that are exactly as good as some bundle x is defined as the set of bundles that are at least as good as x but not strictly better than x .

The assumption of continuity itself can be defined in two completely equivalent ways:

4. Continuity: $\forall x \in X$, $Wpref(x)$ is closed.

4. Continuity: $\forall x \in X$, and $\forall \{x^s\}_{s \in \mathbb{N}} \subset Wpref(x)$, if $\lim_{s \rightarrow \infty} x^s = \bar{x}$, then $\bar{x} \in Wpref(x)$.

It is easy to see that under continuity, there is an indifference curve through every bundle in the consumption set, and that all of these curves are continuous.

Sequences, Convergence, and Limits

Sequence: A sequence in $\mathbb{R}^N$ is an infinite, ordered, set of points, $\{x^s\}_{s \in \mathbb{N}}$, such that $\forall s \in \mathbb{N}$, $x^s \in \mathbb{R}^N$.

Convergence: A sequence in $\mathbb{R}^N$, **converges** to $x \in \mathbb{R}^N$ if $\forall \varepsilon > 0$, $\exists \hat{s} \in \mathbb{N}$ such that $\forall t > \hat{s}$ it holds that $\|x^t - \bar{x}\| < \varepsilon$.

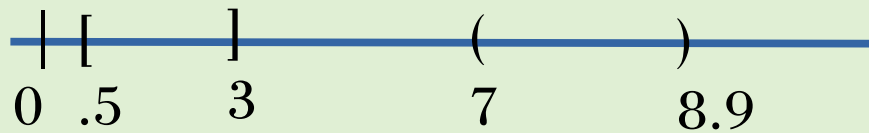
Limit: $\bar{x}$ is the **limit** of a sequence $\{x^s\}_{s \in \mathbb{N}}$ if the sequence converges to $\bar{x}$, denoted:

$$\lim_{s \rightarrow \infty} x^s = \bar{x} \text{ or } x^s \rightarrow \bar{x}.$$

The second definition of continuity above is a precise statement of what it means for the set $Wpref(x)$ to be closed.

Examples of Sequences, Convergence, and Open and Closed Sets

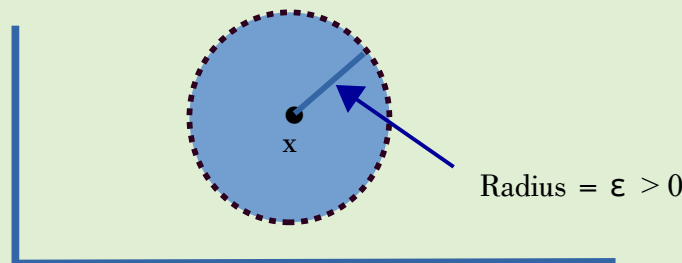
A **closed interval** in the real line, $\mathbb{R}$, is a continuous segment that includes its endpoints, and is usually indicated by square brackets: $[.5, 3]$. If both endpoints are excluded, then we have an open interval, usually indicated by parenthesis: $(7, 8.9)$.



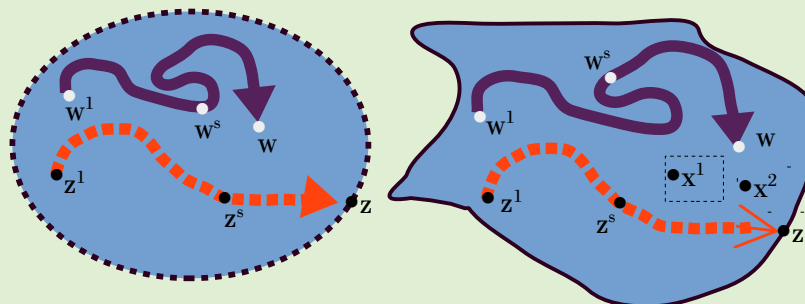
An **open epsilon ball around a point** $x \in \mathbb{R}^2$ consists of vectors that are strictly closer than epsilon to the vector x as measured by Euclidean distance:

$$B_\epsilon(x) \equiv \{z \in \mathbb{R}^N \mid \|x - z\| < \epsilon\}.$$

Note that the boundary is not included, and that epsilon can be large or small. We sometimes say “a proposition is true for every **neighborhood of x**”. By this we mean that the proposition is true for every open epsilon ball around x , no matter how small epsilon gets. In three dimensions, an epsilon ball is really a sphere. In four dimensions or more, it is called a “hyper-sphere,” but this is more than I can envision.



Below are two sets, each with two sequences of vectors contained in the sets, and represented by arrows. The set on the right side is closed. This is because the limit of every **convergent sequence** of vectors is contained in the set. The set of the left is not closed. While the purple sequence (top) converges to a point w contained in the set, the red sequence converges to z which is not in the set.



Note that not all sequences converge. For example, the sequence $(x^1, x^2, x^1, x^2, \dots)$ is an oscillating sequence. No matter how far out you go in this sequence, you will always be outside of a small neighborhood of either point half of the time.

The ideas of open sets, closed sets, boundary points, and interior points in Euclidean spaces like $\mathbb{R}^N$ are both fundamental, and useful in understanding Economics.

Boundary and Interior Points of Sets in $\mathbb{R}^N$

Epsilon Ball: An open epsilon ball around a point, $x \in \mathbb{R}^N$

$$B_\varepsilon(x) \equiv \{z \in \mathbb{R}^N \mid \|x - z\| < \varepsilon\}.$$

We will also call this an **open neighborhood of x** or simply a **neighborhood of x**.

Interior: $x \in S \subseteq \mathbb{R}^N$ is an **interior point** of a set S if $\exists \varepsilon > 0$ such that $B_\varepsilon(x) \subseteq S$:

$$\text{interior}(S) \equiv \{x \in S \mid \exists \varepsilon > 0 \text{ and } B_\varepsilon(x) \subseteq S\}.$$

Boundary: $x \in S \subseteq \mathbb{R}^N$ is a **boundary point** of S if $\forall \varepsilon > 0, \exists y \in B_\varepsilon(x)$ such that $y \notin S$:

$$\text{boundary}(S) \equiv \{x \in S \mid \forall \varepsilon > 0, \exists y \in B_\varepsilon(x) \text{ and } y \notin S\}.$$

Open Set: $S \subseteq \mathbb{R}^N$ is an **open set** if $S = \text{interior}(S)$.

Closed Set: $S \subseteq \mathbb{R}^N$ is a **closed set** if $\text{boundary}(S) \subseteq S$.

Complement of a Set: The **complement** of a set S relative to $\mathbb{R}^N$ (or any $T \subseteq \mathbb{R}^N$) consists of all elements of $\mathbb{R}^N$ (or T) that are not also in S :

$$\text{complement}(S) \equiv \{x \in \mathbb{R}^N \mid x \notin S\}.$$

Bounded: A set $S \subseteq \mathbb{R}^N$ is **bounded** if $\exists \bar{\beta} \in \mathbb{R}$ such that $\forall x \in S, \|x\| < \bar{\beta}$.

Bounded Above: A set $S \subseteq \mathbb{R}^N$ is **bounded from above** if $\exists \bar{z} \in \mathbb{R}^N$ such that $\forall x \in S, x \leq \bar{z}$

Bounded Below: A set $S \subseteq \mathbb{R}^N$ is **bounded from below** if $\exists \bar{z} \in \mathbb{R}^N$ such that $\forall x \in S, x \geq \bar{z}$.

Compactness: A set $S \subseteq \mathbb{R}^N$ is **compact** if it is both *closed* and *bounded*.

Set Subtraction: The **set subtraction** of $T \subseteq \mathbb{R}^N$ from $S \subseteq \mathbb{R}^N$ removes any elements in T from S and is denoted: “/”.

$$S/T \equiv \{x \in S \mid x \notin T\} \equiv \{x \in S \mid x \notin S \cap T\}.$$

If not elements in T are also in S then $S/T = S$.

A few fun facts about open and closed sets in $\mathbb{R}^N$:

- The union of an arbitrary number of open sets is also an open set.
- The intersection of a finite number of closed sets is also a closed set.
- The set of all open sets in $\mathbb{R}^N$ is defined as follows:

$$\mathcal{O} \equiv \{ S \subseteq \mathbb{R}^N \mid \forall x \in S, \exists \varepsilon > 0 \text{ and } B_\varepsilon(x) \subseteq S \}.$$

- The interior of a set S is equal the union of all open sets contained in S :

$$\text{interior}(S) \equiv \bigcup_{T \in \mathcal{O}} T,$$

or, equivalently:

$$\text{interior}(S) \equiv \{ x \in S \mid \exists T \subseteq S, \text{ such that } x \in T \text{ and } T \text{ is an open set} \}.$$

- The boundary of a set S is equal to the intersection of its closure and the closure of its complement:

$$\text{boundary}(S) \equiv \text{closure}(S) \cap \text{closure}(\text{complement}(S)).$$

- The closure of a set S is equal to the intersection of all closed sets containing S .
- A closed set is the complement of an open set.
- Both closed and open sets can be unbounded. For example, the positive orthant without the axes is an open set, while if we add the axes and the origin, it becomes a closed set.
- The empty set ($\emptyset$) and the entire space ($\mathbb{R}^N$) are technically both open and closed.

In the figure below we see an example of a possible indifference curve. All the bundles strictly above any bundle on the curve through A are strictly preferred to A . All the bundles weakly above any bundle on the curve through A are weakly preferred to A . The difference between the weakly and strongly preferred sets to A is the line shown in the figure.

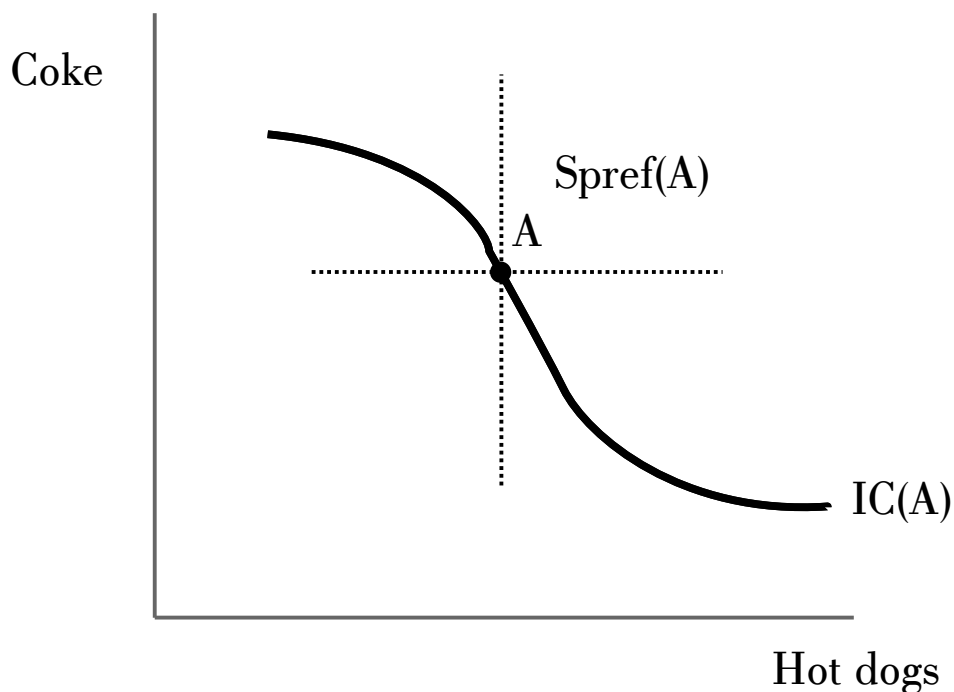


Figure 5: The set of bundles indifferent to bundle A

Now that we know what indifference curves are, we would like to know something about their shape. Consider the following questions:

- Which would you rather have, four slices of pizza only, four glasses of beer only, or two slices and two glasses?
- Would you prefer to own two shirts, or two pairs of pants, or one of each?
- Suppose you have not eaten. How much would you enjoy eating your first chicken wing? How about your second, your third, your 12th, your 30th?

In general, we find two equivalent things:

AVERAGE BUNDLES ARE PREFERRED TO EXTREME BUNDLES.

THE MARGINAL VALUE OF A GOOD DECLINES THE MORE YOU CONSUME.

A formal way of saying this is to assume that preferences are convex:

5A. Strong Convexity: $\forall \lambda \in (0,1), \forall x \in X$ and $\forall \bar{x}, \hat{x} \in Wpref(x)$ such that $\bar{x} \neq \hat{x}$:

$$\lambda \bar{x} + (1-\lambda) \hat{x} \in Spref(x) .$$

Strongly convex preferences have the property that all non-trivial linear combinations of two consumptions bundles (a weighted average where $\lambda \neq 0$ and $\lambda \neq 1$) are strictly preferred to the extremes. Note that strongly convex preferences are also weakly convex, but weakly convex preferences may not be strongly convex.

Strong convexity implies that none of the indifference curves can have any “flat” spots. To see this, take any pair of bundles on a flat part of an indifference curve, and draw a line between them. Clearly, this line is entirely within the flat section between the two bundles. In other words, the line stays on the boundary of the weakly preferred set rather than going into the interior. These weighted averages are therefore not in the strongly preferred set.

This immediately implies that if the preferences of an agent with strongly convex preferences are monotonic at all, they must be strongly monotonic. If they were only weakly monotonic and agents were satiated in some goods, then the indifference curves would have horizontal or vertical parts (both of which would be flat).

If we want to allow for the possibility that agents are satiated in some goods, we have to weaken the convexity assumption as follows:

5B. Weak Convexity: $\forall \lambda \in [0, 1], \forall x \in X$ and $\forall \bar{x}, \hat{x} \in W_{\text{pref}}(x)$
 $\lambda \bar{x} + [1 - \lambda] \hat{x} \in W_{\text{pref}}(x) .$

Weakly convex preferences have the property that averages are at least as good as extremes. Note that strongly convex preferences are also weakly convex, but weakly convex preferences may not be strongly convex.

More generally in mathematics, a (weakly) convex set in a Euclidean space is defined as follows:

Convex Set: $S \subseteq \mathbb{R}^N$ is **convex** if $\forall \lambda \in [0, 1], \forall x, \bar{x} \in \mathbb{R}^N$ it holds that $\lambda x + (1 - \lambda)\bar{x} \in S$.

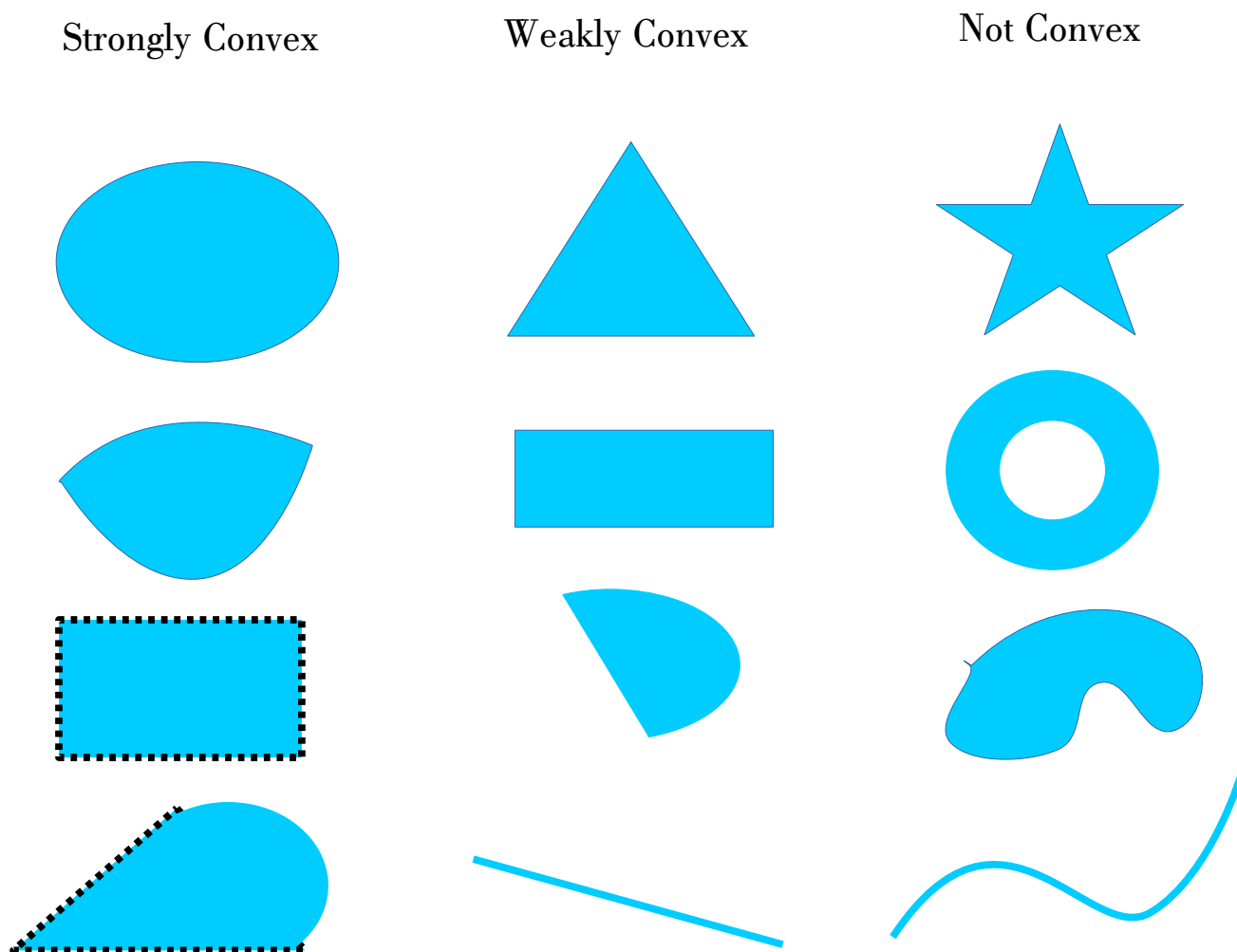


Figure 6: Examples of convex and nonconvex sets

In the figures the first column are examples of strongly convex two-dimensional shapes, or sets. Notice that if I choose any two points in one of these shapes and connect them with a line, the interior of the line is in the interior of the shape. The bottom two examples are cases of sets that do not include some or any of their boundaries. Even though these would not be strongly convex sets if the boundary was included, when we exclude the boundary, you can verify that they become strongly convex.

The second column are examples of shapes that are weakly, but not strongly convex. Notice that if I choose any two points in one of these shapes and connect them with a line, the line stays completely within the set. It may be that the line includes boundary points, but it never leaves the set. Of course, all the strongly convex shapes shown are also weakly convex.

Finally, the third column gives examples of sets that are not convex. Notice that it is possible to choose two points in any one of these shapes and connect them with a line, but have that line go outside the set.

Some Properties of Convex Sets

Extreme Point: A vector $x \in \mathbb{R}^N$ is an **extreme point** of a convex set $S \subseteq \mathbb{R}^N$ if it cannot be expressed as $x = \lambda y + (1-\lambda)z$ for any $y, z \in S$ and $\lambda \in (0, 1)$.

Carathéodory's Theorem: Let $S \subseteq \mathbb{R}^N$ be a convex and compact set. Then every $x \in S$ can be expressed as a convex combination of at most $N + 1$ extreme points.

(N-1)-Dimensional Unit Simplex: The **simplex** in $\mathbb{R}^N$ is defined as follows:

$$\Delta^{N-1} \equiv \{p \in \mathbb{R}_+^N \mid \sum_{n \in \mathcal{N}} p_n = 1\}.$$

Convex Hull:

$$\text{con}(S) \equiv \{z \in \mathbb{R}^N \mid \exists x_1, \dots, x_N \in S, \text{ and } \lambda \in \Delta^{N-1} \text{ such that } z = \sum_{n \in \mathcal{N}} \lambda_n x_n\}.$$

It turns out that the convex hull of a set is the smallest convex set that contains original set.

Another key idea related to the convexity of indifference curves is the **marginal rate of substitution between two goods** x_1 and x_2 , (MRS). Informally:

$$\text{MRS}_{1,2} = \left| \frac{\Delta x_1}{\Delta x_2} \right|_{\text{Agents are just as well off}}$$

For example, suppose you start out with some consumption bundle, say five apples and five bananas. Next I take one of your bananas away. This is sad. You end up on a lower indifference curve. Now I feel bad. I decide I want to make it up to you.

Unfortunately, I have already eaten the banana. Instead, I start giving you slices of apple until I see that you are just as happy as you were before my piece of social engineering. Perhaps I have to give you 1.5 apples to do so. Thus, if I take away 1 banana but give you 1.5 apples, you end up on the same indifference curve.

In other words, starting from the bundle (5, 5), 1.5 apples substitutes perfectly in your preferences for that last banana. This is why we say that your MRS of apples for bananas is 1.5, the ratio of these two numbers. Now how does this relate to convexity?

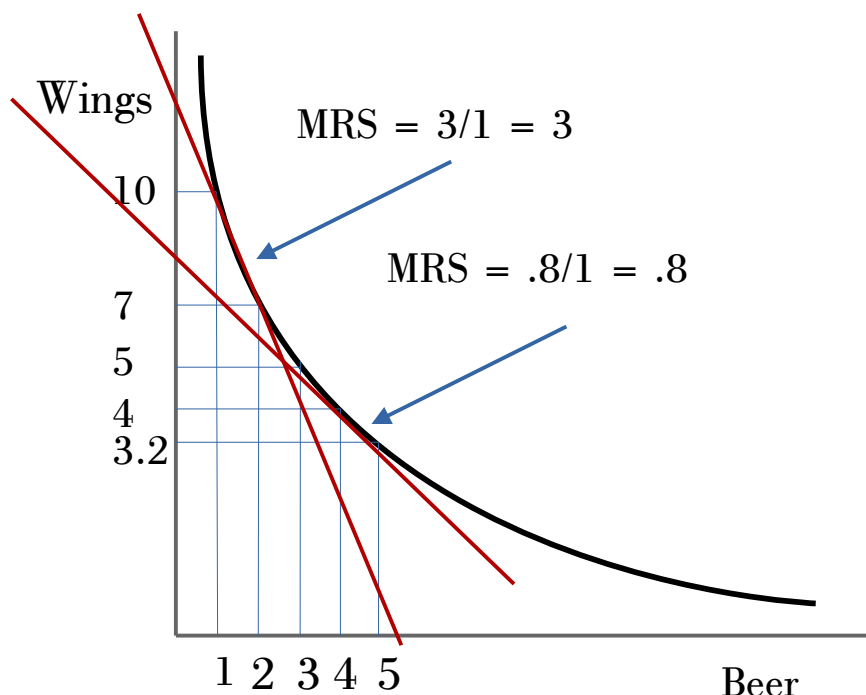


Figure 7: Convex preference and diminishing marginal rates of substitution

Consider the figure above which shows one of the indifference curves over beer and wings for some agent. Initially, the agent is offered 10 wings and 1 beer. Consuming this bundle leaves him full but thirsty. How many wings do you think he might give up in order to have one more beer? You can see that 7 wings and 2 beers is on the same indifference curve. Thus, he would be willing to substitute 3 wings for 1 beer. Thus, the marginal rate of substitution equals 3.

Now starting from this bundle, how many wings would he be willing to give up to have a third beer? You can see that the answer is 2 wings. In general, the indifference curve shown is consistent with the idea that more glasses of beer an agent drinks, the fewer wings he is willing to give up in exchange for yet another beer. In other words, the agent's MRS diminishes as he consumes more of a good. This is a direct implication of the convexity of the indifference curves.

For example, if the agent has 4 beers, and 4 wings, he is only willing to give up .8 beers for an additional win, and so $MRS = .8$.

Having said all this, we can easily imagine situations in which indifference curves are not convex. For example, white wine and red wine are both nice, but you would not want them mixed together in the same glass.

Consider the figure below. You might prefer bundle B which has one glass of white wine, or bundle A which has one glass of red wine, to bundle A which has 1/2 of a glass of red and 1/2 of a glass of white wine mixed together. Even though all three bundles give you one glass of wine, both of the extreme bundles (only white or red) are preferred to the average bundle. Note, however, that you still prefer more wine to less wine, and so the lack of convexity does not affect the monotonicity of your preferences.

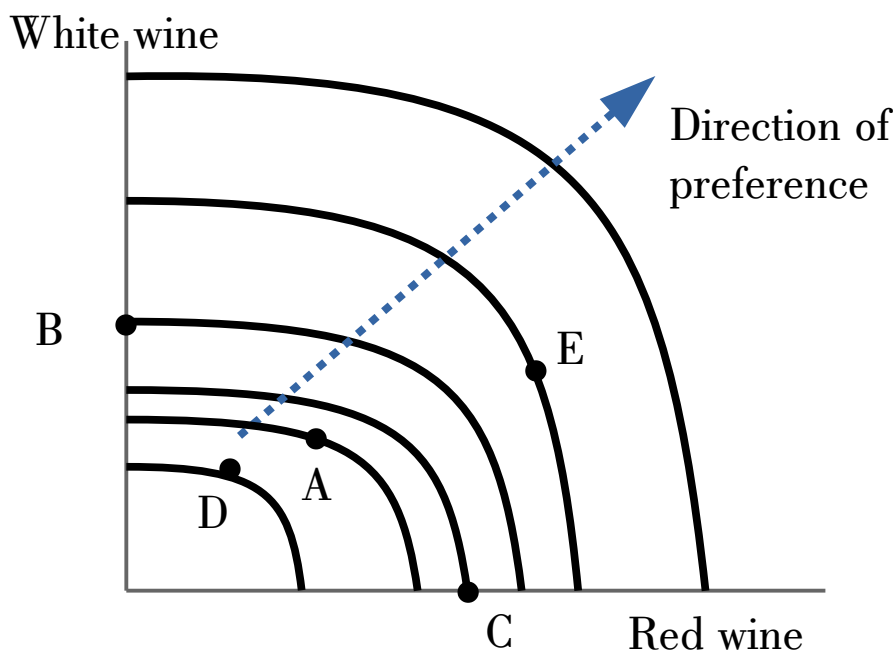


Figure 8: An example of nonconvex preferences

A similar example might be seasons of various shows on Netflix. Would you rather watch all seasons of two shows or half of each of the season four shows? The second choice would give you more hours of content, but most people would prefer to finish what they started even at the cost of having fewer total shows to watch.

Given this example, why do we think that convexity is a good assumption? There are three basic reasons.

- From an empirical standpoint, convexity is probably satisfied by most agents for most goods.
- Even if an agent has non-convex tastes at any given instant, over longer periods, preferences tend to become more convex. Even if you prefer only white or only red wine on any given night, over the course of a year, you prefer to drink some of each. Similarly, if there were an infinite number of seasons of Game of Thrones, eventually, the marginal utility would diminish enough that you would choose not to continue to watch and find a different show instead.
- Even if all individual consumers in an economy happen to have nonconvex preferences, large groups of different non-convex consumers seem to average each other out. That is, their collective behavior is similar to what we would expect from a group of consumers with convex

preferences. In fact, as the number of consumers increase, convexity becomes less and less key to our conclusions.

Section 2.4. Indifference Curves

Continuity implies that we can represent a consumer's preferences graphically with indifference curves. All the bundles on a given indifference curve are equally good. All the bundles on a higher indifference curve are strictly better. All the bundles on lower indifference curves are strictly inferior. (These three statements are implied by transitivity and monotonicity). There is an indifference curve through every point in the consumption set (given completeness) and every indifference curve is continuous (given continuity). Let's spend a little time exploring this useful tool.

First, can indifference curves cross? No! Here is the proof:

Suppose that two indifference curves cross at some bundle A. Take some bundle B on the first indifference curve and some other bundle C on the second indifference curve. Since, B and C are on different indifference curves, one must be strictly better than the other. Without loss of generality, suppose that $B \succ C$. Since A and B are on the same indifference curve, $B \sim A$. Since A and C are on the same indifference curve, $C \sim A$. Then by transitivity (and continuity), $B \sim C$, contradicting the hypothesis that $B \succ C$.

The figure below illustrates:

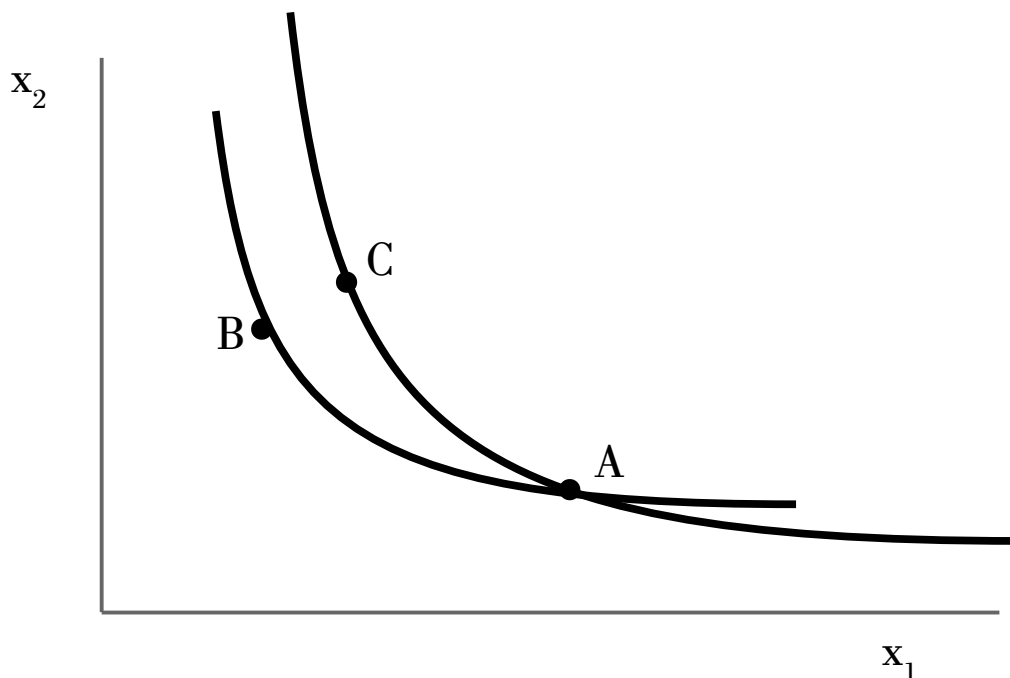


Figure 9: Indifference curves cannot cross

Below, we give examples of several important classes of indifference curves.

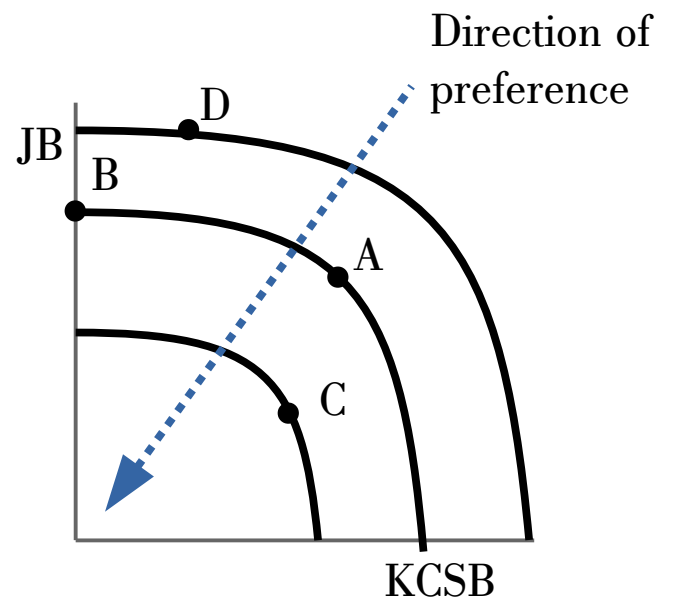
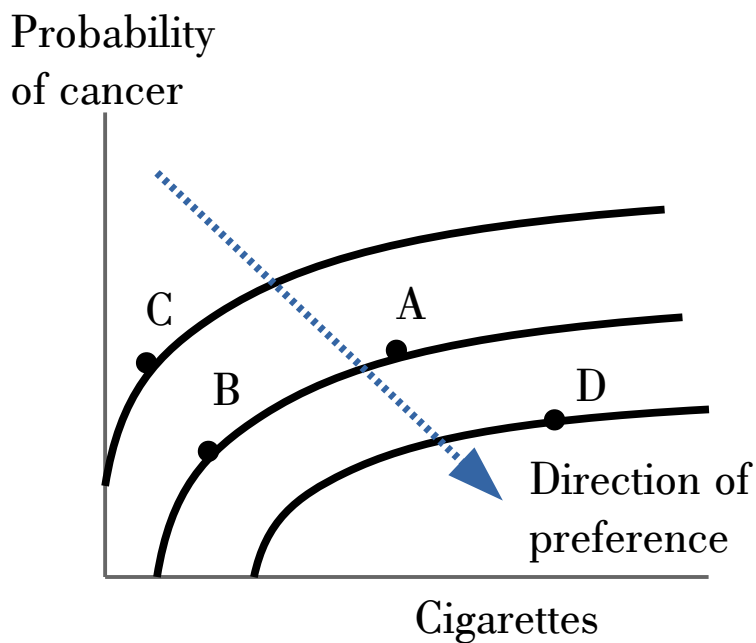
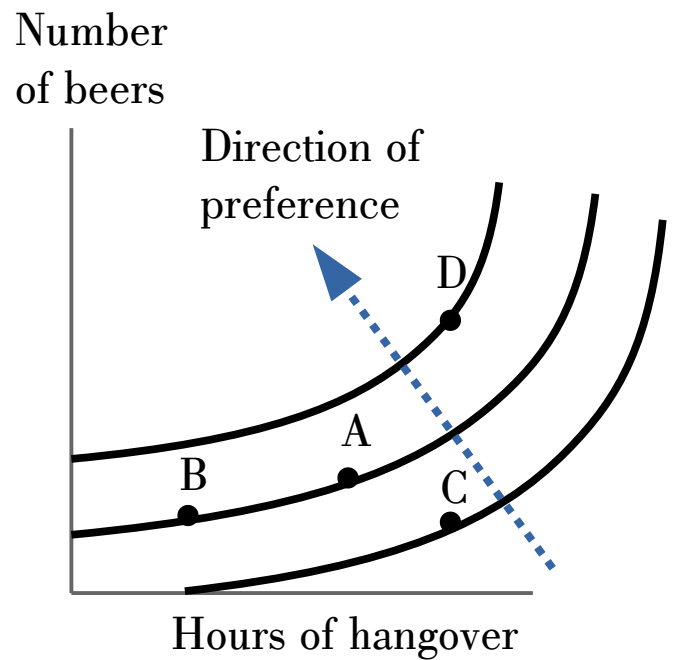
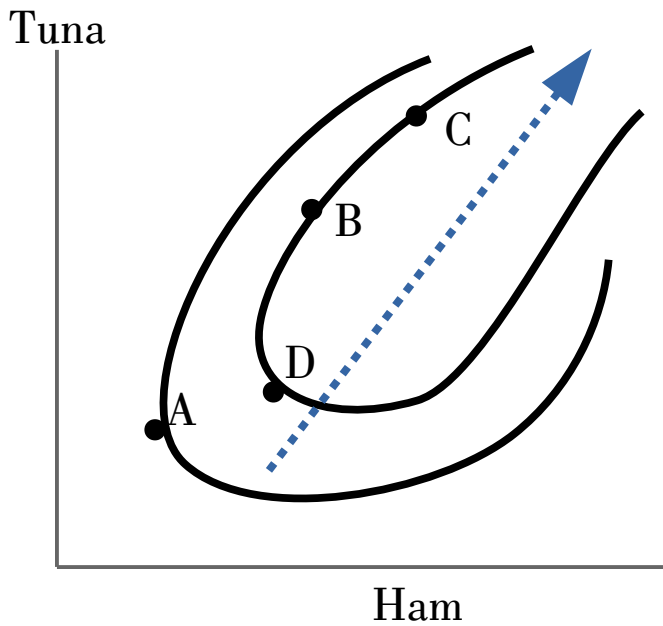


Figure 10: Examples of indifference curves

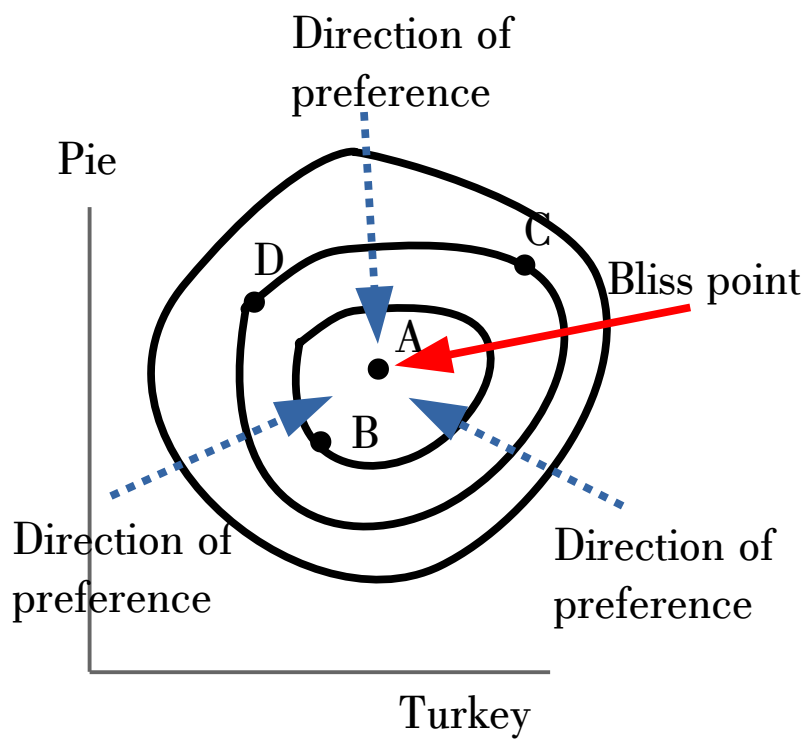
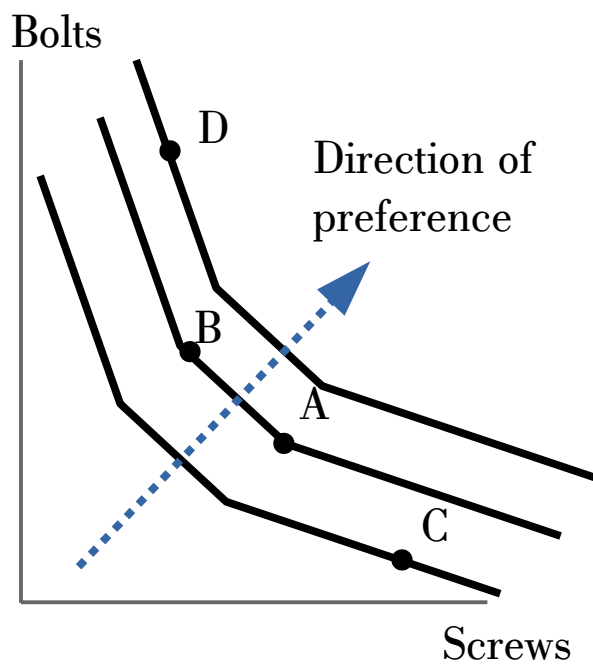
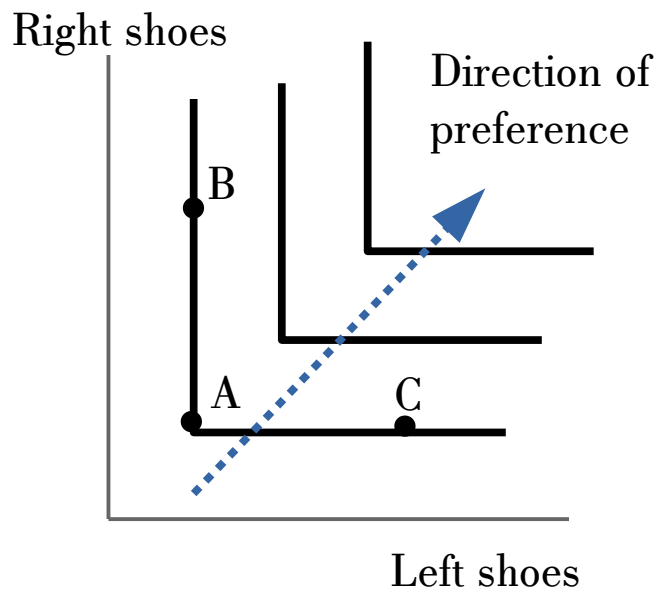
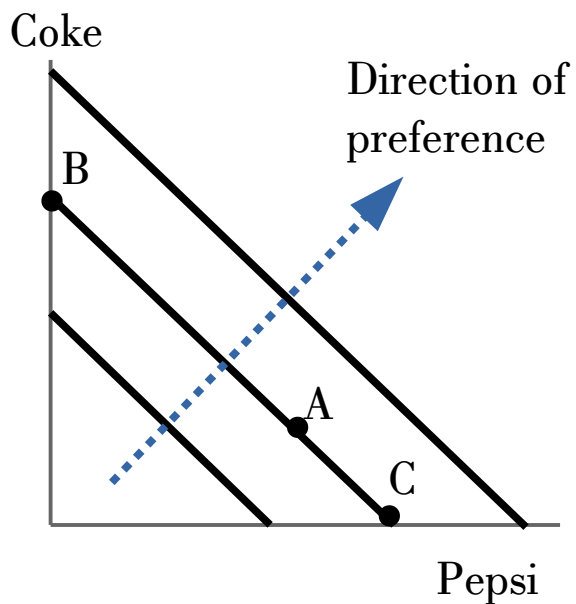


Figure 10: More examples of indifference curves

Perfect Substitutes: Many people cannot tell the difference between Coke and Pepsi in blind taste tests. If you are one of these people, Coke perfectly substitutes for Pepsi. Bundles with six cans of cola are all equally good regardless of what fraction is Coke or Pepsi. Note that the one-to-one ratio is not key. One two liter bottle of cola might be a perfect substitute for six individual cans. These indifference curves are strictly monotonic and weakly convex.

Perfect Complements: Right shoes are only useful to most people if they have a matching left shoe. A right shoe perfectly complements a left shoe. These are goods that are always consumed together. If you have one right shoe and one left shoe, additional left shoes do not improve your welfare (neither would they harm you). Thus, you are satiated in both left shoes and right shoes separately. Having more of both types of shoes puts you on a higher indifference curve. In general, these ratios may not be one-to-one, of course. Four tires perfectly complement one car, for example. These indifference curves are weakly monotonic and weakly convex.

Kinky Preferences: The picture shows indifference curves that start with a slope of -2 , then suddenly change to a slope of -1 and then finally to a slope of $-\frac{1}{2}$. What is happening in real terms is that the rate at which you are willing to exchange one good for the other changes suddenly instead of continuously.

As an example, suppose you are building a science project, and you have to connect several parts made of wood and plastic together. Bolts are best for plastic to plastic connections. It takes twice as many screws to do the same job. Screws are best for wood to wood connections. It takes twice as many bolts to do the same job. Bolts and screws are equally good at connecting plastic to wood.

You start by using screws make the wood to wood connections you need. If you run out, you have to use bolts and so bolts are only half as useful to you as screws would be. If you still have screws left over after all the wood to wood connections are done, you start in on the plastic to wood connections. If you run out of screws part way through, you switch to bolts which do an equally good job. Thus, you could trade bolts for screws at a one-to-one ratio and still be on the same indifference curve.

Finally, once you have finished all the wood to plastic connections, you start in on the plastic to plastic ones. Screws do a poor job, and you would happily trade two screws for one bolt at this point. In this rather far-fetched example, the indifference curves are strictly monotonic and weakly convex. We could also draw examples of kinky preferences that were strongly convex or only weakly monotonic.

Satiation: Suppose you are sitting down to Thanksgiving dinner. You can eat all the turkey and pie that you want. Bundle A on this graph is called a **Bliss Point**, the most preferred bundle in the entire consumption set. Consumption Nirvana, if you will. Moving in any direction makes you worse off. If your grandmother forces you to eat bundle C, you become bloated and end up on a lower indifference curve. This is just as bad as arriving so late that you only have a chance to reach bundle D before all the pie and turkey are gone. In the example, the indifference curves are strictly convex (but they do not have to be in general). They are neither strongly nor weakly monotonic, and do not even satisfy local non-satiation since there is a bliss point.

Goods and Bads: Smokers choose to buy and consume cigarettes. To smokers, cigarettes must be a good. Nevertheless, everyone agrees that a chance of getting cancer is a bad thing. Thus, agents would prefer to move in the direction of more cigarettes but less cancer. They are willing to trade these two things off and so a bundle B with a low risk of cancer, but very few cigarettes might be just as good as a bundle like A with a higher cancer risk but more cigarettes. You can check that average bundles are preferred to extreme ones in this example, and that these curves exhibit diminishing MRS. The example shows indifference curves that are strictly convex, and strictly monotonic in the direction of preference. Of course, it is possible to draw examples that satisfy only the weak versions of both of these.

Bads and Goods: Beer—Good! Hangover—Bad! This is just like the example above, but since the world is still spinning, the good and bad have switched axes. Note that the direction of preference has also reversed. Now go sleep it off.

Non-monotonic Preferences: Have a look at the indifference curve through bundle A. Note that bundles B and C are on the same indifference curve. However, bundle C lies strictly above bundle B! That is, bundle C has more ham and more tuna, but is exactly as good as bundle B. The example shows indifference curves that are strictly convex. However, we can clearly see that they bend backwards away from the commodity axes, and so they are neither strongly nor weakly monotonic. On the other hand, there is no violation of local non-satiation. There appears to be a better bundle in the neighborhood of every other bundle.

Bads and Bads: Your crazy roommate ties you to a chair and keeps you there for eight full hours. He has only two songs downloaded from Spotify: “Baby”, and “That’s the way (uh huh, uh huh) I like it”. He plays music continuously, but he lets you choose which song. Clearly, both Justin Bieber and KC and the Sunshine Band quickly become bads to you. You would prefer to move in the direction of no music at all (the origin). However, averages are preferred to extremes since each time you listen to the same song, it gets worse (you get more disutility on the margin). This is just as it is for goods: each time you consume a good it gets worse (you get less utility on the margin). Thus, the example shows strictly convex indifference curves that are strictly monotonic in the direction of preference.

Section 2.5. Needs, Wants, and Preferences

Let's finish by stepping back a bit in order to understand what preferences do and do not tell us. People often say that they "need" or "want" something. Although, this might seem to be closely related to an agent's preferences, the ideas of "need" and "want" have no operational meaning in economics.

What do we mean when we say something like: "I want to go on vacation", or "I want to have steak for dinner"? At the most abstract level, it might simply be equivalent to noting that these things would make us happier than we currently are, or more formally, would move us to a higher indifference curve.

In other words, saying we want something is the same as saying that this thing is a "good" and having more of it would therefore improve our well-being. In other, other words: preferences are monotonic. Saying you want something in this context, therefore, provides no information besides that you desire the commodity.

A more sophisticated meaning might be that I want to go on vacation, and taking into account the opportunity cost of doing so, this is the best use I can think of for my savings. Again, this certainly might be true, but if my financial circumstances change or the price of airfare increases, I might make a different choice. Thus, saying that a good is "wanted" tells us nothing new about the good itself beyond where it might rank in our preferences.

The notion of "need" is even harder to grasp. What does it mean when your daughter says she needs a new dress to go to a dance, or when your mechanic tells you that you need a new car, or that you think your brother-in-law needs to get a job, or that your mother-in-law needs to mind her own business? Often this is just a more emphatic statement of want, and want is just a statement of monotonicity.

Unfortunately, we have no way of measuring how intense feelings of want, desire, pleasure, or satisfaction. It may very well be the case in some abstract moral sense, a cultured person gets greater pleasure from listening to an opera, than I do from watching *Monster Trucks*. There is no objective metric we know of that can verify that my pleasure is less worthy, sincere, or intense than the pleasure experienced by more thoughtful and cultured people.

A more sophisticated meaning might be that something is necessary to accomplish an objective, or at least would make accomplishing the objective easier or more satisfying. Your daughter might be embarrassed to wear an old dress, it might save you money in the long run to buy a new car instead of continuing to pour money into fixing the old one, bailing your brother-in-law out of jail might be seriously cutting into your beer money, and your mother-in-law's meddling may be destroying your marriage.

Even if these are all factual statements about the world, there is no economic basis for thinking that your daughter can, or should be happy, that you should save money on cars, drink more beer, or have a successful marriage. These may be good things or bad things, and may or may not be worth having, depending upon the opportunity costs. You need to eat if you plan to survive, but this is a biological fact, not an economic one.

In short, to say you “need” something is either part of an if-then statement or a statement of strongly felt want. There is nothing in positive economics that lets us decide whether the “then” part of the statement or the satisfaction of an agent’s wishes are important, necessary, or even desirable. All we know is that at least one person thinks so.

As a father, my daughter’s happiness is important to me, and as a husband, I value my marriage. This means that I will expend the resources that I have under my control in these directions because they are better than the alternative uses available to me. From a philosophical or religious standpoint, I think it is terrible that anyone should not have enough food to survive. I may choose to give to the hungry or try to convince my neighbors to vote in favor of government programs to do the same. These are expressions of choice that come from my preference relation which, in turn, is informed by my own normative view of what is good and right. There is nothing wrong with this, but it is not positive economics either.

Preferences are simply a ranking of different bundles by a particular agent. They allow us to analyze an agent’s behavior, however, they tell us nothing about how much an agent “wants” something or what his “needs” are. These are social, psychological, or philosophical notions, not economic ones. To sum up:

**THERE IS NO SUCH THING AS AN ABSTRACT WANT OR NEED IN
ECONOMICS, THERE ARE ONLY PREFERENCE RANKING OVER
ALTERNATIVES.**

Section 2.6. Utility Functions

So far we have represented preferences by a symbol, $\succ_i$, and by indifference curve pictures. Note that the subscript here denotes that this is the preference relation of some individual $i \in \mathcal{I}$. We can also represent them by a **utility function**. All three representations are fully equivalent.

Formally, a utility function is a mapping from an agent's consumption set to the real numbers:

$$u_i : X_i \Rightarrow \mathbb{R}.$$

In words this means that a utility function assigns numbers to each bundle in an agent's consumption set. (We consider the behavior of a single agent for the rest of this chapter, and so we will drop the subscript for simplicity in the sections that follow.)

What makes this useful is that the utility function encodes the information contained in an agent's preference relation by assigning the same utility number to all bundles that are equally good. In addition, if any given bundle is strictly preferred to another, the utility function will assign it a larger utility number. A typical utility function might look like the following:

$$u(x_1, x_2) = (x_1)^{\frac{1}{2}} (x_2)^{\frac{1}{2}}.$$

In general, consider two bundles $(x_1, x_2), (\bar{x}_1, \bar{x}_2) \in X \subset \mathbb{R}^2$. The following statements are equivalent:

- (x_1, x_2) is on the same indifference curve as $(\bar{x}_1, \bar{x}_2)$
- $(x_1, x_2) \sim (\bar{x}_1, \bar{x}_2)$

Note that we can graph an indifference curve by finding all the bundles that have the same utility value. It might help to visualize this as a three-dimensional shape. Look at a corner of the room you are in and think of it as the origin of a three-dimensional coordinate system. The x_1 and x_2 dimensions are on the floor and are measured along the axes formed by the intersections of the floor and the two walls, respectively. We measure utility in the vertical dimension along the axis formed by the intersection of the two walls.

Thus, if two bundle, $(x_1, x_2), (\bar{x}_1, \bar{x}_2) \in X \subset \mathbb{R}^2$, share a utility value of $\bar{u}$, we place two points in the air at coordinates $(x_1, x_2, \bar{u})$ and $(\bar{x}_1, \bar{x}_2, \bar{u})$. We would place all other bundles with the same utility value at this same height of $\bar{u}$ above the floor.

Together, all these points would link up to form what is called a **level set of a function**. This is just like a contour line on a map which represents all the points on the ground that have the same elevation. In fact, if we mapped out a bunch of level sets with different heights/utility levels and

then looked down on the shape we created in the corner of the room (call it “Mount Utility”) from directly above, it would look just like the two-dimensional indifference curve pictures we have been drawing in this book. Maybe that did not help.

At a more formal level, if a utility function is concave, then it represents convex preferences (I am sorry about the counter-intuitive mathematical terminology). A concave function has a convex sub-graph. If you picture “Mount Utility” as described above, the mountain would be the sub-graph and this would be convex if it were derived from a concave utility function. Concavity turns out to be more than we need for utility functions, but it will be required when we introduce production functions, below.

Concave and Convex Functions

Concave Function: $f : \mathbb{R}^N \Rightarrow \mathbb{R}$ is a **concave function** if $\forall x, y \in \mathbb{R}^N$ and $\forall \lambda \in (0,1)$:

$$\lambda f(x) + (1-\lambda)f(y) \leq f(\lambda x + (1-\lambda)y).$$

Strictly Concave Function: $f : \mathbb{R}^N \Rightarrow \mathbb{R}$ is a **strictly concave function** if $\forall x, y \in \mathbb{R}^N$ and $\forall \lambda \in (0,1)$:

$$\lambda f(x) + (1-\lambda)f(y) < f(\lambda x + (1-\lambda)y).$$

Convex Function: $f : \mathbb{R}^N \Rightarrow \mathbb{R}$ is a **convex function** if $\forall x, y \in \mathbb{R}^N$ and $\forall \lambda \in [0,1]$:

$$\lambda f(x) + (1-\lambda)f(y) \geq f(\lambda x + (1-\lambda)y).$$

Strictly Convex Function: $f : \mathbb{R}^N \Rightarrow \mathbb{R}$ is a **strictly convex function** if $\forall x, y \in \mathbb{R}^N$ and $\forall \lambda \in (0,1)$:

$$\lambda f(x) + (1-\lambda)f(y) > f(\lambda x + (1-\lambda)y).$$

Quasi-Concave Function: $f : \mathbb{R}^N \Rightarrow \mathbb{R}$ is a **quasi-concave function** if $\forall x, y \in \mathbb{R}^N$ such that $f(x) \geq f(y)$, and $\forall \lambda \in [0, 1]$:

$$f(\lambda x + (1-\lambda)y) \geq f(y).$$

Note that the utility function: $u(x) = x_1 x_2$ is not concave, since, for example:

$$\frac{1}{2}u(2,2) + (1-\frac{1}{2})f(1,1) = \frac{1}{2}4 + \frac{1}{2}1 = 2.5 > u(\frac{1}{2}(2,2) + \frac{1}{2}(1,1)) = u(1.5, 1.5) = 2.25,$$

You can check that if a utility function is quasi-concave, then the weakly preferred sets are convex. That is, $\forall x, y, z \in \mathbb{R}^N$ such that $\forall U(x), U(y) \geq U(z)$, and $\forall \lambda \in [0,1]$:

$$U(\lambda x + (1-\lambda)y) \geq U(z).$$

Now, consider what would happen if we multiplied the example utility function by two? What if we added one? The utility value assigned to different consumption bundles would change, but if

two bundles had the same utility value before we multiplied or added, they would have the same (higher) value afterward. Thus, the shape of the indifference curves would be unchanged, but the new utility function would assign a higher utility value to each.

These utility numbers are only useful in the sense that bundles with higher utility values are preferred, lower utility values are inferior, and the same utility value are indifferent. The fact that I get one more “util” or twice as many “utils” from one bundle as opposed to another has no interpretation.

We have no way of measuring what it means to be one unit happier or twice as happy. We can tell that you are happier or sadder, perhaps, but not by how much. Economists summarize this by saying that a utility function carries only **ordinal** and not **cardinal** information.

Ordinal: Containing information about ordering. When we say that utility functions are ordinal, we mean that they order bundles in the sense that they tell which bundles are better or worse than one another.

Cardinal: Containing information about absolute differences. When we say that utility functions are cardinal, we mean that they tell us how much better one bundle is compared to another in some absolute sense. This is not something we are equipped to determine with the tools we have available in economics.

The fact that utility functions are ordinal has a number of mathematical implications.

First, we only need utility functions to be quasi-concave. Each level set of “Mount Utility” has to be convex, but the mountain itself can gain height in a nonconvex way.

Second, any monotonic transformation of a particular utility representation of a given specific preference relation is an equally good way to represent it. This is because the indifference curve implied by any member of a family or monotonic transformations of the same underlying utility function will have the same shape and the same order.

Recall that $f : \mathbb{R} \Rightarrow \mathbb{R}$ is a monotonic function if $\forall x, \bar{x} \in \mathbb{R}$

$$f(x) > f(\bar{x}) \Leftrightarrow x > \bar{x}.$$

Note that this implies the equivalent statements for “ $\geq$ ” and “ $=$ ”. For example, let $k \in \mathbb{R}^1$ be a real number, also called a constant or a scalar in this context. Here are some examples of monotonic and non-monotonic functions.

- $f(u) = u + k$ is monotonic
- $f(u) = u^k$ for $k > 0$ is monotonic
- $f(u) = ku$ for $k > 0$ is monotonic
- $f(u) = k - u$ is not monotonic

- $f(u) = \frac{k}{u}$ for $k > 0$ is not monotonic

Third, not only are we unable to make cardinal comparisons of the utility value of different consumption bundles for any given agent, we are even less able to make interpersonal comparisons of utility levels. That is, if we transfer some consumption from one agent to another, the first gets lower utility and the second gets higher utility. However, it is impossible to compare utility levels across people either before or after the transfer. This makes thinking about issues like equity or fair distribution quite difficult.

Monotonic Functions

Monotonic Function: $f : \mathbb{R} \Rightarrow \mathbb{R}$ is a **monotonic function** if

$$\forall x, y \in \mathbb{R}, x \geq y \Leftrightarrow f(x) \geq f(y).$$

Note that if f is differentiable, monotonicity implies $\forall x \in \mathbb{R}, \frac{\partial f}{\partial x} \geq 0$.

Strictly Monotonic Function: $f : \mathbb{R} \Rightarrow \mathbb{R}$ is a **strictly monotonic function** if

$$\forall x, y \in \mathbb{R}, x > y \Leftrightarrow f(x) > f(y).$$

Note that if f is differentiable, strict monotonicity implies $\forall x \in \mathbb{R}, \frac{\partial f}{\partial x} > 0$.

Strictly monotone functions have the property that they preserve order. Thus, if we compose a strictly monotone function f with another function:

$$g : \mathbb{R}^N \Rightarrow \mathbb{R} : f \circ g(x) \equiv f(g(x))$$

we say that the resulting composite function is a **monotonic transformation** of g since:

$$\forall x, y \in \mathbb{R}^N, g(x) > g(y) \Leftrightarrow f(g(x)) > f(g(y)).$$

Previously, we gave an informal definition of the marginal rate of substitution as the rate at which an agent would be willing to trade one good for another. Graphically, this can be understood as the slope between any two bundles on the same indifference curve.

The more formal definition of MRS takes the slope of the line connecting these indifferent trades as the distance between them goes to zero. In other words, the trades become smaller and smaller until they are infinitesimals. You can see that in the limit, this is the slope of the indifference curve at the initial consumption bundle.

We know that utility functions are ordinal and so their numerical value has no economic meaning. Nevertheless, consider the derivatives of a utility function with respect to any given good. We call these derivatives the **marginal utility of the good (MU)**:

$$MU_n \equiv \frac{\partial u}{\partial x_n}.$$

In words, the marginal utility gives the rate at which utils increase in response to an infinitesimal increase in the consumption of a good. For example, if the marginal utility of apples is g and we added or subtracted the smallest possible quantity ε of apple to the agent's consumption bundle, his utility level would go up or down by 4ε utils, respectively. Of course, this also has no economic meaning. However, suppose that the marginal utility of oranges is 2. Then if we took ε apples away but added 2ε oranges to the agent's consumption, he would first lose 4ε utils, but then gain 4ε utils. This would leave him at his initial utility level and be back to his initial indifference curve.

This means that on the margin, apples give twice the utility of oranges. Thus, the agent could trade tiny amounts of apples for oranges at a ratio of one to two and be just as well off. Putting this together, the ratio of the marginal utilities gives the rate which the agent would be willing to substitute one good for another starting from a specific initial consumption bundle. This ratio equals the *negative* of the slope of the indifference curve at that bundle.

$$\frac{\partial u / \partial x_m}{\partial u / \partial x_n} \equiv \frac{MU_m}{MU_n} \equiv MRS_{m,n}.$$

We finish this section with a list of several classes of useful utility functions:

Perfect substitute utility functions: $u(x) = \sum_{n \in \mathcal{N}} \alpha_n x_n$ where $\forall n \in \mathcal{N}, \alpha_n > 0$.

Perfect substitute utility functions give completely linear indifference curves. If then the slope of $\forall n \in \mathcal{N}, \alpha^n = 1$, the implied indifference curves is uniformly -1 . A good example might be series 2002 \$20 bills, and series 2004 \$20 bills. All you care about is how many \$20 bills are in your wallet, not the particular year in which they were printed.

Perfect complement utility functions: $u(x) = \min_{n \in \mathcal{N}} \{ \alpha_1 x_1, \dots, \alpha_n x_N \}$ where $\forall n \in \mathcal{N}, \alpha_n > 0$.

Perfect complement utility functions give right-angle indifference curves. If $\forall n \in \mathcal{N}, \alpha^n = 1$ then the corners line up along the 45° line where $\forall m, n \in \mathcal{N}, x_m = x_n$. A good example might be nuts and bolts. A bolt without a nut (or the inverse) is useless. Each bolt needs one and only one nut to allow it to secure parts together. Thus, all you care about is the number of nut/bolt pairs you own. Having a surplus of either nuts or bolts makes you neither better nor worse off. Thus, if you are at a bundle where you have the same number of nuts and bolts, you are satiated in both goods individually in the sense that increasing your allocation of either nuts or bolts, while keeping your allocation of the complementary item the same, leaves you on the same indifference curve.

Cobb-Douglas utility functions: $u(x) = \prod_{n \in \mathcal{N}} (x_n)^{\alpha_n}$ where $\forall n \in \mathcal{N}, \alpha_n \geq 0$, and $\sum_{n \in \mathcal{N}} \alpha_n = 1$.

Cobb-Douglas utility functions give strictly convex, and strictly monotonic indifference curves. These indifference curves are everywhere differentiable and are asymptotic to the axes (that is, they satisfy what are called “Inada conditions”). This makes them very convenient to work with algebraically.

Quasilinear utility functions: $u(x) = x_1 + v(x_2, \dots, x_N)$.

Quasilinear utility functions are linear in the first good (called the “numéraire good” or sometimes the “transferable good”). The remainder of the goods have a “sub-utility function” which could have any properties, but is generally assumed to be convex and monotonic. This implies that once we have mapped out one indifference curve, we can get all the rest of the indifference curves by shifting it right or left in the plane. This in turn implies that we can give agents any amount of the numéraire good without changing the MRS between any two of the other goods. Put another way, there are no income effects associated with the numéraire good. (In the next section, we will explore budget constraints, income effects and income elasticities.) This will often be a very useful case to examine. Perfect substitute utility functions are also quasilinear.

Homothetic utility functions: $u(kx) = ku(x) \quad \forall k > 0$, and $x \in X$.

Homothetic utility functions are homogeneous of degree A . (or are monotonic transformations of functions that are homogeneous of degree 1). Homothetic utility functions have the useful (though not realistic) trait that the income elasticities are uniformly equal to one, and so the income expansion paths are lines of different slopes starting from the origin. This in turn means that income distribution has no effect on aggregate demand. Perfect substitute and perfect complement utility functions are also homothetic.

Homogeneous, Linear and Affine Functions

Linear Function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}^N$ is a **linear function** if $\forall x \in \mathbb{R}^N, \forall \alpha \in \mathbb{R}_{++}^N$:

$$f(\alpha x) = \alpha f(x).$$

Affine Function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}^N$ is an **affine function** if $\forall x \in \mathbb{R}^N, g(x) \equiv f(x) - f(0)$ is a linear function.

Homogeneous Function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}^N$ is **homogeneous of degree k** , if $\forall x \in \mathbb{R}^N, \forall \alpha \in \mathbb{R}_{++}^N$, and some $k \in \mathbb{N}$, $f(\alpha x) = \alpha^k f(x)$.

For example, let $\alpha \in \mathbb{R}_{++}^N$ and $\beta \in \mathbb{R}^N$, then $f(x) = \alpha x$ is a linear function and $f(x) = \alpha x + \beta$ is an affine function. Of course all linear functions are affine, but not the reverse. You can also verify that all linear functions, but not all affine functions, are homogeneous of degree 1.

We finish with something called **lexicographic preferences**, sometimes called “dictionary ordering”. In a two-dimensional goods space, the idea is that agents care about the first good to the exclusion of the others. Thus, if bundle A has more good 1 than bundle B, an agent prefers it regardless of how much of good 2 there is in each bundle. Only if the bundles have exactly the same amount of good 1 does the agent consider good 2 at all relevant to his preference ranking. The agent prefers a bundle with more good 2 if and only if the both bundles have the same amount of good 1.

For example, suppose you are traveling from London to Paris by high-speed train with you family. If there is an accident, you might care first and foremost that all of your family survives a train wreck, and no amount compensation from Eurostar would be worth even one death. Given that unfortunate fact that a specific number of family members did not survive, you would prefer to get the highest possible monetary compensation. Thus, you care about the lives of your family lexicographically more than money.

In the case of two dimensions, we can define such a preference relation as follows:

$$x \succ \bar{x} \Leftrightarrow \{x_1 > \bar{x}_1\} \text{ or } \{x_1 = \bar{x}_1 \text{ and } x_2 > \bar{x}_2\}$$

and illustrate it thus:

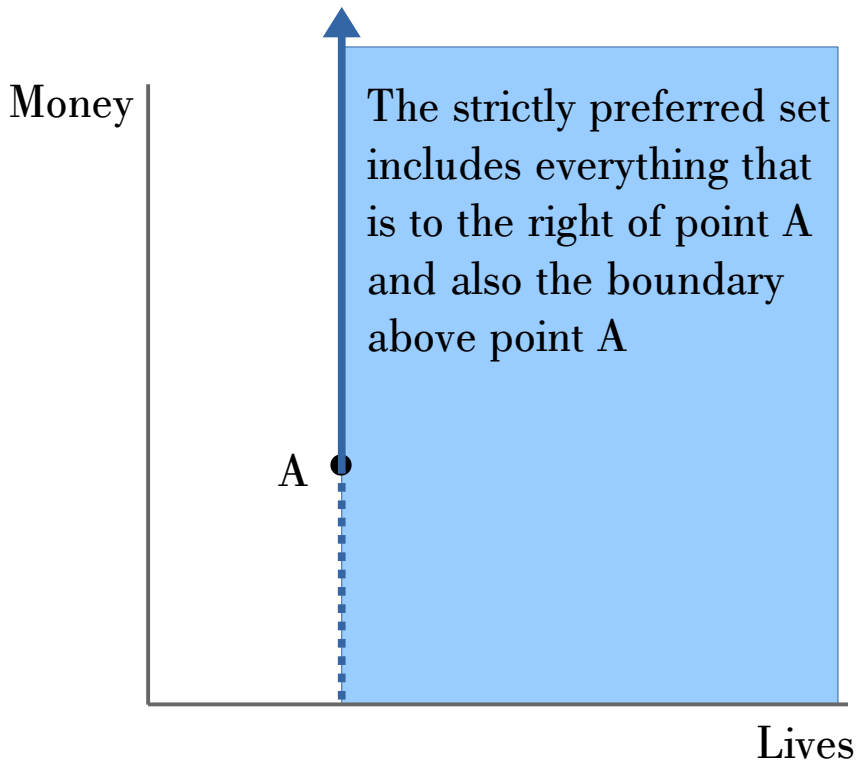


Figure 11: Lexicographic preferences

Note that we don't define a utility function. This is because lexicographic preferences are not continuous, and so are an example of **nonrepresentable preferences**. To see this, consider bundle A in the figure. What bundle are exactly as good as bundle A? That is what points are on the same indifference curve as bundle A? Try to identify the weakly and strongly preferred sets. Is the weakly preferred set closed? You can check that each indifference set consists of a single bundle. The preferred sets to bundle A are unbounded above and continue in the upward/rightward direction forever.

In n-dimensions, lexicographic preferences are defined as follows:

Lexicographic preferences: $\forall x, \bar{x} \in X, x \succ \bar{x}$ if and only if:

$$\exists m \in \mathcal{N}, \text{ such that } \forall n \leq m, x_n > \bar{x}_n$$

or

$$\exists m \in \mathcal{N}, \text{ such that } \forall n < m, \hat{x}_n = \bar{x}_n \text{ and } x_m > \bar{x}_m .$$

Note that this means that agents care lexicographically more about goods with lower indexes (that is, good 1 is most important, followed by good 2 and so on up to good N).

Section 2.7. Constraints

Having fully described consumers' objectives, we now turn to their constraints. The most basic constraint is called a budget constraint. Simply put, agents have a certain initial income and must choose a consumption bundle that costs no more than their income given prevailing prices.

To illustrate, consider the simple case of two goods, x_1 and x_2 . Suppose that the prices are p_1 and p_2 and that the agent's income or wealth is $w \in \mathbb{R}_+$. Then agent must choose a consumption bundle that satisfies the following budget constraint:

$$p_1 x_1 + p_2 x_2 \leq w.$$

In words, this says that the price of a good times quantity consumed of the good is the expenditure on the good in dollars, and the sum of expenditure on all goods added together can be no more than an agent's income or wealth in dollars. Any bundle in the consumption set X that satisfies this inequality is said to be in the **budget set**. Any bundle in the consumption set $\alpha, \beta \in \mathbb{R}$ that satisfies the constraint with equality is said to be on the **budget line**.

Have a look at this graph and answer the following:

- What happens if p_1 changes?
- What happens if p_2 changes?
- What happens if p_1 and p_2 both change?
- What happens if w changes?

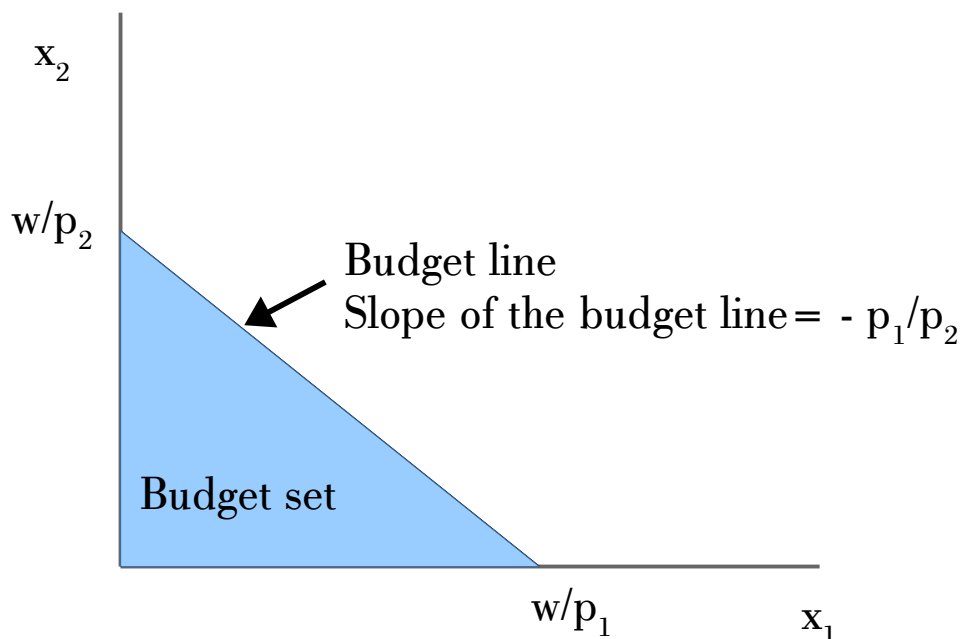


Figure 12: The budget constraint

More generally, the budget set is a special case of a feasible set. A **feasible set** simply represents the choices available to an agent. Here are some examples:

Discrete choices: Sometimes you have a list of feasible options. For example, if Princeton, the University of Chicago, and Vanderbilt accept you as an undergraduate, you have a feasible set consisting of three choices. Similarly, there may be a finite set of houses, jobs, or dateable people, to choose from. This situation is not representable as a budget set like the one above (what would it mean to consume two Princetons or one half of a University of Chicago?), but preferences can still be well-defined over sets like this. Maybe you like house number one better than house number two, for example.

Discrete choices over a hedonic space: Sometimes discrete choices have can be described as having different combinations of desirable, or undesirable characteristics. For example, houses differ in square footage, age, number of bathrooms, distance from work, and so on. We can think of a **hedonic** (things that bring pleasure) **decomposition** along these dimensions as representing the space of things we truly have preferences over. For example, suppose you had four job offers that had different salaries, but involved longer or shorter daily commutes (obviously, a bad):

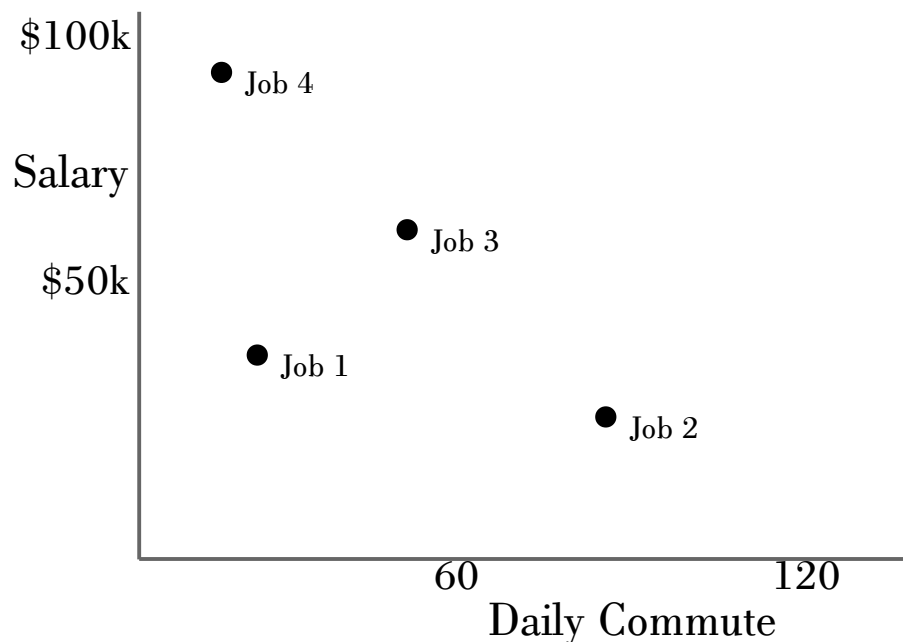


Figure 13: The budget constraint over the hedonic characteristics of four job offers.

Goods endowment: What if the agent started off with a bundle of goods he could trade at market prices rather than money income? For example, suppose he was initially endowed with five chickens and ten rabbits. What would his budget line look like? Suppose the initial relative prices are such that one rabbit is worth one chicken. First, we graph the agent's endowment in the goods space. This is a feasible choice for him regardless of whether he chooses to trade or not. He can always take his bundle and go home. The steeper budget line shows the trades he can make away from his endowment under the prevailing relative prices. Now suppose that the price of chickens doubles. The shallower budget line represents the trades available to the agent under the new prices. Note that the budget line pivots through the endowment point but now has a shallower slope to represent the new relative prices.

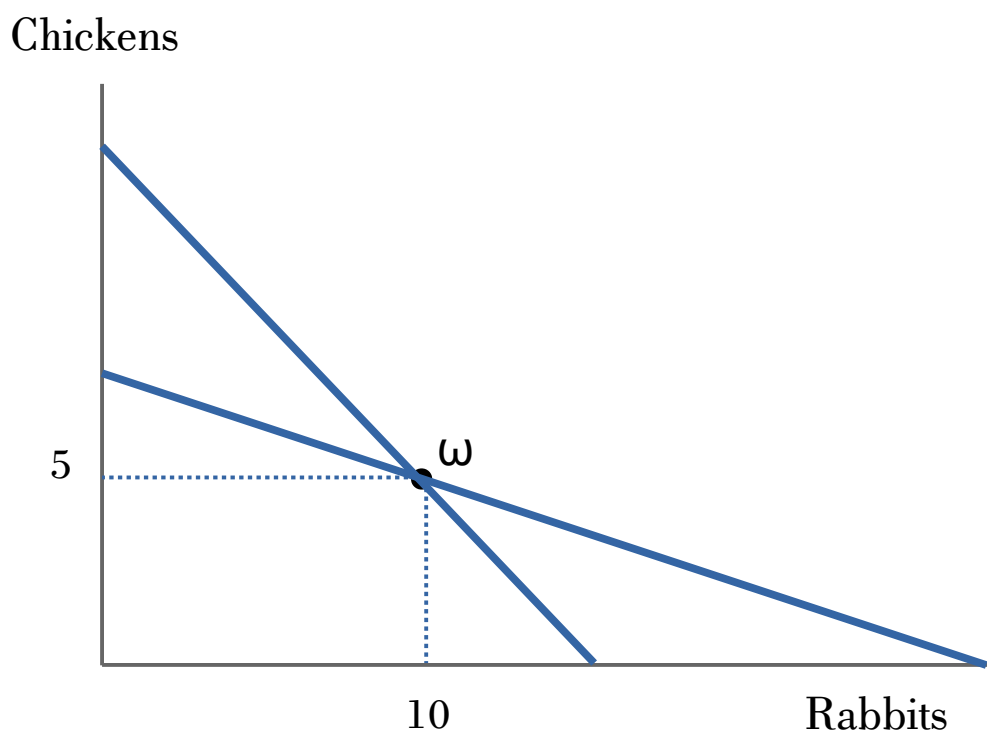


Figure 13: The budget constraint with an endowment of goods instead of income

Food stamps/EBT: The federal government gives poor people a EBT debit card that contains an amount of credit that depends on income, number of children, and so on. This credit can only be used to buy certain kinds of food but not alcohol, cigarettes, paper products, rent, clothes, gas, etc. This alters the budget line as follows: The agent can spend any money income he happens to have on either food, or something besides food. We think of this alternative use of income as purchasing a composite commodity we call **All Other Goods** (AOG). One dollar spent on food means one dollar less is available for all other goods. Thus, the “budget line” has a slope of -1 . Suppose the agent gets \$100 of food stamps. Then the budget line shifts out in the food direction by \$100. If the stamps are not spent on food, they are simply lost. Thus, the feasible set has a kink as shown.

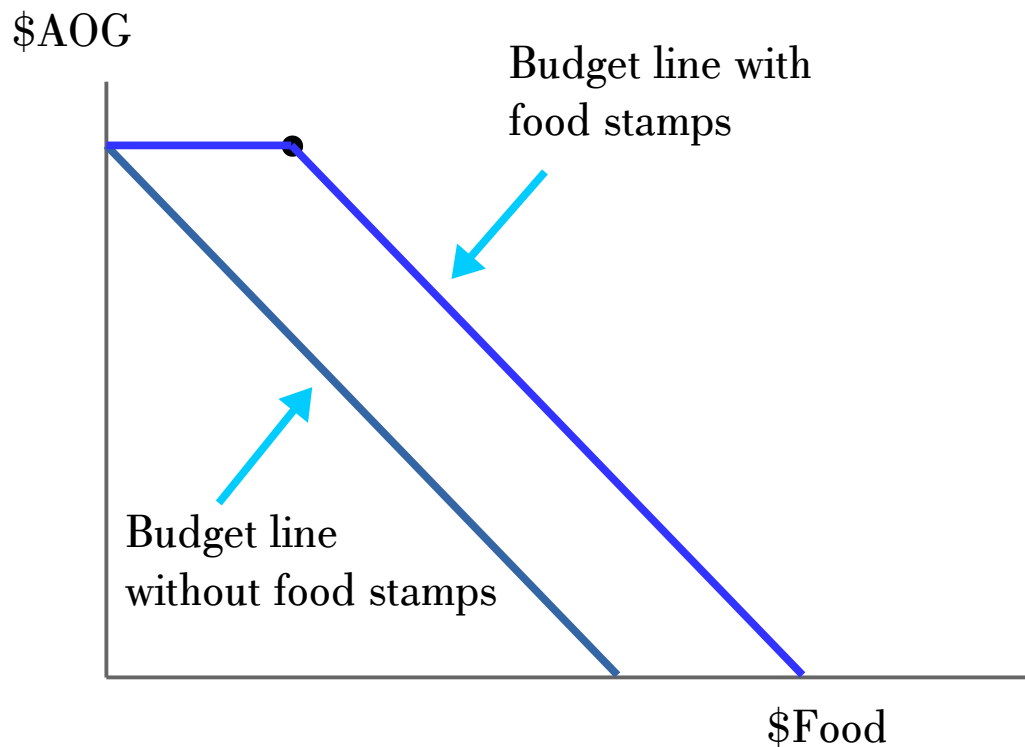


Figure 14: A nonlinear budget constraint with food stamps

Progressive income tax: Suppose an agent can work as many as 2000 hours per year at a wage rate of \$25 per hour. Any time not spent working is taken as leisure. The top most budget line shows this “labor-leisure” trade off. Now suppose the government imposes the following progressive income tax: (a) the first \$2500 the agent makes in gross income is not taxed at all, (b) all gross income between \$2500 and \$25,000 is taxed at a rate of 10% and (c) all income above \$25000 is taxed a rate of 50% You can see that this causes there to be two kinks in the feasible set.

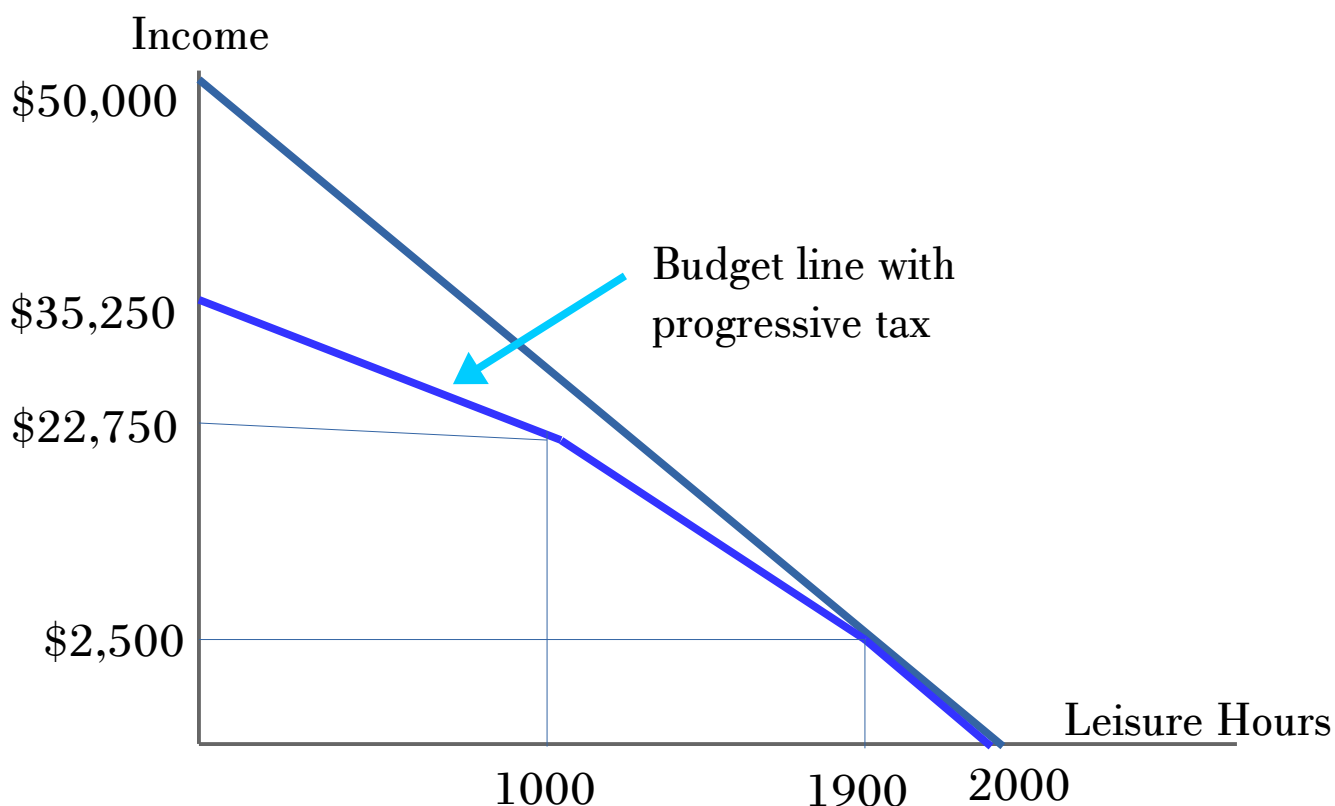


Figure 15: A nonlinear budget constraint with a progressive income tax

How would you graph the budget constraint for the following situations?:

Head tax: You are endowed with \$5000 and you are choosing between pizza and AOCs. Suppose you have to pay a tax of \$500 completely independently of anything else you do. This **lump sum tax** does not depend on any other consumption choice.

Intertemporal consumption: You are endowed with \$5000, and you are trying to decide how much you should consume this year, and how much to consume next year. Suppose the interest rate is 10%

Quantity discount: You love to go to movies, and so you must decide how many to see, and how much of your \$5000 endowment to save for other goods. Suppose that movie tickets cost \$20 each, but you can buy a booklet containing 10 tickets for \$150.

Glossary

Affine Function: We say $f: \mathbb{R}^N \Rightarrow \mathbb{R}^N$ is an affine function if $g(x) \equiv f(x) - f(0)$ is a linear function. (see the definition below).

All Other Goods (AOG): A fictional composite commodity measured in dollars. It allows us to turn a many dimensional consumer problem into a two-dimensional choice between some specific good, and dollars to be spent on a combination of all the other goods.

Bad: A commodity that is not desired by an agent and so lowers his utility level.

$$Spref(x) \equiv \{z \in X | z \succ x\}$$

Binary Relation: A set of ordered pairs drawn from a larger set. If (x, y) , is one of these ordered pairs that defines a binary relation R , we write $x R y$, and say that “ x stands in the relation R to y ”. Examples include preference relations ($\sim$, $\succcurlyeq$, and $\succ$) on an agent’s consumption set X , the vector relations ($\geq$, $>$, and $\gg$) on $\mathbb{R}^N$ and even notions like “is taller than” which is a well-defined binary relation over people.

Bliss Point: The most preferred bundle in an agent’s entire consumption set. An agent is completely satiated at a bliss point in all commodities, and moving in any direction makes him worse off.

Budget Line: The set of consumption bundles that cost exactly as much as an agent’s income under prevailing prices: $\sum_{n \in \mathcal{N}} p_n x_n = w$.

Budget Set: The set of consumption bundles that cost no more than an agent’s income under prevailing prices: $\sum_{n \in \mathcal{N}} p_n x_n \leq w$

Cardinal: Containing information about absolute differences. When we say that utility functions are cardinal, we mean that they tell us how much better one bundle is compared to another in some absolute sense. This is not something we are equipped to determine with the tools we have available in economics.

Closed Set: In the Euclidean space $\mathbb{R}^N$, a set that includes all of its boundary. More generally, the complement of an open set.

Cobb-Douglass Utility Function: Cobb-Douglas utility functions give strictly convex, strictly monotonic indifference curves that satisfy the Inada conditions. A Cobb-Douglas utility function takes the form:

$$u(x) = \prod_{n \in \mathcal{N}} (x_n)^{\alpha_n} \text{ where } \forall n \in \mathcal{N}, \alpha_n \geq 0, \text{ and } \sum_{n \in \mathcal{N}} \alpha_n = 1.$$

Commodity: A homogeneous class of items that can be exchanged with other agents. This includes goods, bads, material items, services, indivisible and divisible items, and even unique items.

Completeness: All bundles can be compared and ranked under an agent’s weak preference relation:(or both). $\forall x, \bar{x} \in X, x \succcurlyeq \bar{x} \text{ or } \bar{x} \succcurlyeq x$

Consumer's Problem: To choose the most preferred consumption bundle from a set of feasible alternatives.

Convex Hull: The smallest convex set containing an arbitrary set $S \in \mathbb{R}^N$:

$$con(S) \equiv \{z \in \mathbb{R}^N \mid \exists x_1, \dots, x_N \in S, \text{ and } \lambda \in \Delta^{N-1} \text{ such that } z = \sum_{n \in N} \lambda_n x_n\}.$$

Consumption Set: The set of all consumption bundles that are physically possible and permit an agent to survive.

Continuity of Preferences: Consumption bundles that are almost identical are almost equally good to any given agent. An agent's utility function is continuous if and only if his weakly preferred sets are closed: $\forall x \in X$, $Wpref(x)$ is closed, or equivalently:

$$\forall x \in X, \text{ and } \forall (x^s)_{s \in \mathbb{N}} \subset Wpref(x), \text{ if } \lim_{s \rightarrow \infty} x^s = \bar{x}, \text{ then } \bar{x} \in Wpref(x).$$

Diminishing Marginal Rate of Substitution: An implication of the convexity of preferences which says that the MRS between good m and good $\lim_{i \rightarrow \infty} x_i = \bar{x}$ goes down if the consumption of good m goes up while the consumption of good $\bar{x} \in Wpref(x)$ (and all other goods) remains the same.

Dot Product or Inner Product: Let $p, x \in \mathbb{R}^N$ be two N-dimensional vectors. Then the dot product or inner product is defined as follows:

$$p \cdot x = p_1 x_1 + p_2 x_2 + \dots + p_N x_N \equiv \sum_{n \in N} p_n x_n.$$

Epsilon Ball: An open epsilon ball around a point $x \in \mathbb{R}^2$ consists vectors that are strictly closer than epsilon to the vector x as measured by Euclidean distance:

$$B_\epsilon(x) \equiv \{z \in \mathbb{R}^N \mid \|x - z\| < \epsilon\}.$$

Euclidean Space: Informally, this is an n-dimensional coordinate space of real-valued vectors. If $N=2$, this is the 2-dimensional concordant plane that we were all introduced to sometime in middle-school. More formally, this is an n-dimensional, real, linear metric space endowed with an inner product operation. See the mathematical appendix for a precise definition and more clarity.

Euclidean Metric: The distance measure for Euclidean spaces often called the **Euclidean norm** denoted $\| \cdot \|$. Thus, Let $x, y \in \mathbb{R}^N$ be two N-dimensional vectors. The Euclidean distance between x and y is:

$$d(x, y) \equiv \|x - y\| \equiv \sqrt{\sum_{n \in N} (x_n - y_n)^2}.$$

Feasible Set: The set of the alternative consumption bundles available for an agent to choose over. The budget set is a special case of a feasible set.

Good: A commodity that is desired by an agent.

Homogeneous of Degree k: We say $f: \mathbb{R}^N \Rightarrow \mathbb{R}^N$ is **homogeneous of degree k**, if $\forall x \in \mathbb{R}^N, \forall \alpha \in \mathbb{R}_{++}^N$, and some $k \in \mathbb{N}$, $f(\alpha x) = \alpha^k f(x)$.

Homothetic Utility Function: Homothetic utility functions are homogeneous of degree 1 (or are monotonic transformations of functions that are homogeneous of degree 1). They have the property that the MRS is constant along any ray drawn from the origin. Homothetic utility functions take the form:

$$u(kx) = ku(x) \quad \forall k > 0, \text{ and } x \in X.$$

Indifference Curve: A set of bundles in an agent's consumption set that are equally good as one another under the agent's preferences. In general, an agent has an infinity of distinct indifference curves, and every bundle in an agent's consumption set is on one and only one indifference curve:

Indifference Relation: If bundle $x \in X$ is identically as good as bundle $\bar{x} \in X$ to agent i , we write:

$$x \sim_i \bar{x}.$$

Lexicographic Preferences: Lexicographic preferences are sometimes called “dictionary orderings” and are not continuous. Agents care about the consumption level of each successive good to the exclusion of remaining goods. Lexicographic preferences are defined as follows:

$\forall x, \bar{x} \in X$, $x \succ \bar{x}$ if and only if:

$$\exists m \in \mathcal{N}, \text{ such that } \forall n \leq m, x_n > \bar{x}_n,$$

or

$$\exists m \in \mathcal{N}, \text{ such that } \forall n < m, \hat{x}_n = \bar{x}_n \text{ and } x_m > \bar{x}_m.$$

Linear Function: We say $f: \mathbb{R}^N \Rightarrow \mathbb{R}^N$ is a **linear function** if

$$\forall x \in \mathbb{R}^N, \forall \alpha \in \mathbb{R}_{++}^N, \quad f(\alpha x) = \alpha f(x).$$

Local Nonsatiation: An agent's preferences satisfy local nonsatiation if in every neighborhood (no matter how small) of every consumption bundle in an agent's consumption set, there exists another consumption bundle that he strictly prefers:

$$\forall x \in X, \text{ and } \forall \varepsilon \in \mathbb{R}_+, \exists \bar{x} \in X \text{ such that } \|x - \bar{x}\| < \varepsilon \text{ and } \bar{x} \succ x.$$

Lump Sum Tax: A tax of a fixed amount that each agent must pay. The tax level is completely independent of consumption, production or any other choice made by the agent, and so does not distort an agent's economic decisions on the margin. While they do not generate any substitution effects, they do have income effects that may affect an agent's decisions. These are also called **head taxes**.

Marginal Rate of Substitution: The rate at which an agent could exchange one good for another and still remain on the same indifference curve. This is defined at every point in the consump-

tion set and is equal to the negative of the slope of the indifference curve at any given point. Mathematically:

$$\frac{\partial u / \partial x_m}{\partial u / \partial x_n} \equiv \frac{MU_m}{MU_n} \equiv MRS_{m,n}.$$

Monotone Function: A function $f: \mathbb{R} \Rightarrow \mathbb{R}$ that is order preserving, That is:

$$\forall x, y \in \mathbb{R}, x \geq y \Leftrightarrow f(x) \geq f(y).$$

Monotonic Transformation: We say that the composition function, $f \circ g(x) \equiv f(g(x))$ is a monotonic function of $g: \mathbb{R}^N \Rightarrow \mathbb{R}$ if $\forall x, y \in \mathbb{R}^N, g(x) > g(y) \Leftrightarrow f(g(x)) > f(g(y))$.

Need: There is no such thing as “need” in economics, there are only preference ranking of alternatives. In informal usage, something may be needed in the sense that it is necessary to accomplish an objective, or at least would make accomplishing the objective easier or more satisfying. Need may also mean that an agent believes that his desire for a good is very strong in some sense.

Nonrepresentable Preferences: A preference ordering that can not be represented with a continuous function. That is, preferences with indifference sets that can not be mapped in a one-to-one correspondence with the level sets of any continuous function.

Open Set: In the Euclidean space $\mathbb{R}^N$, a set that includes none of its boundary. More generally, a set of sets is simply defined to be “open” and gives a foundation for what is called a “topology”.

Ordinal: Containing information about ordering. When we say that utility functions are ordinal, we mean that they order bundles in the sense that they tell which bundles are better or worse than on another.

Perfect Complements: Two commodities are perfect complements if they must be consumed in a specific ratio. Agent are satiated in both commodities individually if they are currently consuming in this ratio. If all goods are perfect compliments for an agent, then his utility function takes the form:

$$u(x) = \min_{n \in \mathcal{N}} \{ \alpha_1 x_1, \dots, \alpha_n x_N \} \text{ where } \forall n \in \mathcal{N}, \alpha_n > 0.$$

Perfect Substitutes: Two commodities that can exchanged with one another at some fixed ratio without changing the well-being of an agent. If all goods are perfect substitutes for an agent, then his utility function takes the form:

$$u(x) = \sum_{n \in \mathcal{N}} \alpha_n x_n \text{ where } \forall n \in \mathcal{N}, \alpha_n > 0.$$

Preference Relation: Binary relations denoted, $\sim$, $\succsim$, and $\succ$, over agents' consumption sets that indicate, respectively, that commodity bundles are equally good, at least as good, or strictly better than another bundle according to a given agent's pattern of tastes.

Price Taker: An agent who assumes that the prices he observes in the market are unaffected by any purchases, sales, or other actions he may take.

Quasilinear Utility Function: Quasilinear utility functions are linear in the first good (called the “numéraire good” or sometimes the “transferable good”). Quasilinear utility functions take the form:

$$u(x) = x_1 + v(x_2, \dots, x_N).$$

Satiation: Having had enough, but not too much of a commodity. If a consumer is satiated in a good, then more of the good makes him neither better off, nor worse off. He is simply indifferent to having more of the good than he has already.

Strong or Strict Preference Relation: If bundle $x \in X$ preferred to bundle $\bar{x} \in X$ by agent i , we write:

$$x \succ_i \bar{x}.$$

Strong or Strict Convexity: Strictly weighed averages of two consumption bundles leave an agent strictly better off and the marginal value of a good strictly declines as an agent consumes more:

$$\forall \lambda \in (0, 1), \forall x \in X \text{ and } \forall \bar{x}, \hat{x} \in Wpref(x) \text{ such that } \bar{x} \neq \hat{x}: \lambda y + (1-\lambda)z \in Spref(x).$$

Strong or Strict Monotonicity: Agents are made strictly better off if you give them strictly more of at least one good while keeping them at least at the same level of consumption of all other goods:

$$\forall x, \bar{x} \in X \text{ such that } x > \bar{x}, \text{ it holds that } x \succ \bar{x}.$$

Strongly or Strictly Preferred Set: The set of bundles in an agent’s consumption set that are strictly preferred to a given bundle:

$$x: Spref(x) \equiv \{z \in X \mid z \succ x\}.$$

Transitivity: Agents' preference relations do not allow for inconsistent cycles of preference ordering:

$$\forall x, \bar{x}, \hat{x} \in X, \text{ if } x \geq \bar{x} \text{ and } \bar{x} \geq \hat{x} \text{ then } x \geq \hat{x}.$$

Utility Function: A utility function assigns a utility number to bundles of goods such that in a way that respects the preference ordering of an agent. All bundles that are equally good are assigned the same utility number and collectively define an indifference curve. Preferred indifference curves are assigned higher utility numbers than inferior ones. The idea of numerical utility is completely artificial, but is nevertheless useful. Formally, utility functions map an agent’s consumption set to the real numbers:

$$u_i: X_i \Rightarrow \mathbb{R}.$$

Vector Addition: Let $x, y, z \in \mathbb{R}^N$ be vectors. Then:

$$x + y = (x_1 + y_1, x_2 + y_2, \dots, x_N + y_N).$$

Scalar Multiplication: Let $\alpha \in \mathbb{R}$ be a real number, called a scalar, and $x \in \mathbb{R}^N$, be a vector. Then:

$$\alpha x = (\alpha x_1, \alpha x_2, \dots, \alpha x_N).$$

Want: There is no such thing as a want in economics, there are only preference ranking of alternatives. In informal usage, when an agent says he wants something, he may mean that, taking into account the opportunity cost, this is the best use of his resources he can find.

Weak Convexity: Weakly weighed averages of two consumption bundles leave an agent weakly better off and the marginal value of a good weakly declines the more an agent consumes more:

$$\forall \lambda \in [0, 1], \forall x \in X \text{ and } \forall \bar{x}, \hat{x} \in Wpref(x), \lambda \bar{x} + [1 - \lambda] \hat{x} \in Wpref(x).$$

Weak Monotonicity: Agents are at least as well off if you give them a bundle that has at least as much of each good, m , and are made strictly better off if you give them a bundle that has strictly more of all goods: $\forall x, \bar{x} \in X$ such that $x \geq \bar{x}$, it holds that $x \succcurlyeq \bar{x}$. In addition, $\forall x, \bar{x} \in X$ such that $x \gg \bar{x}$, it holds that $x \succ \bar{x}$.

Weak Preference Relation: If bundle $x \in X$ is at least as good as bundle $\bar{x} \in X$ to agent i , we write:

$$x \succcurlyeq \bar{x}.$$

Weakly Preferred Set: The set of bundles in an agent's consumption set that are weakly preferred to a given bundle, x :

$$Wpref(x) \equiv \{z \in X \mid z \succcurlyeq x\}.$$

Problems

1. Logic: Consider the set all ancient Greeks denoted G . Let g denote a typical Greek, thus $g, g_1, \dots$, are specific Greeks who are all members of the set all Greeks G . Thus, $g_m, g_n \in G$. Some Greeks are very smart, and all Greeks have their IQ tattooed on their foreheads. Let IQ_n denote the IQ of Greek person $g_n \in G$. If a Greek has an IQ of at least $\overline{IQ}$, he is a philosopher. The smartest philosopher is called the philosopher king. The set of philosophers is denoted P and philosopher king. Unfortunately, to be a king of any kind in Ancient Greece, you have to be male. The set of all males is denoted M . Write the following in formal mathematical notation.
 - a. Give a definition of the set of philosophers
 - b. Give a definition of the philosopher king.
 - c. Will the set of philosophers always be non-empty
 - d. If the set of philosophers is not-empty, will there always be a philosopher king?
 - a. What assumptions would be sufficient to guarantee the existence of a philosopher king?
2. Suppose that there are four commodities in the economy. Commodities 1 and 2 are goods, but commodities 3 and 4 are bads.
 - a. Write a formal definition of strong monotonicity that takes this fact into account.
 - b. Write a formal definition of weak convexity that takes this fact into account.
3. You go an all-you can eat restaurant called the Ample American. It is glorious! All the Apple pie and hot dogs you can eat and not a French fry in sight (freedom fries and tater-tots are available, however). Unfortunately, you can only eat so much without exploding. At some point, pushing in one more hot dog becomes painful. What would your indifference curves over apple pie and hot dogs look like in this situation?
4. Suppose that your utility function is: $u(x) = x_1 + \sqrt{x_2}$.
 - a. What is the formula for the marginal rate of substitution? What is the value of the MRS at the consumption points $(3, 5)$, $(7, 5)$, and $(5, 7)$? What is your utility level at each point?
 - b. Sketch the Indifference curves through the three points above.
 - c. Suppose your best friend's utility function is $u(x) = (x_1)^2 + 2x_1\sqrt{x_2} + x_2$. Answer part (a) and (b) for your BFF's utility function. What do you notice. How do you explain your observation?
5. Suppose Covid-19 has you worried, and you are afraid to leave your house. You think about ordering your favorite restaurant, Nemo Sushi. Rather than risk having to touch another person. You have \$30 to spend. Uber eats charges you a flat fee of \$2 for delivery, however you notice that while Nemo Sushi charges \$5 per roll if you pick it up, the menu price is \$7 if you order through the Uber app. You also notice that if you order \$28 worth of food, Nemo will give you an extra roll for free (this is true for both pickup and Uber deliveries). Draw your budget set for pickup and Uber Eats.

Chapter 3. Consumers: Optimization and Demand

Section 3.1. Optimization

We described both the preferences and the constraints of consumers in Chapter 2. In this Chapter we put these together and show how an agent makes the best consumption choice available to him. Formally, the agent is simply solving a constrained optimization problem. To get a sense of what this means, consider the figure below.

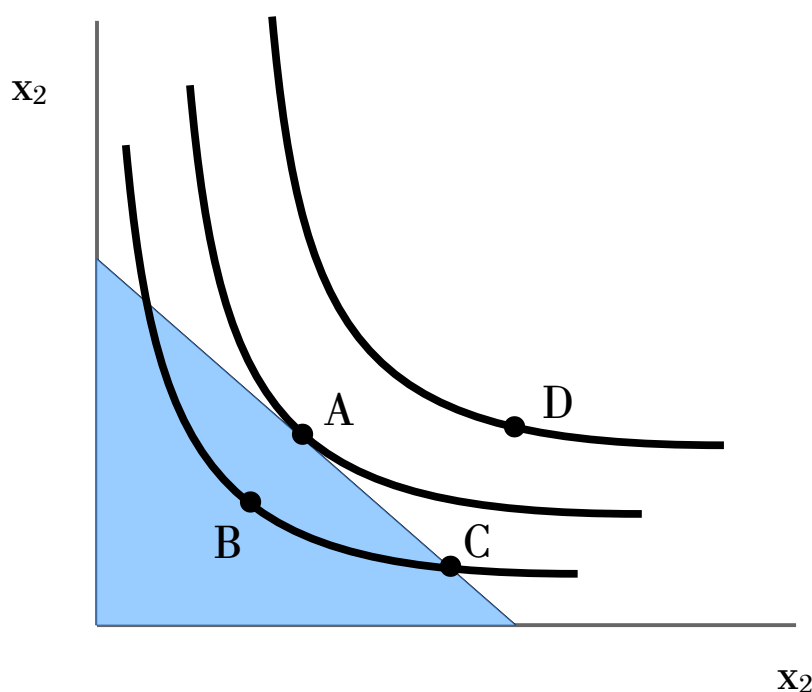


Figure 16: The optimal consumption bundle in a budget set

Which of these bundles is optimal, and why? We claim it is consumption bundle A. To see this, consider the alternatives. Bundle B is in the budget set (like bundle B) but is on a lower indifference curve. Bundle C is on the budget line rather than merely in the budget set, but again bundle C is C on a higher indifference curve. Finally, bundle D is on a higher indifference curve than bundle B but bundle D it is not within the budget set and so is not an affordable choice. Bundle C is therefore not an optimal choice. Putting this together, we can make the following observations:

- Bundle A is the best choice because A is on the highest indifference curve which still makes contact with budget set.
- This implies that the optimal choice will be where an indifference curve is *tangent to the budget set*.
- This in turn implies that it is a *necessary condition* that the slopes of the indifference curve and the budget line are the same at an optimal consumption choice. That is: $MRS_{1,2} = \frac{p_1}{p_2}$.
- Equivalence of slopes between the indifference curve and the budget line, however, is not a *sufficient condition* for optimality. Observe that the slope of the indifference curve at bundle A is the same as the budget line and yet B is not an optimal choice.
- If a bundle is optimal, there will be no intersection between the feasible set and the strongly preferred set. Observe that all bundles above the indifference curve through bundle B are not in the budget set. In other words, all strictly better bundles are infeasible, and all feasible bundles are no better than bundle A. ***This is the gold standard! An empty intersection between the feasible and strictly preferred sets is necessary and sufficient for optimality.***
- Separation: At a more general level, suppose we have two convex sets, such as a budget set and a strongly preferred set. Then at any optimal choice we can draw a line through the optimal bundle such that all the feasible choices are below this line and the strictly preferred bundles are above this line. We call this a **line of separation**. In most economic contexts this line touches, but does not penetrate either the feasible set or the preferred set. In fact, it is a **line of tangency** to both. In other words, the slope of the budget line, the indifference curve through the optimal bundle, and the line of tangency are all the same at the optimal bundle. Since these lines touch the sets but do not go into the interior, they are also called **lines of support**. (Lines of separation do not always touch the sets being separated, but lines of support do separate such sets.) In more than two dimensions, lines of support are called **supporting hyperplanes**.
- **Walras' Law:** If a consumer has monotonic preferences, then at an optimal consumption choice, that is, agents with monotonic preferences will always spend their entire income at an optimal consumption choice. Although, the budget constraint requires only that the agent spend no more than his income (or the value of this endowment), monotonicity implies that the agent will always choose to completely exhaust his income and choose a bundle on the budget line and not in the interior on the budget set.

Hyperplanes and Separation Theorems

Hyperplane: Let $p \in \mathbb{R}^N$ with $p \neq 0$ and $k \in \mathbb{R}$, then the **hyperplane generated by p and k** is defined as:

$$H_{p,k} \equiv \{z \in \mathbb{R}^N \mid pz = k\}.$$

Upper Half Space: The upper half space of $H_{p,k}$ is defined as: $\{z \in \mathbb{R}^N \mid pz \geq k\}$.

Upper Lower Space: The lower half space of $H_{p,k}$ is defined as: $\{z \in \mathbb{R}^N \mid pz \leq k\}$.

Separation: Consider two sets $S, T \subset \mathbb{R}^N$. A hyperplane $H_{p,k}$ is said to **separate S from T** if:

$$\forall x \in S, \text{ and } y \in T, \quad px \leq k \leq py.$$

Separating Hyperplane Theorem: Let $S \subset \mathbb{R}^N$ be convex and closed, and $x \notin S$ be a vector in $\mathbb{R}^N$. Then $\exists p \in \mathbb{R}^N$ with $p \neq 0$ and $k \in \mathbb{R}$ such that $px > k$ and $\forall y \in S, py < k$.

More generally:

Minkowski Separation Theorem: Let $S, T \subset \mathbb{R}^N$ be disjoint convex sets, $S \cap T = \emptyset$. Then $\exists p \in \mathbb{R}^N$ with $p \neq 0$ and $k \in \mathbb{R}$ such that $H_{p,k}$ separates S from T .

Supporting Hyperplane Theorem: Let $S \subset \mathbb{R}^N$ be a convex set, and $x \in \text{boundary}(S)$. Then $\exists p \in \mathbb{R}^N$ with $p \neq 0$ and $k \in \mathbb{R}$ such that $\forall y \in S, k = px \geq py$.

Examples of Hyperplanes and Separation Theorems

The figure below illustrates these Theorems. We have two convex sets, S and T , two points y and x in their interiors, respectively, and z , a point on the boundary of S . The supporting/separating hyperplane of S at z has slope of -4 and so $p = (4,1)$ is a vector that is orthogonal or perpendicular to this hyperplane. (Note that in two dimensions, the hyperplane is really just a line.) We will set $k = 14 = pz = (4,1)(3,2)$.

To understand the Separating Hyperplane Theorem, first note that for the point $x \notin S$, it is the case that $px = (4,1)(4.5,3.5) = 21.5 > 14 = k$. This is true because x is “above”, and thus, separated from S by $H_{p,k}$. we can always find a hyperplane like this that puts all of S “below” x if $x \notin S$. On the other hand, $py = (4,1)(1,3) = 7 < 14 = k$. By a similar argument, this must be true for all $py < k$.

The Minkowski Separation Theorem extends this to disjoint sets like S and T . Suppose that S is closed. Then we can find a tangent hyperplane to S that does not intersect T since both are convex. Then all points $x \in S$ will satisfy $px \leq k \leq py$ while all points $y \in T$ will satisfy $py \leq k \leq px$. Think about whether something similar will hold if S is not closed.

Finally, the Supporting Hyperplane Theorem says if we have a point on the boundary of a convex set, then we can “support” it, meaning that we can find a tangent hyperplane to S at the point. In the figure, z is a point on the boundary of S , and $pz = 14$. You can see that every point $y \in \text{interior}(S)$ will have the property $py < 14$.

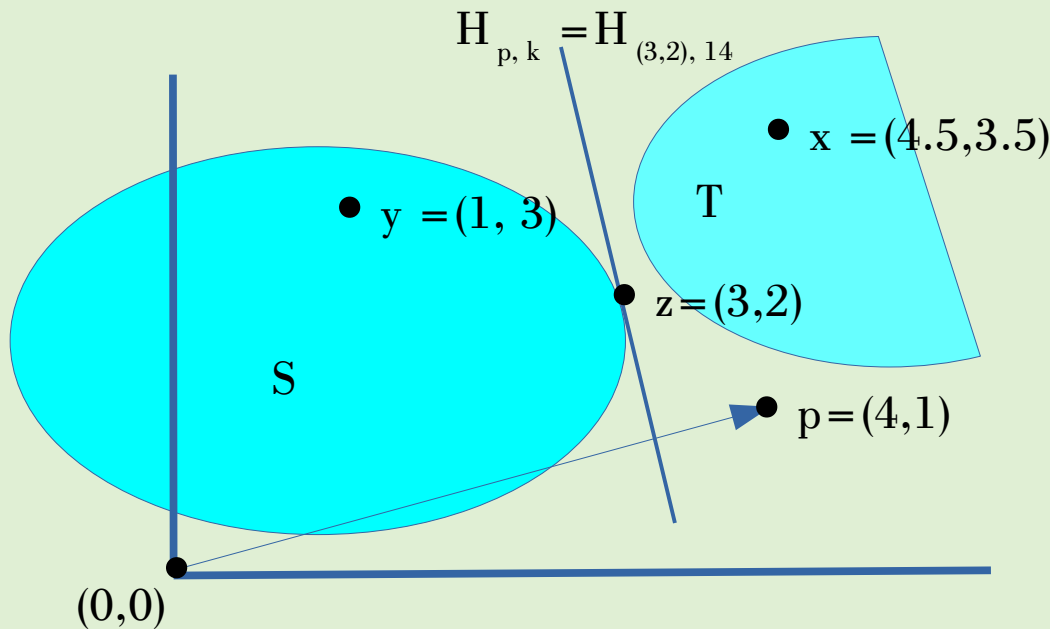


Figure 17: Separation and support theorems

We finish this section with a few examples of optimal choice using other indifference curves for practice. Note that the dashed lines represent budget lines with different prices for the good on the A -axis.

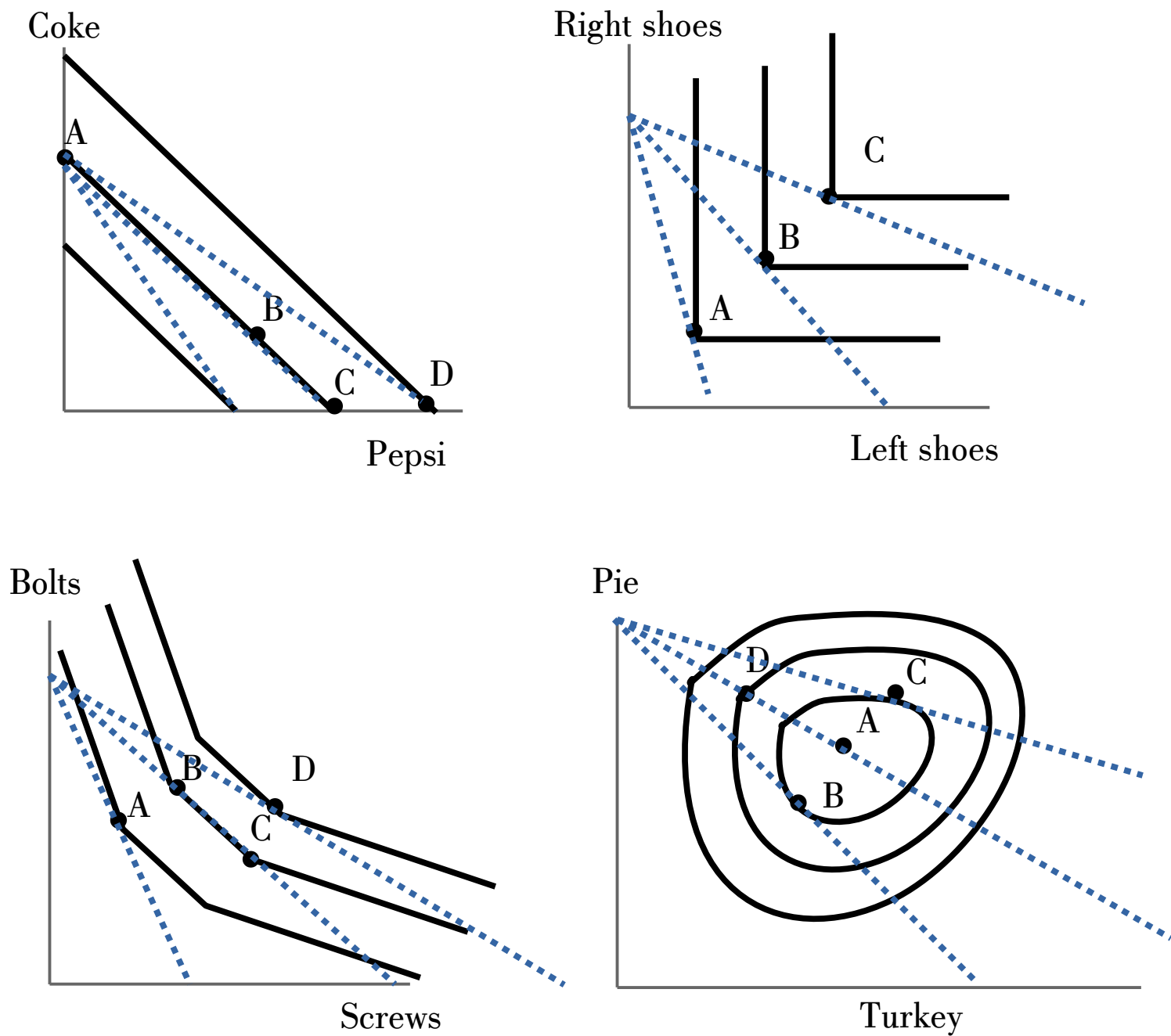


Figure 18: Examples of optimal choices with different types of preferences

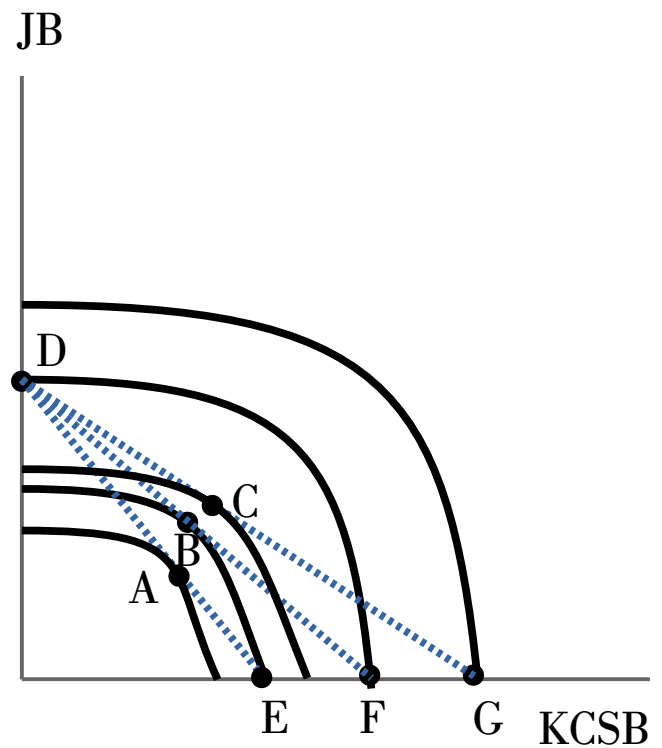
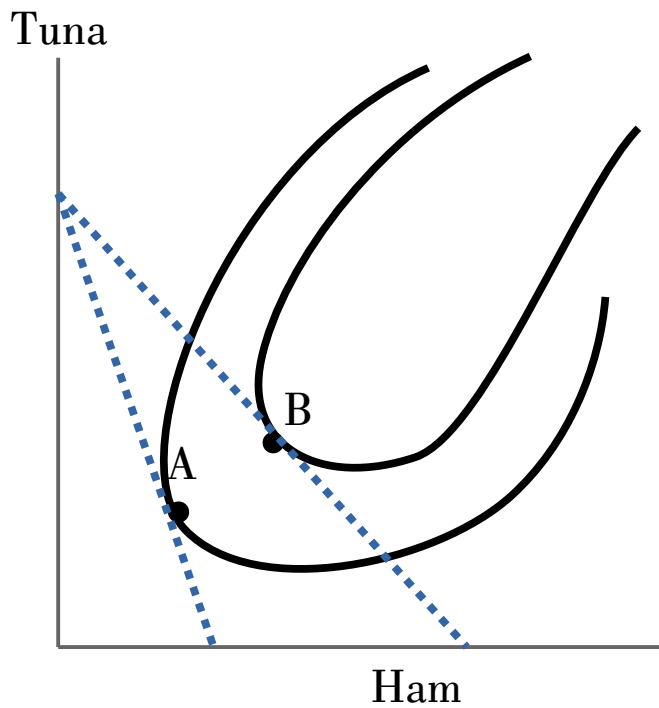
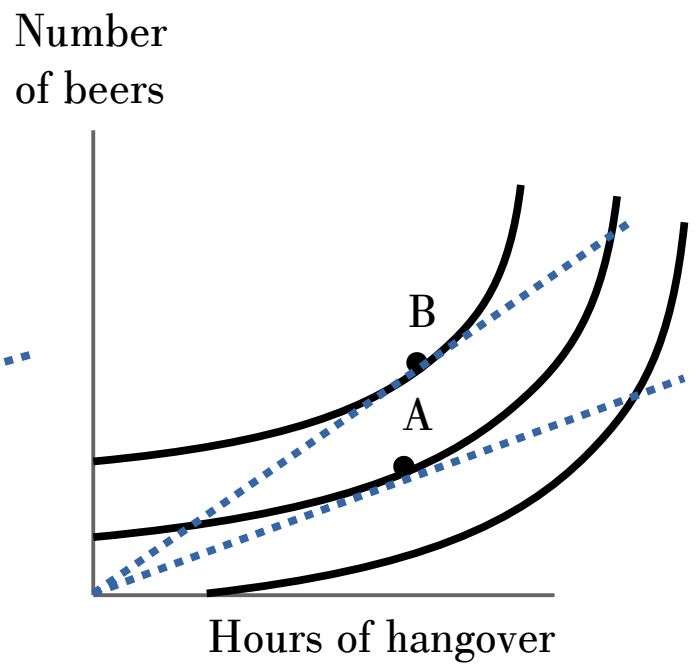
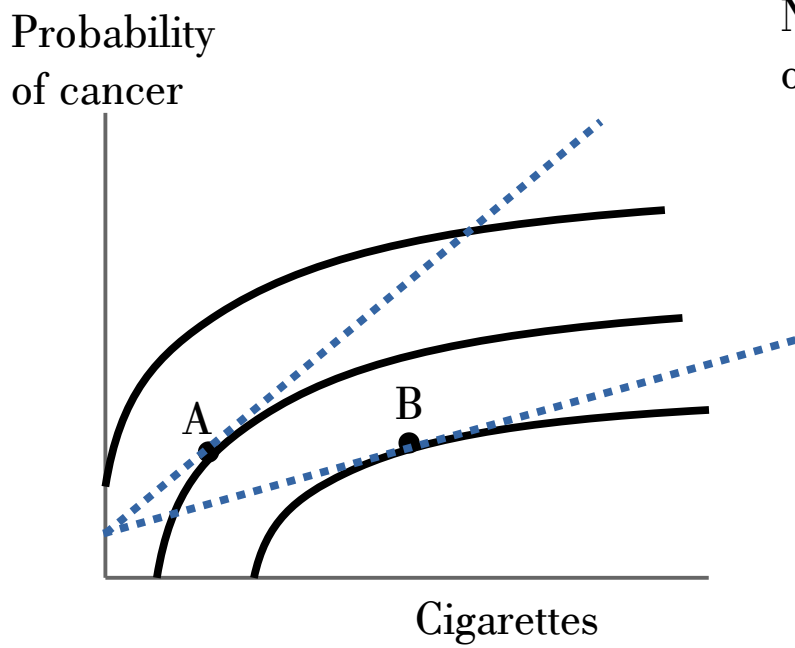


Figure 19: More examples of optimal choices with various preferences

Perfect substitutes: Optimal choices switch very suddenly as the price changes. At the highest price for Pepsi, the consumer consumes only Coke and chooses bundle A. We call this a **corner solution** since the agent chooses a consumption bundle which has as little as possible of at least one good. At the middle price when the budget line lies on top of the indifference curve C, and all the bundles on the budget line/indifference curve are equally good, affordable and optimal choices. At the lowest price for Pepsi, the consumer specializes in Pepsi and chooses bundle D. This is another corner solution.

Perfect complements: The optimal choice is always at the kinks of these indifference curves such as bundles like A, B, and C. The only time the consumer would choose allocations on the flat parts would be if the price of right or left shoes was zero, and even then, he would be no better off than if he chose to consume at the kink.

Kinky preferences: At the highest price for screws, the consumer chooses the first kink, bundle A. This is because the relative price of bolts to screws is less than 2, but greater than 1, in the example. At the middle prices (a relative price of 1), all the bundles between B and C are optimal. At the lowest price, the agent chooses the second kink, D.

Satiation: At the highest price for turkey, the consumer chooses a bundle like B but would be better off if he could afford to eat more. At the middle price, the bliss point, A, becomes affordable, and so the agent chooses it. At the lowest price, A is still affordable. In fact, it is inside the budget set and not on the budget line. Bundle A is preferred to bundle A (and every other bundle) and so again, is chosen by the agent. Note that the agent is therefore not spending his entire income. Without local non-satiation, Walras' law fails.

Goods and bads: A flatter budget line represents a “cheaper” price for cigarettes in terms of cancer risk. It might be that tobacco companies find a way to make cigarettes safer. Note that the budget constraint does not start at zero since even if you do not smoke any cigarettes, you still might get cancer for other reasons. In any event, the agent chooses bundle A and consumes very few cigarettes when they are more likely to cause cancer, and bundle C, and a much larger consumption of cigarettes when they become safer. Note that if you made a habit of inhaling cleaning solution, your cancer risk would be higher even if you kept the number of cigarettes that you smoked the same. Thus, you can choose to consume any bundle on or above the budget line, and so this describes the budget set. Of course, inhaling cleaning solution is just as stupid and suboptimal as throwing away some of your income. If preferences are monotonic, a version of Walras' law applies, and we can be certain the agent will choose to consume a bundle on the budget line and not above it.

Bads and goods: If you drink cheap beer, you may be getting extra chemicals with your brew. More expensive beer uses pure mountain spring water. Thus, you can drink more expensive beer and have the same level of hangover than if you drank less cheap beer. The lower line is the cheap beer/hangover trade-off. The upper line is the expensive beer/hangover trade off. Note that if you added shots of tequila to your beer, you would suffer an even great hangover with the same number of beers. Thus, the budget set includes all the bundles below the budget line (and so away from the direction of preference).

Non-monotonic preferences: There is nothing different here from the standard case. You simply want to find a bundle such as A that has the property that any strongly preferred bundles (the bundles above/within the indifference curve through this bundle) are not within the budget set.

Bads and bads: The main difference here is that the budget set is the area *above* the budget line. If you wanted to, you could continue to listen to JB and KCSB even after your crazy roommate untied you. The preferred set, on the other hand, is in the direction of the origin. Thus, B and C are the optimal choices at the high, middle and low prices, respectively. For example, at the lowest price, all the bundles that are preferred to B (that is, the bundles that are “above” the indifference curve through B) are not feasible since you must listen to something for the whole eight hours you are tied to the chair.

For completeness, consider the case of nonconvex preferences over goods:

White wine

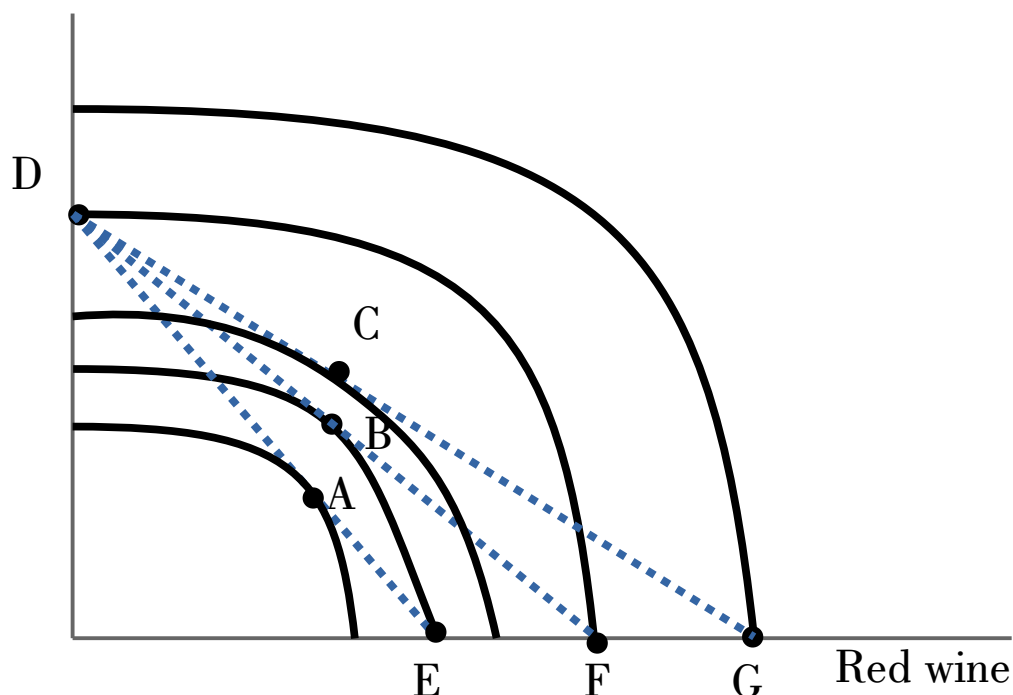


Figure 20: Optimal choices with nonconvex preferences

Nonconvex preferences over goods: We see that this looks a lot like the example above with two bads and convex preferences. The difference is that the feasible set is now *below* instead of *above* the budget line and direction of preference is now in the *positive* direction instead of *toward the origin*. Given this, the optimal choice at the highest price for red wine is bundle A. This contrasts with the case of convex preferences over bads where bundle A would have been optimal. At the middle price, bundles A and F are both optimal choices since they are on the same highest indifference curve. Bundle B would have been the optimal choice in the case of bads. Finally, at the lowest price, bundle G is optimal rather than bundle C. We see that the optimal choice is always a corner solution if preferences are nonconvex. This is similar to the case of perfect substitutes. Agents jump from specializing in the consumption in one good to specializing in the other as relative prices change. The only difference is that at the transition price (the middle price in the example above), only the two extreme bundles are optimal under nonconvex preferences while with perfect substitute preferences, all the bundles on the budget line between these extreme bundles are also optimal.

Section 3.2. Income and Substitution Effects

In this section we build on the notion of consumer optimization to begin to understand consumer behavior in the market.

We start by exploring how an agent responds to price changes. In the figure below, we fix an agent's income and the price of all goods except one (in this case, we have two goods, and we fix the price of good 2 while changing the prices of good 1). We then find the optimal consumption choice for the budget sets drawn with different prices. Linking these together gives us something we call the **price expansion path** or **PEP**.

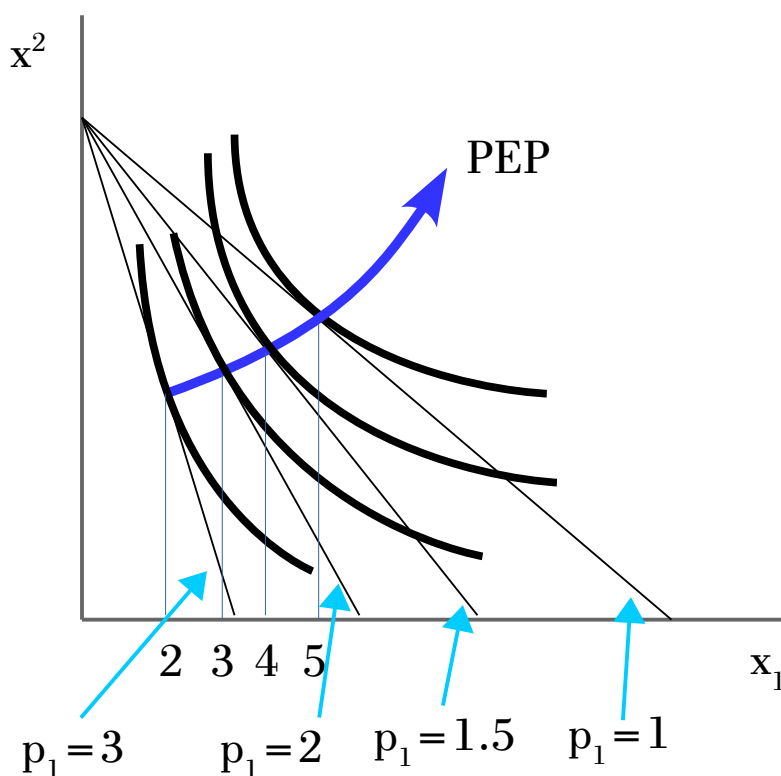


Figure 21: The price expansion path as the locus of optimal choices

From here, we can directly derive the demand curve for good 1 for the agent. For example, at a price of 1 the optimal bundle contains 5 units of good 1, at a price of 2 it contains 3 units of good 1, and so on. Graphing these combinations of price and optimal quantity gives us the demand curve shown below:

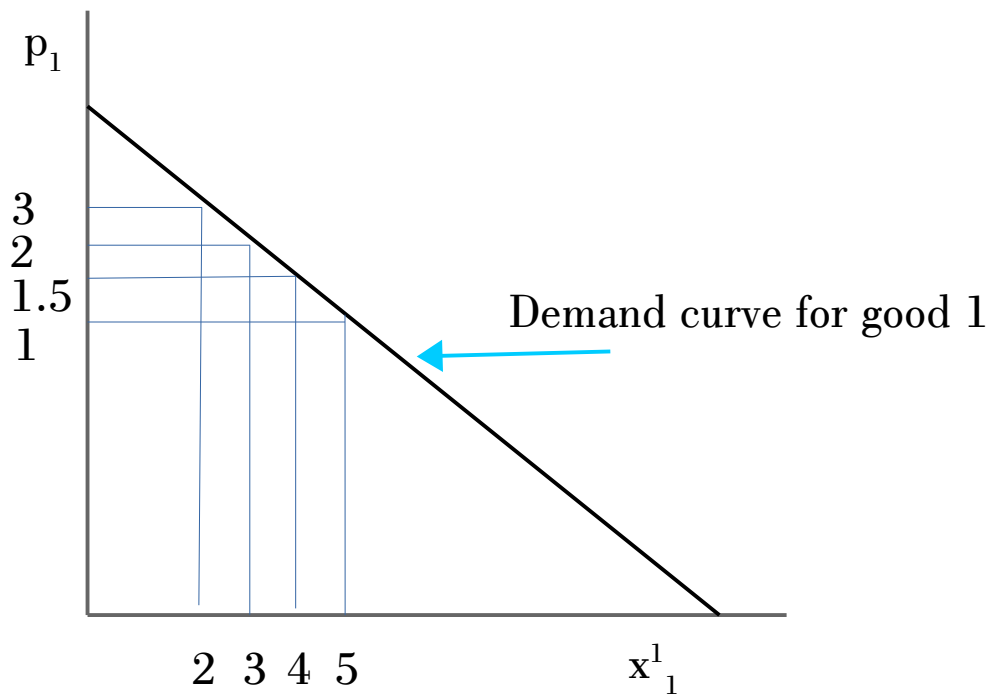


Figure 22: The ordinary demand curve as derived from the price expansion path

Note that there are two reasons that consumption changes as the price of good 1 , declines

- First, good 1 is getting cheaper. Naturally the agent wants to consume more of the relatively cheaper good: This is called the **substitution effect**.
- Second, the budget set is expanding. The agent has more choices as the price goes down. In other words, his income is worth more (even though the dollar amount is the same) and he can get to a higher indifference curve. In effect, the agent is wealthier and his consumption choice changes as a result. This is called the **income effect**.

We can separate these graphically:

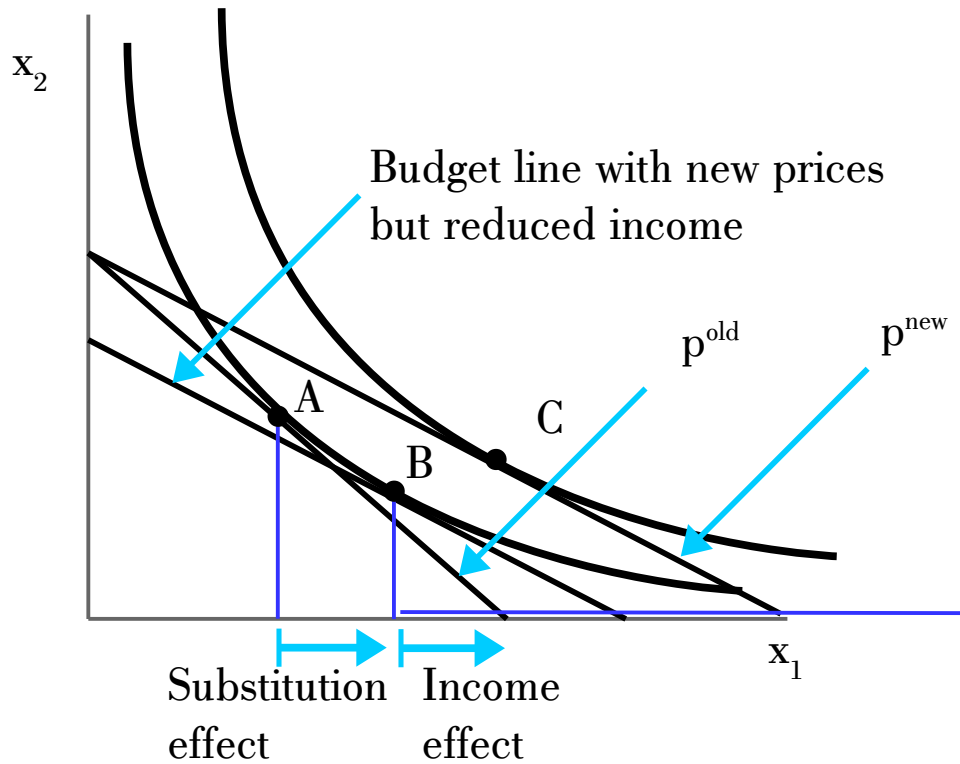


Figure 23: Income and substitution effects

As the price of good 1 goes down from the old price to the new price, the agent changes his optimal consumption choice from A to C.

To separate the income and substitution effects, we do an artificial thought experiment in which income is taken away from the agent until the old indifference curve is just barely affordable. Thus, the agent faces the new prices, but has only enough income to choose an optimal bundle on the old indifference curve through bundle A. The optimal choice in this case is bundle B. Thus, the movement from A to B is a pure substitution (or price) effect since it is motivated purely by the price change and not by any increase in welfare.

Next we consider what would happen if prices stayed at the new levels, but we added back the income we just took away in the thought experiment above. This would bring us back up to the old money income level on the new budget line. Graphically, this is a parallel upward shift of the budget line. Thus, the movement from B to C is a pure income (or wealth) effect, motivated by an increase in income or welfare, and not by any change in relative prices.

The demand curve we derived above includes both the income and substitution effects. This is called the **uncompensated** or **Marshallian demand curve**, because we do not compensate (positively or negatively) the agent for the fact that changes in price leave him better or worse off than he was initially. It is also what we actually observe in real life.

Next we consider the income and substitution effect separately. As we change income, the budget line shifts up or down, but the slope stays the same. Just as we did for price changes, we can find the optimal consumption bundle as income changes.

The **income expansion path (IEP)** is the locus of optimal choices as income expands, but all prices stay the same.

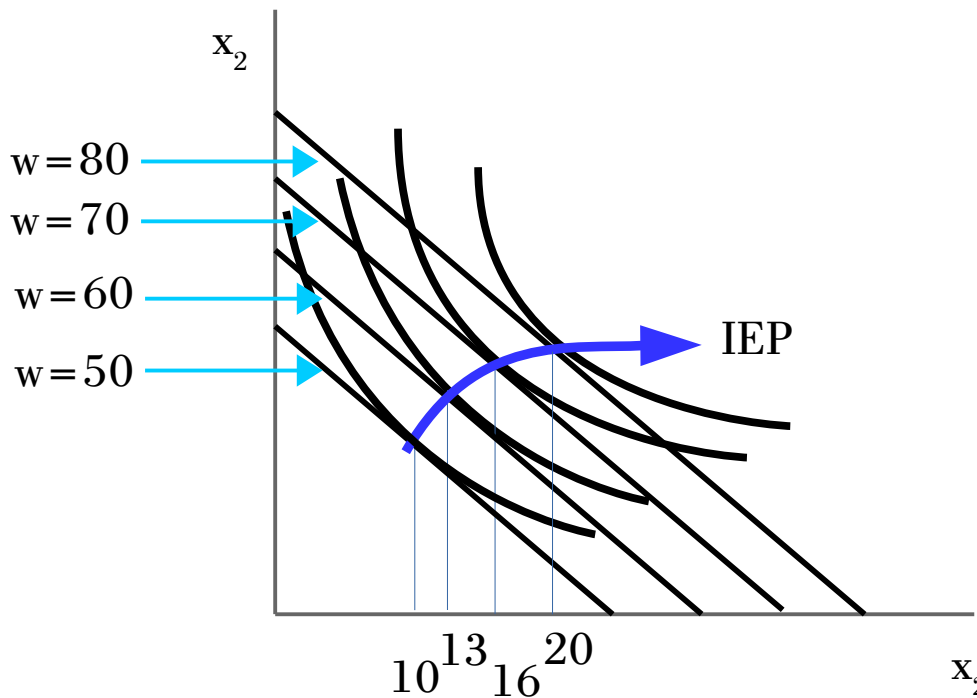


Figure 24: The income expansion path as the locus of optimal choices

We can graph the optimal quantity of good 2 (or good x_1) against income. This gives us something called the **Engel curve**.

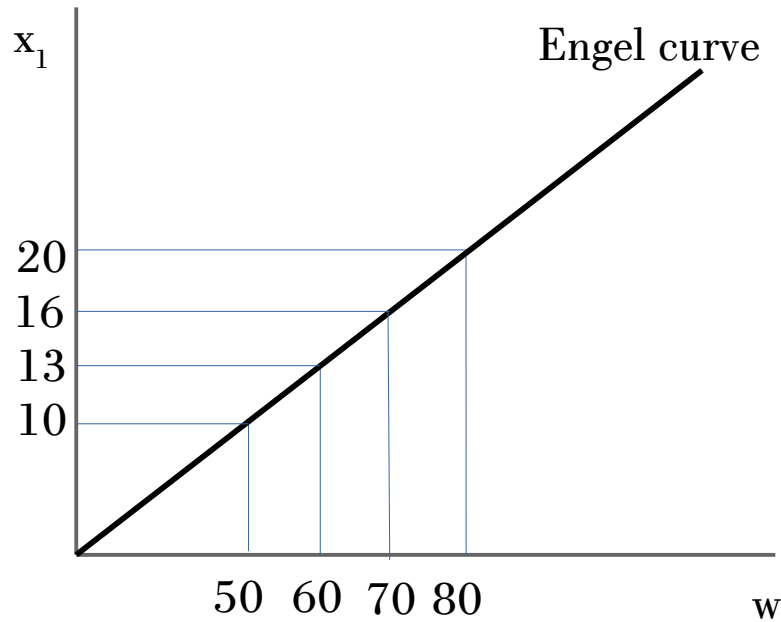


Figure 25: The Engel curve as derived from the income expansion path

Note that both the Engel curve and the IEP are upward sloping in the figures above. This means as income goes up, the consumption of each good goes up. Does this have to be true? No! The IEP may bend backwards and the Engel curve could be downward sloping.

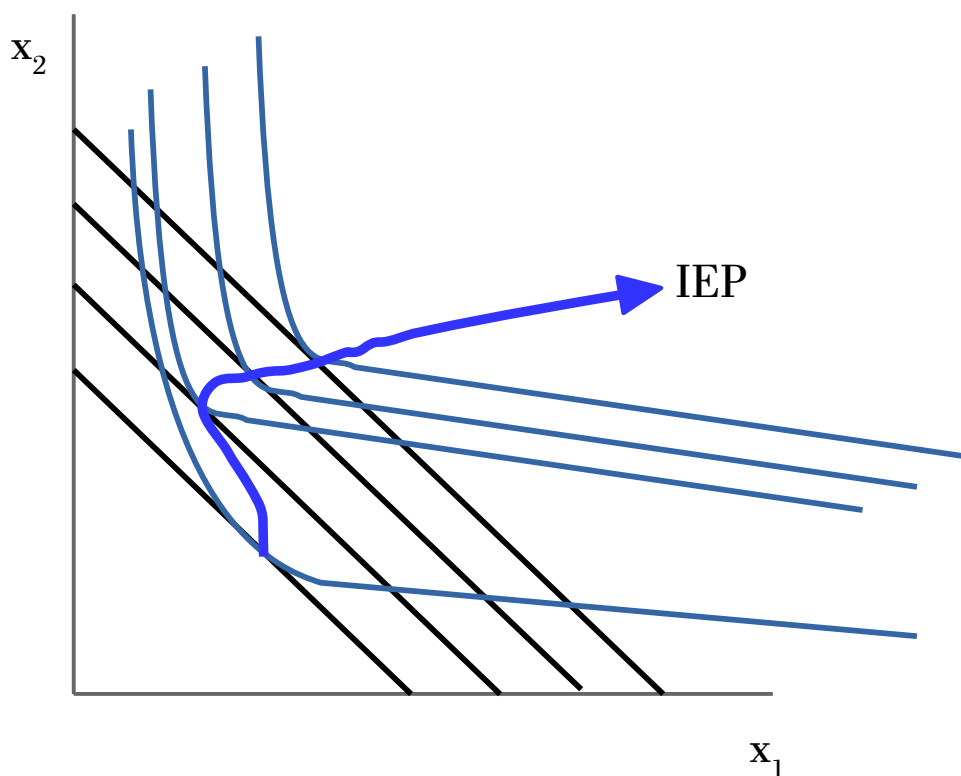


Figure 26: The income expansion path can be upward or downward sloping

Between the two lowest income levels in the figure above, the IEP slopes downward. This means as income goes up, the consumption of good 1 goes down. Consider generic beer, for example. If you get rich, you might switch to imported beer and therefore consume less generic beer. You might also consume fewer Ford Escorts, but more Cadillac Escalades, and fewer hamburgers, but more steak.

Inferior Good: A good is **inferior** if an income increase results in a decrease in quantity demanded (and an income decrease results in an increase in quantity demanded).

Normal Good: A good is **normal** if an income increase results in an increase in quantity demanded (and an income decrease results in a decrease in quantity demanded)

It is important to remember that normality and inferiority *are properties of a good at a specific set of prices and income, not a general property of a good*. Goods can go back and forth between being inferior or normal at different combinations of price and income.

When people talk loosely about a good being inferior or normal, they really mean that the good is inferior or normal at the current prices and income.

We have considered substitution and income effects together (the uncompensated demand curve) and income effects separately (the Engel curve). It remains to consider substitution effects in isolation. When we do so, we get something called the **compensated** or **Hicksian demand curve**. We call it this because it gives the quantity of good demanded at different prices *after we*

compensate the agent for the changes in utility level by either increasing or decreasing his money income. Thus, the agent expresses a demand under the new prices, but ends up back on the old indifference curve.

Note that this creates a kind of duality: Marshallian demand holds income constant but lets utility vary as prices change, while Hicksian demand holds utility constant but lets income vary as prices change.

To see this graphically, consider the figure below. The price of good 1 goes down from the old level p_1^{old} to the new level p_1^{new} . As a result, the uncompensated demand at the new prices is x_1^M . This includes both substitution and income effects. Compensating for the increase in utility that this lower price permits requires us to take away income until we find a tangency under the new relative price but on the old indifference curve. Thus, x_1^H is the quantity of good 1 chosen at the new prices along the Hicksian demand curve.

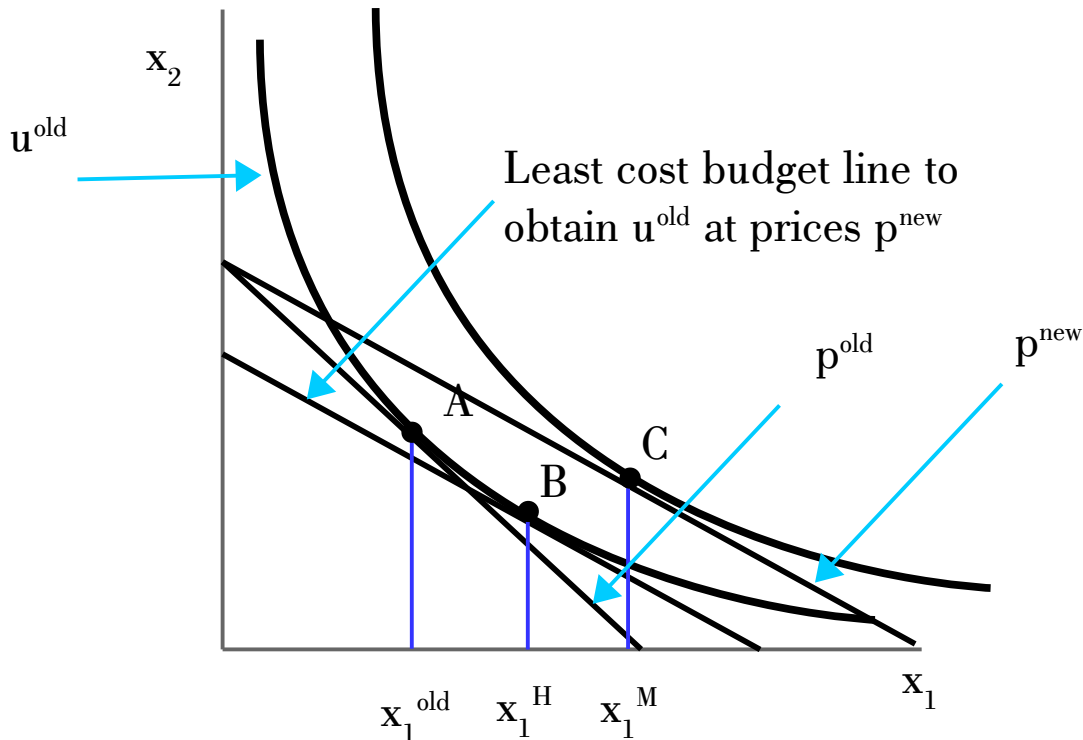


Figure 27: Compensating for income effects when prices change

The Hicksian demand curve is constructed by graphing bundle $(x^1, x^2, \bar{u})$ with the new price instead of bundle C . Since the Hicksian demand curve only includes substitution effects, Hicksian demand curves will always slope downward. On the other hand, it is theoretically possible for Marshallian demand curves to slope upward (but not very likely in real life).

You might wonder how the Hicksian and Marshallian demand curves relate to one another. The first thing to notice is that they agree at one point. We will explain why and where this happens in more detail in the next two sections. Moving away from point that they cross, which demand curve is steeper. It turns out to depend on whether the good in question is normal or inferior.

For example, suppose the price of the good goes down from old to new prices, and the good is normal. Starting at the Marshallian demand level x_1^M we use the following logic: (a) price goes down, (b) the consumer is therefore better off, (c) to compensate for this (and get the Hicksian demand level) we must therefore take away income, (d) but if we take away income, and the good is normal, the agent consumes less, (e) therefore, $x_1^M > x_1^H$.

Using similar logic we start from the Marshallian level of demand and find how the Hicksian level of demand compares as follows:

- $p \downarrow \Rightarrow u \uparrow \Rightarrow$ to compensate $w \downarrow \Rightarrow$ if normal, $x \downarrow \Rightarrow x_1^M > x_1^H$
- $p \downarrow \Rightarrow u \uparrow \Rightarrow$ to compensate $w \downarrow \Rightarrow$ if inferior, $x \uparrow \Rightarrow x_1^M < x_1^H$
- $p \uparrow \Rightarrow u \downarrow \Rightarrow$ to compensate $w \uparrow \Rightarrow$ if normal, $x \uparrow \Rightarrow x_1^M < x_1^H$
- $p \uparrow \Rightarrow u \downarrow \Rightarrow$ to compensate $w \uparrow \Rightarrow$ if inferior, $x \downarrow \Rightarrow x_1^M > x_1^H$

We summarize this in the following picture:

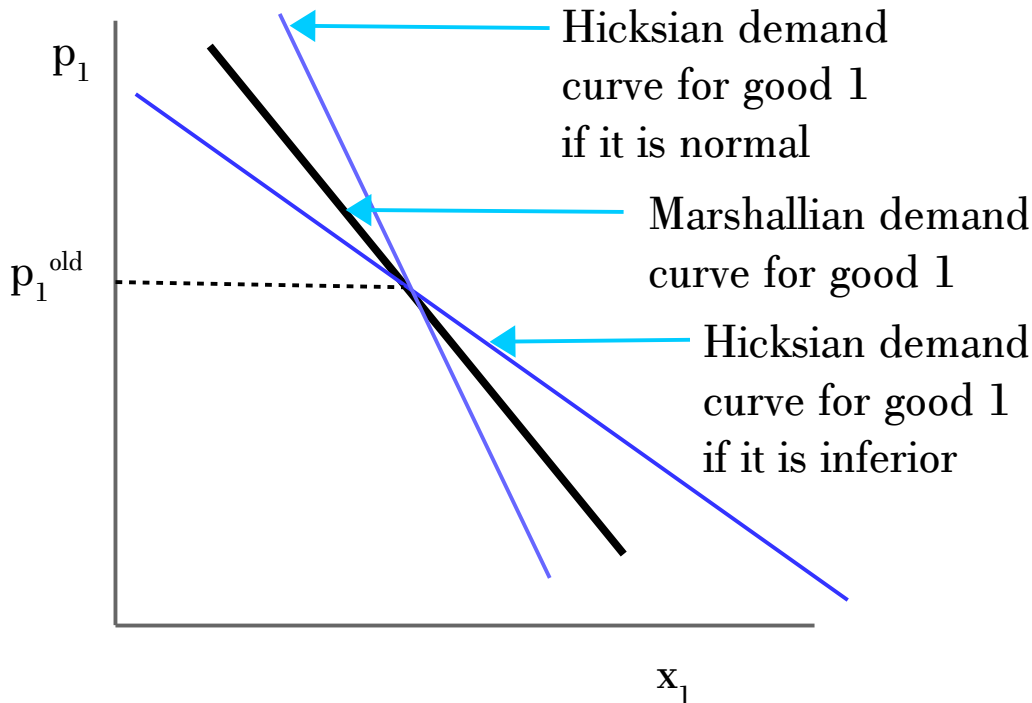


Figure 28: The relationship between the Hicksian and Marshallian demand curves

As a final summary of income and substitution effects, consider the figure below. Initially, the consumer optimizes at bundle x and $\bar{x}$: In this case, the new price of good 1 is higher than the old price. Bundle C is the new optimal choice for the consumer.

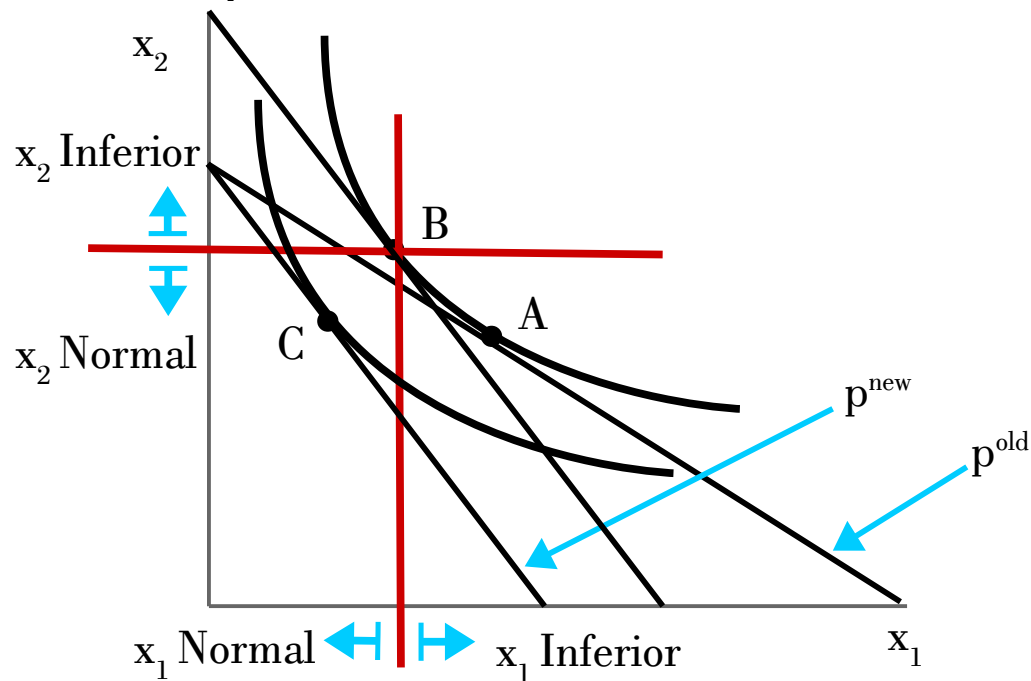


Figure 29: The direction of income effects

To isolate the income effect, the first thing to do is to remove the substitution effect. We do this by finding a bundle on the old indifference curve (through the initial choice, bundle A) that is tangent to a budget line with the new prices. In the figure, this is bundle B . Thus the movement from A to B is the substitution effect due to this price change. As expected, the consumer chooses more of the relatively cheaper good 1, and less of the relatively more expensive good 2

To see the direction of the income effect, first note that since the new budget line is everywhere below the old budget line, the consumer is worse after the prior change. This allows us to conclude the following about the income effect (that is, the position of bundle C relative to bundle B).

- If good 1 is normal, the agent consumes less good 1.
- If good 1 is inferior, the agent consumes more good 1.
- If good 2 is normal, the agent consumes less good 2.
- If good 2 is inferior, the agent consumes more good 2.

In the figure, both goods happen to be normal, and the agent chooses bundle C .

Section 3.3. Giffen Goods

We mentioned that the Law of Demand has an empirical foundation but that it is theoretically possible that demand curves might slope upwards. The figure below illustrates this strange case. As you can see, there is range of prices between p_1^{low} and p_1^{high} over which the quantity demanded increases as the price increases. Formally:

Giffen good: A good for which the demand curve slopes upwards.

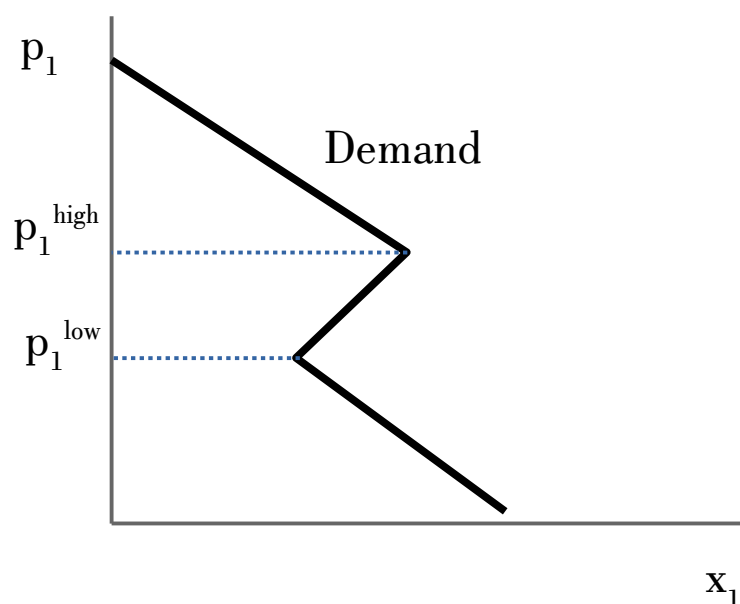


Figure 30: Giffen goods

What might drive this result? Here are some commonly proposed possibilities:

Snob effect: The price of designer clothes is high and this makes them exclusive. Could it be that as the price increases, they get more exclusive and so more people buy them? No! If this happened then they would be less exclusive since more people would buy them! (At least that is how it always works out in every observation we have.) This would mean that fewer people would want them at this higher price. Designers might make more money at higher prices, (this is an example of the Laffer Curve) but they do not sell more clothes.

Price as a signal of quality: Suppose you want to buy a high-end audio system. You listen to magnetostrictive transducer and electrostatic speakers, vacuum tube amplifiers, preamplifiers, and various digital to audio converters. It is difficult to tell which system will sound best for the widest range of music in your own media room, much less how well the components will hold up over time and integrate with new components. You get what you pay for, right? So, just buy the most expensive stuff, and you will end up with the best system! Easy! If consumers really do

take price as indicator of quality, a high price would seem to create its own demand in some sense. Can this happen? Maybe. In general, however, it turns out that more people are discouraged by the high price than are encouraged by any quality signal it might convey. If this were not the case, then we would have a price war in which all producers raised prices on their products to increase demand. The fact that we don't see this going on implies that producers know that this signaling effect is quite limited. To the extent that there are quality signals in prices, they simply reduce the degree to which demand curves slopes downward.

Income effects: Now we are getting somewhere. Although the substitution effect is always negative, income effects can go either way. When price goes up, the substitution effect causes us to consume less of the good and choose other products instead. The higher price makes us worse off which decreases our effective income. If a good is inferior, this decrease in welfare causes us increase our consumption of the good. Thus, we see the income and substitution effects of an increase (or decrease) in price working in opposite directions. If this income effect is strong enough, then a higher price could in fact lead to an increase in quantity demanded and therefore an upward sloping demand curve.

This is illustrated in the figure below for the case of a price decrease. Initially the price is p_1^{high} , and the consumer chooses bundle A. Suppose that the price decreases to p_1^{low} . The substitution effect causes the optimal choice to move from A to B. Note that B is the tangency of a budget line with the new lower prices and the initial indifference curve through A. The lower price makes the consumer better off.

When we add back the income we took away to isolate the substitution effect, the consumer ends up at bundle C. You can see the good is inferior since he chooses a lower quantity of good 1 even though he is now on a higher budget line and better indifference curve. In the figure, the income effect was so strong that it overwhelms the substitution effect. The optimal choice at lower prices (bundle C) has a smaller amount good 1 than optimal the choice at the higher price (bundle A). Good 1 is therefore a Giffen good.

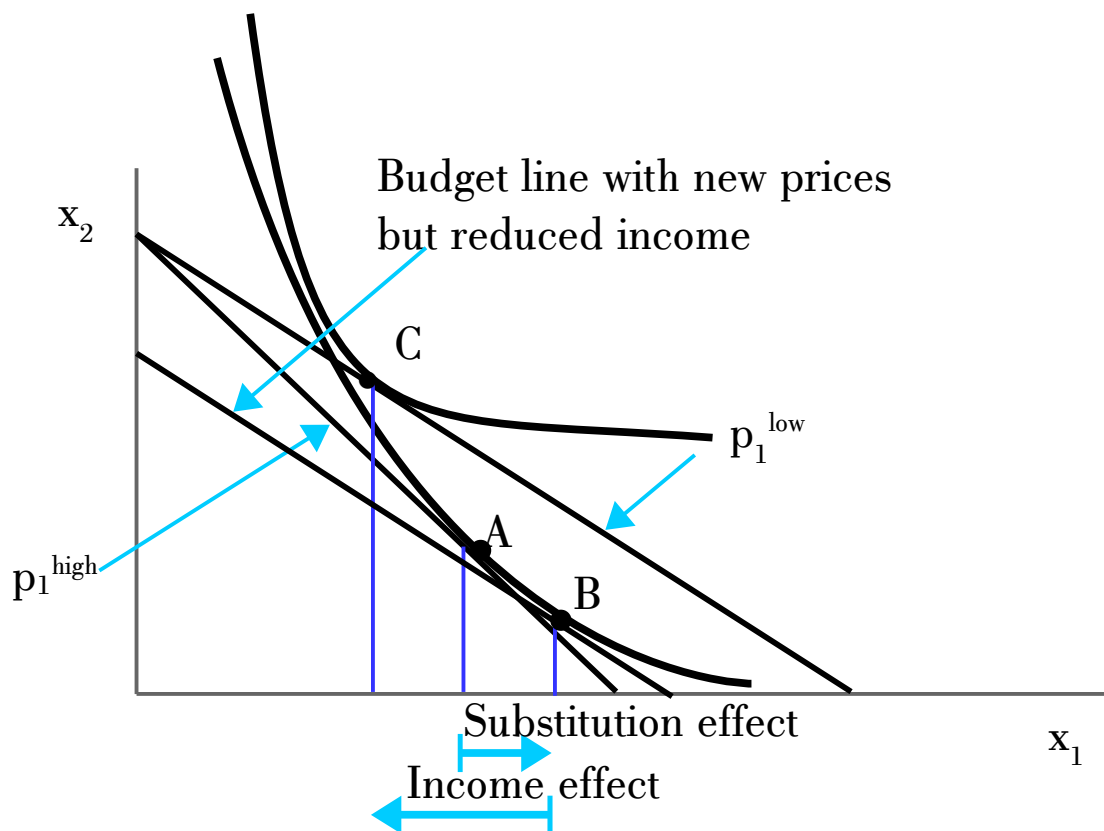


Figure 31: Giffen goods require strong negative income effects in demand

Although we cannot exclude this at a theoretical level, let's consider the circumstances that would be required for this to happen in real life. For example, suppose the price of salt went up and, that salt was a strongly inferior good. It is true that you are worse off as a result of the price increase, but how much has your effective income really decreased?

This depends on how big a part of your budget salt was in the first place. I will venture a guess that you do not spend more than of your income on salt (if you do, run, do not walk, to the nearest emergency room!). If the price doubles up, you might have to spend .01% of your income on salt. Even big changes in the price of salt have only a small impact on your budget. Thus, even if the good is strongly inferior, the change in your effective income is so small that the net effect on quantity demanded is also small and unlikely to offset the substitution effect.

Therefore, to be a Giffen good, (a) the good must be strongly inferior and (b) represent a large fraction of your spending. For example, if the cost of housing, or food, or tuition went up, you might be significantly worse off. These kinds of big ticket items are almost always normal goods, however, and therefore cannot be Giffen goods.

The one instance where economists thought they might have found a Giffen good was the case of potatoes in nineteenth century Ireland. Potatoes were a large fraction of the budget of Irish peas-

ants, but were strongly inferior since peasants replaced potatoes with bread if they could afford to. The potato famine of 1845-52 more than doubled the price of potatoes, and so we have a case of a strongly inferior good, a big change in price, and a large impact on effective income. Closer investigation, however, showed that the negative income effect was still not strong enough to overcome the substitution effect.

We conclude that all Giffen goods are inferior, but not all inferior goods are Giffen. Giffen goods are never observed in real life, it as it turns out, but they do provide a useful exercise to help students understand income and substitution effects.

Section 3.4. The Primal Problem and Marshallian Demand

The problem of the consumer is to maximize his objectives given his constraints. We showed above how the utility function is a representation of a consumers preferences, the objective he seeks to maximize. Given prices and the consumer's income, we can state his problems formally as follows:

$$\max u(x_1, \dots, x_N) \quad \text{subject to} \quad \sum_{n \in \mathcal{N}} p_n x_n \leq w$$

We call this the **primal problem** since we are maximizing the objective subject to the constraint. Below we define the associated dual problem.

In general, the consumer's problem is a constrained optimization, and we would use the method of Lagrange to solve it. Fortunately, it is possible to gain an intuition for what goes on here by considering a simple two good example which we can convert into an unconstrained maximization problem. Suppose the agent has the following symmetric Cobb-Douglas utility function:

$$u(x_1, x_2) = x_1^{1/2} x_2^{1/2}.$$

This gives us the following consumer problem:

$$\max x_1^{1/2} x_2^{1/2} \quad \text{subject to} \quad p_1 x_1 + p_2 x_2 = w.$$

Rearranging the budget constraint we find:

$$x_2 = \frac{w - p_1 x_1}{p_2}.$$

We can substitute this into the objective to produce an unconstrained optimization problem in one dimension:

$$\max x_1^{1/2} \left(\frac{w - p_1 x_1}{p_2} \right)^{1/2}.$$

Next, we take the derivative with respect to x_1 and set it equal to zero. We will have to employ both the product and the chain rule to take this derivative.

$$x_1^{-1/2} \left(\frac{w - p_1 x_1}{p_2} \right)^{1/2} - \left(\frac{p_1}{p_2} \right) x_1^{1/2} \left(\frac{w - p_1 x_1}{p_2} \right)^{-1/2} = 0.$$

Thus,

$$x_1^{-1/2} \left(\frac{w - p_1 x_1}{p_2} \right)^{1/2} = \left(\frac{p_1}{p_2} \right) x_1^{1/2} \left(\frac{w - p_1 x_1}{p_2} \right)^{-1/2}$$

Collecting terms we get:

$$x_1 \left(\frac{w - p_1 x_1}{p_2} \right) = \frac{w - p_1 x_1}{p_2},$$

and then multiplying by p_2^2 gives us

$$2 p_1 x_1 = w,$$

which allows us to conclude:

$$x_1 = \frac{w}{2 p_1}, \text{ and by symmetry, } x_2 = \frac{w}{2 p_2}.$$

Optimization without Constraints

Consider a function $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ and suppose that all the first derivatives exist. We will often be interested in finding maxima and minima of such functions.

Formally:

Local Maximum: $x^* \in \mathbb{R}^N$ is a **local maximum** of f if for all x in a small enough neighborhood of x^* , $f(x^*) \geq f(x)$.

Local Minimum $x^* \in \mathbb{R}^N$ is a **local minimum** of f if for all x in a small enough neighborhood of x^* , $f(x^*) \leq f(x)$.

Extreme Points: The union of all local maxima and minima of a function f .

Global Maximum: The largest local maximum of a function f , if it exists.

Global Minimum: The smallest local minimum of a function f , if it exists.

In the simple case where $f: \mathbb{R} \Rightarrow \mathbb{R}$ is a function of only one variable, $x^* \in \mathbb{R}$ is a local extreme point of f only if:

$$\frac{\partial f(x^*)}{\partial x} = 0.$$

This is called the **first order condition (FOC)** and is a **necessary condition** for x^* to be a local maximum or a local minimum.

You may wish to have a look at the appendix for every more fascinating details of how to optimize more general situations. [Appendix Section 2.5: Optimization](#)

We have just solved for the **Marshallian demand functions**. If we were to write the solutions to the general problem, they would look like this $\forall n \in \mathcal{N}$:

$$D_n^M(p_1, \dots, p_N, w).$$

The Marshallian demand functions give the optimal quantities of each good for any combination of prices and any level of income. The fact that the solutions we found in the example happen show that the optimal quantity of each good depends only on income and the good's own price (but not at all on the price of the other good) is an artifact special to the Cobb-Douglas form of the utility function we used. In general, the optimal quantity would depend on all prices, but in the Cobb-Douglas case, the solution shows us that it is optimal to divide spending on each good proportionally to the exponent on the good in the utility function.

The items we solved for graphically in the previous discussion turn out to be two-dimensional slices of the Marshallian demand function we just derived algebraically. In particular:

Marshallian demand curve: $D_n^M(\bar{p}_1, \dots, p_n, \dots, \bar{p}_N, \bar{w})$ – Hold income and all prices constant except good n 's. This makes the quantity of good n a function of the price of good n alone, $x_n = D_n^M(p_n)$.

Engel curve: $D_n^M(\bar{p}_1, \dots, \bar{p}_n, \dots, \bar{p}_N, w)$ – Hold all prices constant and vary income so that good n 's quantity is function of income alone. $x_n = D_n^M(w)$.

Price expansion path (PEP): $(D_1^M(\bar{p}_1, \dots, p_n, \dots, \bar{p}_N, \bar{w}), \dots, D_N^M(\bar{p}_1, \dots, p_n, \dots, \bar{p}_N, \bar{w}))$ – Hold income and all prices but p^n constant, and then form a vector that gives the optimal quantity of each of the N goods as the price of good n varies.

Income expansion path (IEP): $(D_1^M(\bar{p}_1, \dots, \bar{p}_n, \dots, \bar{p}_N, w), \dots, D_N^M(\bar{p}_1, \dots, \bar{p}_n, \dots, \bar{p}_N, w))$ – Hold all prices constant, and then form a vector that gives the optimal quantity of each of the N goods as the income varies.

We can also use these demand functions to define something new called the indirect utility function. This simply gives the maximum utility level an agent can obtain when he optimizes given some set of prices and income level. Thus, utility is an indirect function of prices and income which assumes the consumer makes optimal choice rather than a direct function of the quantities he consumes. More precisely:

Indirect utility function: $v(p, w) = u(D_1^M(p, w), \dots, D_N^M(p, w))$ is the utility of the utility maximizing consumption bundle at prices p and income w . $f \circ g(x) \equiv f(g(x))$ You can verify that the indirect utility function is homogeneous of degree zero. You may wish to look at this section of the appendix: [B.4.3: Homogeneous, Linear, and Affine Functions](#).

Algebra of Exponents, and Rules for Calculus

Multiplication Rule:

$$x^s x^t = x^{s+t}$$

Division Rule:

$$\frac{x^s}{x^t} = x^{s-t}$$

Power Rule:

$$(x^s)^t = x^{s \times t}$$

Inverse Rules:

$$\frac{1}{x^t} = x^{-t} \text{ and } \frac{1}{x^{-t}} = x^t$$

Constant Rule:

$$\frac{\partial k}{\partial x} = 0$$

Power Rule:

$$\frac{\partial x^s}{\partial x} = s x^{s-1}$$

Derivatives of Exponents:

$$\frac{\partial k^x}{\partial x} = k^x \ln(k)$$

Derivatives of Natural logs:

$$\frac{\partial \ln|x|}{\partial x} = \frac{1}{x}$$

Derivatives of Exponential Functions:

$$\frac{\partial e^x}{\partial x} = e^x$$

Product Rule:

$$\frac{\partial f(x)g(x)}{\partial x} = \frac{\partial f(x)}{\partial x}g(x) + \frac{\partial g(x)}{\partial x}f(x)$$

Quotient Rule:

$$\frac{\frac{\partial f(x)}{\partial x}}{\frac{\partial g(x)}{\partial x}} = \frac{\frac{\partial f(x)}{\partial x}g(x) - \frac{\partial g(x)}{\partial x}f(x)}{g(x)^2}$$

Chain Rule:

$$\frac{\partial f(g(x))}{\partial x} \equiv \frac{\partial f \circ g(x)}{\partial x} = \frac{\partial f(g)}{\partial g} \frac{\partial g(x)}{\partial x}$$

Recall that the MRS between two goods is the rate at which one good can be exchanged for each other in such a way that the agent is left just as well off. Put another way the MRS is the *utility rate of exchange* between two good at any point in the consumption set. We showed above that this is equal to the negative of the slope of the indifference curve through the point.

The slope of the budget lines is economic *rate of exchange* between two goods at a given set of market prices. Thus, a point of the budget line can only be optimal if the utility and economic rate of exchange are the same. Otherwise, there is an arbitrage opportunity. This implies the following:

$$\frac{\partial u / \partial x_m}{\partial u / \partial x_n} \equiv \frac{MU_m}{MU_n} \equiv MRS_{m,n} = \frac{p_m}{p_n} = - \text{ Slope of the budget line .}$$

The mathematical derivations we just completed may seem complicated at first. There is some calculus involved, after all. Let me encourage you to look upon this as really very simple instead. After all, economics is easy, right? Look at the problem again. All it says is that consumers maximize objectives (the utility function) within their constraints.

In other words, this is simply the calculus of decision-making. We do this every day. None of us can have everything we want. Money and time are scarce resources. Choosing to consume one thing necessarily means we must give up something else. If you do not understand these trade-offs, you cannot make good decisions. Balancing these trade-offs, in turn, requires evaluating how each of the available choices helps us achieve our objectives. The utility function is just an exact and formal statement of an agent's objectives.

It is probably the case that most people do not take derivatives when choosing what to eat for dinner. Most consumers do not know their utility function, much less how to apply the chain rule. However, if they did know both of these things, they should *do the math* because they would be happier as a result.

To whatever extent they can approximate this by balancing the marginal benefits and costs of different choices, they are better off. In the real world, perfect decision-making is limited by incomplete information, bounded rationality, uncertainty, and so on. These are all things that are considered in more advanced statements of the consumer's problem. They complicate the decision, but do not change its essential nature.

In short, the message here is that as, sophisticated students, you should try to think of most economic problems you face in real life (and many non-economic problems as well) as trying to find your best choice from a limited set of possibilities.

By thinking carefully about what your objectives and constraints truly are, you are likely to make better choices even if you are not able to make fully optimal decisions. The theory developed above is a clean and idealized version of what we all do in everyday life. It should be seen as a model to approximate what individual consumers actually do, but one which does seem to capture what groups of consumers do on the average.

Section 3.5. The Dual Problem and Hicksian Demand

To find the Hicksian demand function, we consider what is called the **dual problem**. Rather than fix income and maximize utility, we fix utility and minimize expenditure (E):

$$\min \sum_{n \in \mathcal{N}} p_n x_n \quad \text{subject to} \quad u(x_1, \dots, x_N) = \bar{u}.$$

Using the same Cobb-Douglas utility function from the previous section as an example, this becomes:

$$\min p_1 x_1 + p_2 x_2 \quad \text{subject to} \quad x_1^{1/2} x_2^{1/2} = \bar{u}.$$

Rearranging the constraint and substituting gives us:

$$\min p_1 x_1 + p_2 \frac{\bar{u}^2}{x_1}.$$

Substituting this into the objective gives us the following one-dimensional, unconstrained optimization problem:

$$\min p_1 x_1 + p_2 \frac{\bar{u}^2}{x_1}.$$

Taking the derivative with respect to x_1 gives:

$$p_1 - p_2 \frac{\bar{u}^2}{x_1^2} = 0 \Rightarrow p_1 = p_2 \frac{\bar{u}^2}{x_1^2} \Rightarrow x_1^2 = p_2 \frac{\bar{u}^2}{p_1}.$$

Thus,

$$x_1 = \sqrt{\frac{p_2 \bar{u}^2}{p_1}}, \text{ and by symmetry, } x_2 = \sqrt{\frac{p_1 \bar{u}^2}{p_2}}.$$

We have just solved for the **Hicksian demand functions**. In general, the solutions look like this for each good:

$$D_n^H(p_1, \dots, p_n, \dots, p_N, u).$$

The Hicksian demand functions give the cheapest consumption bundle for any combination of prices that allow an agent to achieve a specified level of utility. When we consider slices of the general Hicksian demand functions, we get the following:

Hicksian Demand Curve: $D_n^H(\bar{p}_1, \dots, p_n, \dots, \bar{p}_N, \bar{u})$ – Hold utility level and all prices constant except good n 's. This makes the quality of good n a function of the price of good n alone, $x_n = D_n^H(p_n)$.

Indifference Curve: $(D_1^H(p_1, \dots, p_n, \dots, p_N, \bar{u}), \dots, D_N^H(p_1, \dots, p_n, \dots, p_N, \bar{u}))$ – Hold utility constant and vary *all* prices, and then form a vector that gives the optimal quantity of each of the N goods for every possible price combination, and you find the indifference curve with utility level $\bar{u}$.

The Hicksian demand functions also allows us to state something new:

Expenditure Function: $e(p, u) = \sum_{n \in N} p_n D_n^H(p, u)$ – The smallest possible income needed to achieve a given utility level under given prices.

You can verify that the expenditure function is homogeneous of degree one in prices (but not fully homogeneous which would require it to be so in both prices and utility). This will turn out to be useful in evaluating the impact of social policies.

We finish this chapter with a graphical comparison of the primal and dual problem. With the primal problem, the point is to find the maximum utility u^* attainable for any given level of income $\bar{w}$. With the dual problem, the point is to find the minimum income w^* required to get any given level of utility $\bar{u}$. These problems are “dual” to each other in the sense that the solution to the primal problem is $\bar{u}$ if and only if the solution to the dual problem is $\bar{w}$.

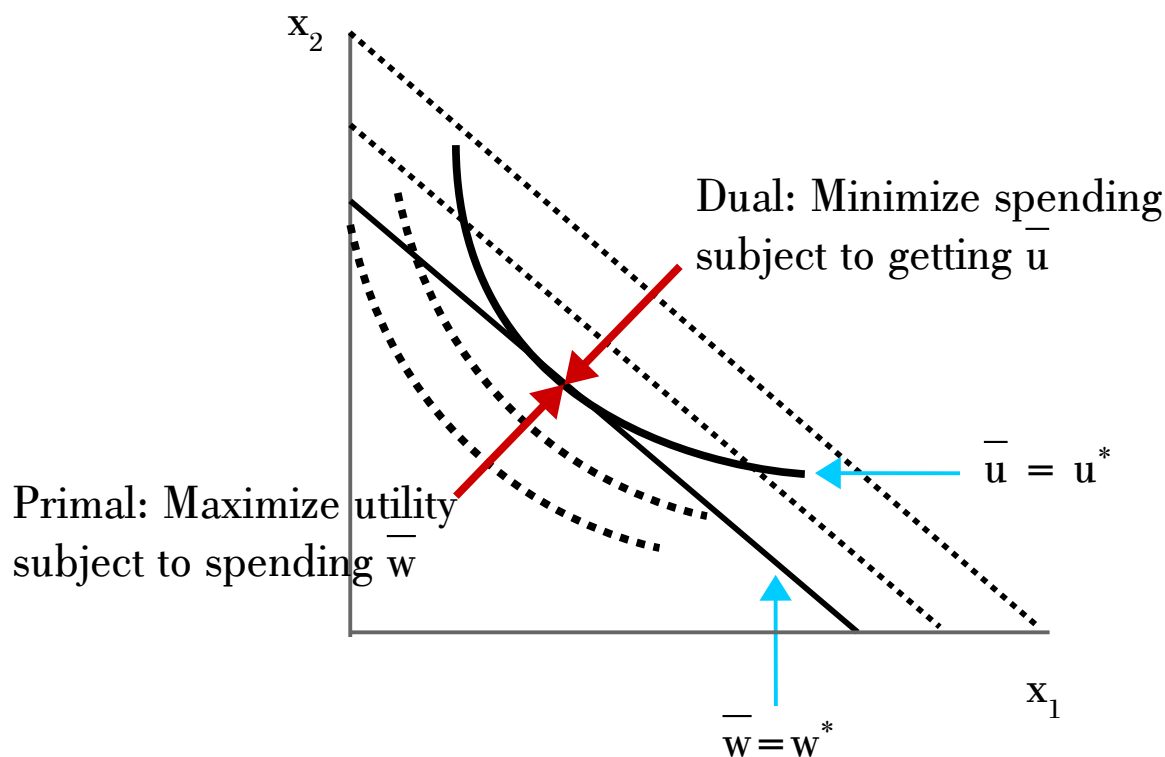


Figure 32: The primal and dual problems

Section 3.6. Elasticity

It is useful to know how sensitive consumers and producers are to changes in price. From a policy standpoint, if agents are very price sensitive then small increases in price may result in large decreases in the size of a market. From a business standpoint, knowing how price sensitive your customers are is very helpful for maximizing prices.

The most obvious candidate to measure price sensitivity is the slope of the demand curve. This is a deeply flawed measure, however. To see this, suppose that I told a firm that lowering prices 10% would double sales. Evidently, consumers are extremely price sensitive. Lowering price just a little would dramatically increase revenue. What if, instead, I told a firm that cutting prices in half would add only 10% to sales? Evidently, consumers are very price insensitive. It might even be tempting for the firm to raise price since it seems that sales would not go down very much as a result.

As an example, suppose that the demand curve is $Q = 11 - P$. Note the following:

- The slope is -1 everywhere.
- If the price starts out at $P = 2$, lowering the price to 1, that is, by 10%, would result in the quantity demanded going from 9 to 10, that is, doubling.
- If the price starts at $P = 10$, then lowering the price to 9, that is, cutting it in half, would result in the quantity demanded going from 1 to 2. that is going up by just a bit over 10%.

The point is that even though the slope is constant in the example, the degree of price sensitivity depends on the where along the demand curve we happen to be. Consumers seem to be more price sensitive at high prices and low quantities, and less sensitive as they move down the demand curve to lower prices and higher quantities.

The slope of the demand curve seems to be a poor measure of price sensitivity. A better question to ask is: what percentage change in quantity is associated with a given percentage change in price?

Thus, suppose we had the following information about a consumer's demand curve.

$$P^{old} = 10, P^{new} = 11, \text{ and } Q^{old} = 100, Q^{new} = 80.$$

Note that the change (the Δ) in the price and quantity are:

$$\Delta P = 1, \quad \Delta Q = -20.$$

To express these as percentage changes we would simply divide them by the starting values (and multiply them by 100). Thus,

$$\% \Delta P = \frac{\Delta P}{P} \times 100 = \frac{1}{10} \times 100 = 10\%, \quad \% \Delta Q = \frac{\Delta Q}{Q} \times 100 = -\frac{20}{100} \times 100 = -20\%.$$

If we take the ratio of these two, we get a number which gives *the percentage change in quantity caused by a one percent change in price (later, also income)*. This is the fundamental idea of price sensitivity which we generally call **price elasticity** and denote by the Greek letter “epsilon”.

In the example above, this gives us a price elasticity of:

$$\varepsilon = \frac{\% \Delta Q}{\% \Delta P} = \frac{\Delta Q/Q}{\Delta P/P} = -\frac{0.2}{0.1} = -2 .$$

We can understand this intuitively as “if the price of the good goes up by $\forall \alpha \in \mathbb{R}_{++}^N$, then the quantity demanded will go down by 2 %” (both from the initial values).

Actually, the equation above is not quite right from two standpoints.

First, what is defined above is actually an **arc-elasticity**, not a point elasticity. Arc-elasticity is a statement about some kind of “average” elasticity as the price changes from 10 to 11 (that is over a nonnegligible arc of the demand curve), rather than the instantaneous elasticity at a single infinitesimal point on the demand curve. The correct way to find the elasticity is to take the derivative of the Marshallian demand function with respect to price (which gives the slope at a point) instead of taking the ratio of “deltas” (which gives the average slope along a section of the demand curve).

Second, there are several kinds of elasticity (own price, cross price and income, not to mention supply elasticities). The equation above really gives the “own price arc-elasticity of demand” since it measures how the quantity of a good changes when its own price changes over an interval on the demand curve. Recall from the “law of demand”, however, that demand curves always slopes down, and so the slope is negative. Rather than carry around this negative sign for own price elasticity, it is traditional to drop it. This also has the advantage of being inconsistent and confusing, which helps keep economists employed. Thus, the true definition is the following:

$$\textbf{Own Price Elasticity of Demand: } \varepsilon_{n,n} \equiv \frac{p_n}{x_n} \left| \frac{\partial x_n}{\partial p_n} \right| \equiv \frac{p_n}{x_n} \left| \frac{\partial D_n^M(\bullet)}{\partial p_n} \right| .$$

Note that we get rid of the negative sign by taking the absolute value of the derivative, that the first superscript on the epsilon indicates the “quantity good” and the second, the “price good”, and that we are considering how quantities change as determined under the Marshallian demand curve. Applying the same idea to the price of other goods gives us:

$$\textbf{Cross Price Elasticity of Demand: } \varepsilon_{m,n} \equiv \frac{p_n}{x_m} \frac{\partial x_m}{\partial p_n} \equiv \frac{p_n}{x_m} \frac{\partial D_m^M(\bullet)}{\partial p_n} .$$

This tells us how the quantity of a good changes when the price of another good changes. Note that the superscripts on the epsilon are now different. Also, we do not take the absolute value of the derivative, so the cross price elasticity can be positive or negative. We can use the cross price elasticity to give formal definitions of two notions discussed informally above:

Substitutes: Goods m and n are called substitutes if and only if $\varepsilon_{m,n} > 0$.

Complements: Goods m and n are called complements if and only if $\varepsilon_{m,n} < 0$.

Goods like tea and coffee, train travel and air travel, or vodka and Prozac, are likely to be substitutes since they fill the same sort of role in consumption. If you drink tea in the morning, take a train to visit your in-laws, and drink vodka to get through the ordeal, you won't need coffee, air travel and Prozac to do the same job.

Goods like bologna and white bread, different seasons of *Breaking Bad*, or gold chains and polyester shirts are likely to be complements. These are goods that go better together and can be thought of as a composite good in some sense. For example, you might enjoy a bologna sandwich, but who would want to eat bologna or white bread separately? How could you watch only one season of *Breaking Bad* without desperately wanting to binge-watch to the end? If you wear a polyester shirt, you pretty much have to have a gold chain, or you ruin the effect.

We can also think about how quantities change when income changes. This is the **income elasticity** and is symbolized by the Greek letter “eta”:

Income Elasticity of Demand: $\eta_{m,n} \equiv \frac{w}{x_n} \frac{\partial x_n}{\partial w} \equiv \frac{w}{x_n} \frac{\partial D_n^M(\bullet)}{\partial w}$.

We can also use the income elasticity to make the following notions precise:

Inferior Good: Good n is called inferior if and only if $\eta_{n,w} < 0$.

Normal Good: Good n is called normal if and only if $\eta_{n,w} > 0$.

Luxury Good: Good n is called luxury if and only if $\eta_{n,w} > 1$.

We finish this section with a few observations about elasticities.

- The reason we choose to call goods for which $\eta_{n,w} > 1$ “luxury” is that when income goes up by some amount Δw , spending on the good goes up by $\eta_{n,w} \Delta w > \Delta w$, that is, the quantity demanded increases more than proportionally to the income increase.
- Can all goods be inferior? No. This would mean that as income goes up, spending on all goods goes down. Thus, if you were spending all of your income before, you would now be spending less while your income is greater. This is a violation of Walras' law.
- Can all goods be luxury? No. This would mean that as income goes up, spending on all goods goes up more than proportionally to your income. Thus, if you were spending all your income before, you are spending more than your total income now. This is a violation of the budget constraint.
- Can all goods be normal? Yes. We can go further. Suppose that all goods had an income elasticity of 1. This means that if your income goes up by 10%, your spending on each and

every good goes up by 10%. This is exactly feasible and puts you back on your new budget line. From this we conclude that an income elasticity of 1 is typical or average. In fact, there is a precise formula that says the appropriately weighted average of income elasticity over all goods must equal 1. This detail is beyond our scope, however.

- Finally, many students ask about goods that have elasticities exactly equal to 0. This is a knife-edge case, and it does not really matter what we choose to call them. Such goods might be termed “weakly inferior” *and* “weakly normal” (or “weak substitutes” *and* “weak complements”) at the same time.

Section 3.7. Price Indices

One of the more public functions of economists is to calculate the cost of living index (also called the consumer price index or CPI) or give an estimate of inflation. These are important and are used for many purposes including calculating the cost of living adjustment (COLA) to be made for social security and other pensions and payments, serving as a basis for management and labor to negotiate new contract conditions, helping to set interest rates, providing a measure of how successful the government's macroeconomic policy has been, and so on.

Let's begin by distinguishing between inflation and an increase in the cost of living.

Pure Inflation: If the absolute price of everything rises in the same proportion such that there is no change in any of the relative prices, then we have a pure inflation of the price level. Alternatively, a pure inflation can be seen as an increase in the relative price of all goods with respect to money, but with no other change in the relative prices. (Note, if prices all go down in the same proportion, we call it deflation.)

Cost of Living Index: A cost of living index attempts to measure how much money income an individual needs to maintain the same standard of living as prices change from one period to the next. Changes in both absolute and relative prices may affect the cost of living.

Clearly, if there is a pure inflation of 10% and nothing else happens, then the cost of living goes up by 10%. However, if many prices change at one time, it is unclear how much better off or worse off one is as a result. To see this, consider the following example.

Suppose that income is $w = 10$ and prices of the two goods go from $(1, 1)$ to $(2, 2)$. The figure below shows both budget constraints. As you can see, if your preferences are like the red dashed indifference curves, you chose bundle A under the old prices but bundle B under the new prices. You are worse off as a result of the price change. On the other hand, if your preferences are like the blue solid indifference curves, you chose bundle C under the old prices but bundle D under the new prices. You are better off as a result.

What is happening is that the price change has both added some choices (above bundle C) and removed some choices (below bundle A). Whether this price change makes it cheaper or more expensive to get to your old indifference curve depends on the type of preferences you have.

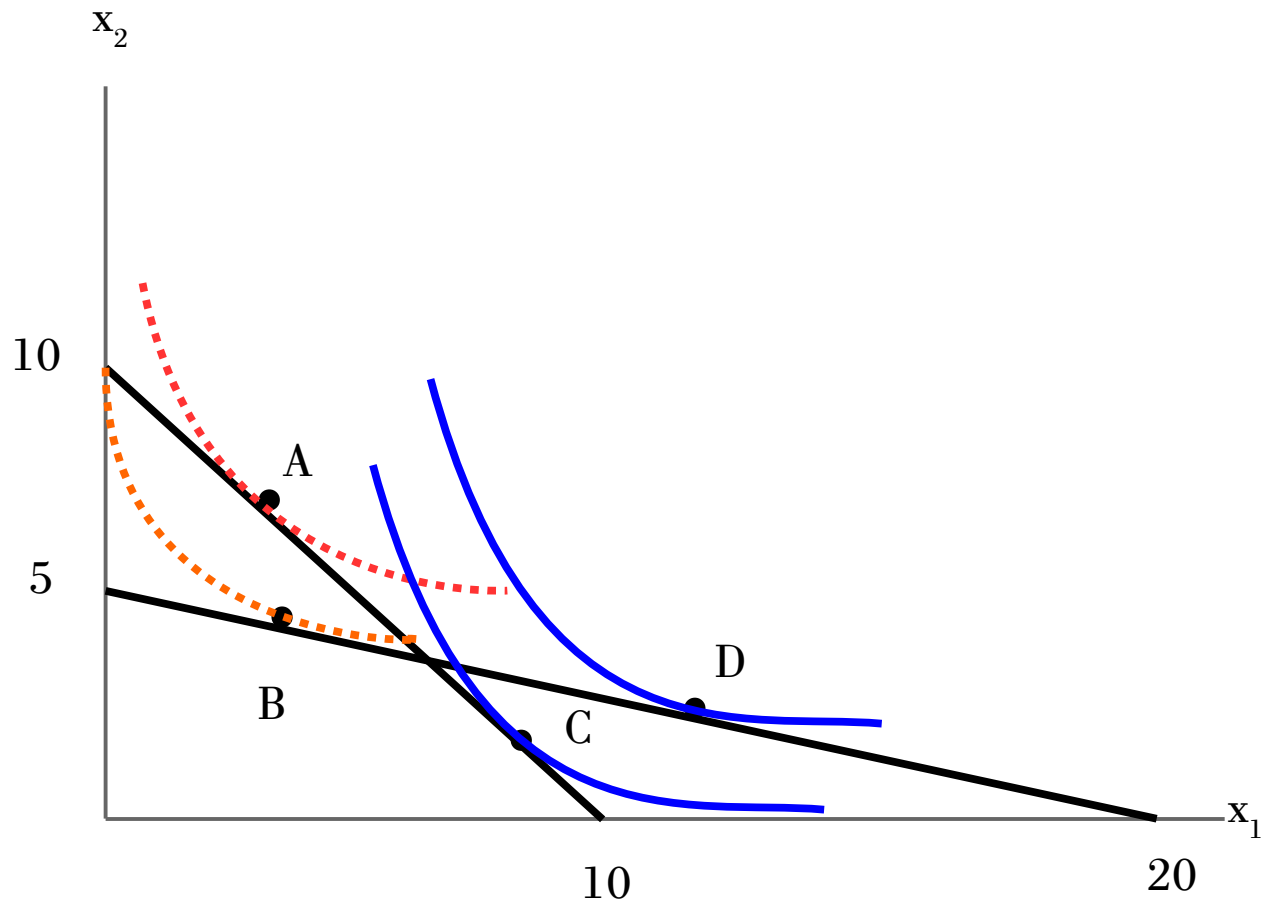


Figure 33: Changes in prices have an ambiguous effect on agents' welfare which depend upon preferences

If we wanted to calculate an **ideal price index** (IPI), we would need to find the level of income required to get back to the old indifference curve under the new prices and then divide this by the original income:

$$\text{IPI} = \frac{w^{\text{IPI}}}{w^{\text{old}}} .$$

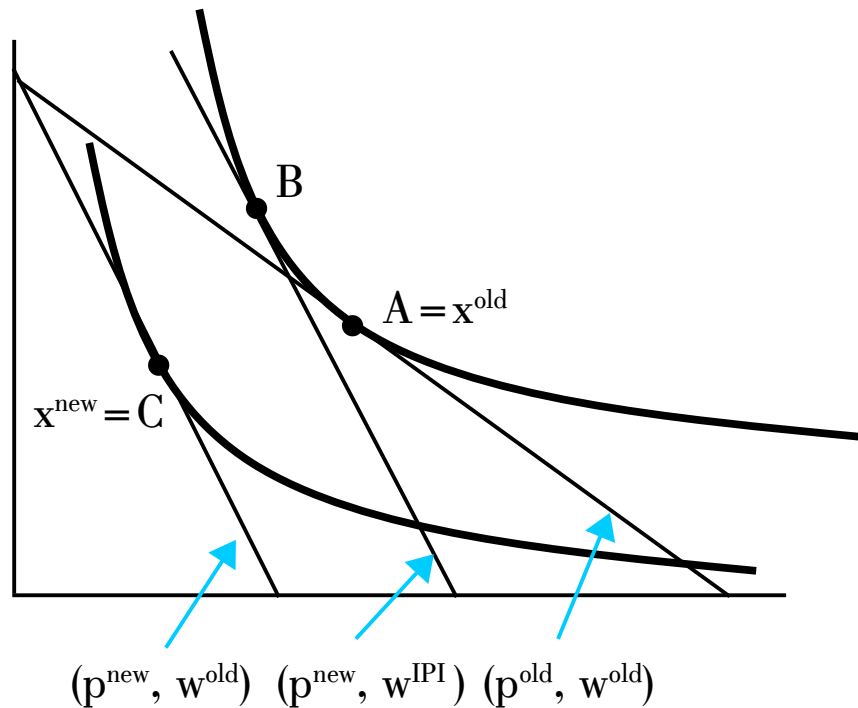


Figure 34: Perfectly compensating for a change in prices

Consider the example above. Under old prices and income, the agent chooses bundle A. The change to the new prices leaves him worse off at bundle C on a lower indifference curve. Thus, we compensate by adding income while keeping the prices at the new level until we find a tangency with the old indifference curve. In this case, it happens at bundle B, and the income required to obtain this budget line under the new prices is the exact compensation required under the IPI.

Unfortunately, unless we know the shape of the agent's indifference curves, we cannot know where this tangency might take place. Thus, the IPI is a completely non-operational idea.

Sorry to have wasted your time.

In real life we calculate the **consumer price index** (CPI) which is in the family of **Laspeyres price indices**.

$$\text{CPI} = \frac{p^{\text{new}} x^{\text{old}}}{p^{\text{old}} x^{\text{old}}}.$$

Note that the $P^{\text{old}} x^{\text{old}} = w^{\text{old}}$, that is, the denominator is just the old income, while the numerator is the cost of the old bundle under new prices. The figure below illustrates this. Instead of adding income in the impossible attempt to find the tangency at B on the original indifference curve, we add income until the old optimal choice, A, is affordable. You might notice that this

makes bundle D a feasible choice and that D is preferred by the agent to bundle A. We will say more about this below.

As an aside, what puts the CPI in the class of Laspeyres price indices that use the *optimal choice under the old prices*, x^{old} , as the reference bundle, which is costed out under both price systems. The alternative is to use a Paasche type price index. The difference is that Paasche price indices use the *optimal choice under the new prices*, x^{new} , as a reference bundle instead of x^{old} . This class of index is very seldom used in practice, so we will not discuss it further.

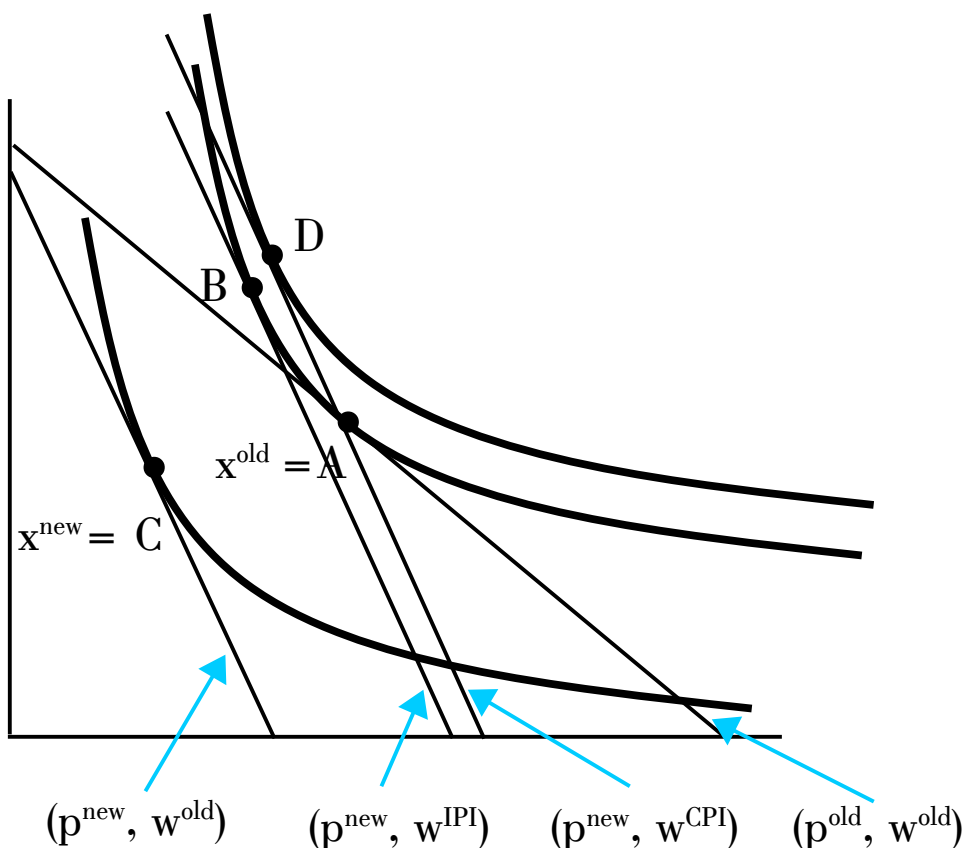


Figure 35: Laspeyres compensation makes the old bundle affordable under new prices

The CPI always over-compensates consumers for the price change. This is because of the following logic:

- Making the old bundle affordable at new prices is geometrically equivalent to pivoting the budget line to a new slope at the old tangency point, A.

- If some bundle A is an optimal choice given some budget line, and you then pivot the budget line to the new slope at the tangent bundle A, the new budget line must *penetrate* the old indifference curve.
- If a budget line penetrates an indifference curve at a bundle A, it must allow the agent to choose some bundle like D which is better than A.

The figure below shows this logic in detail.

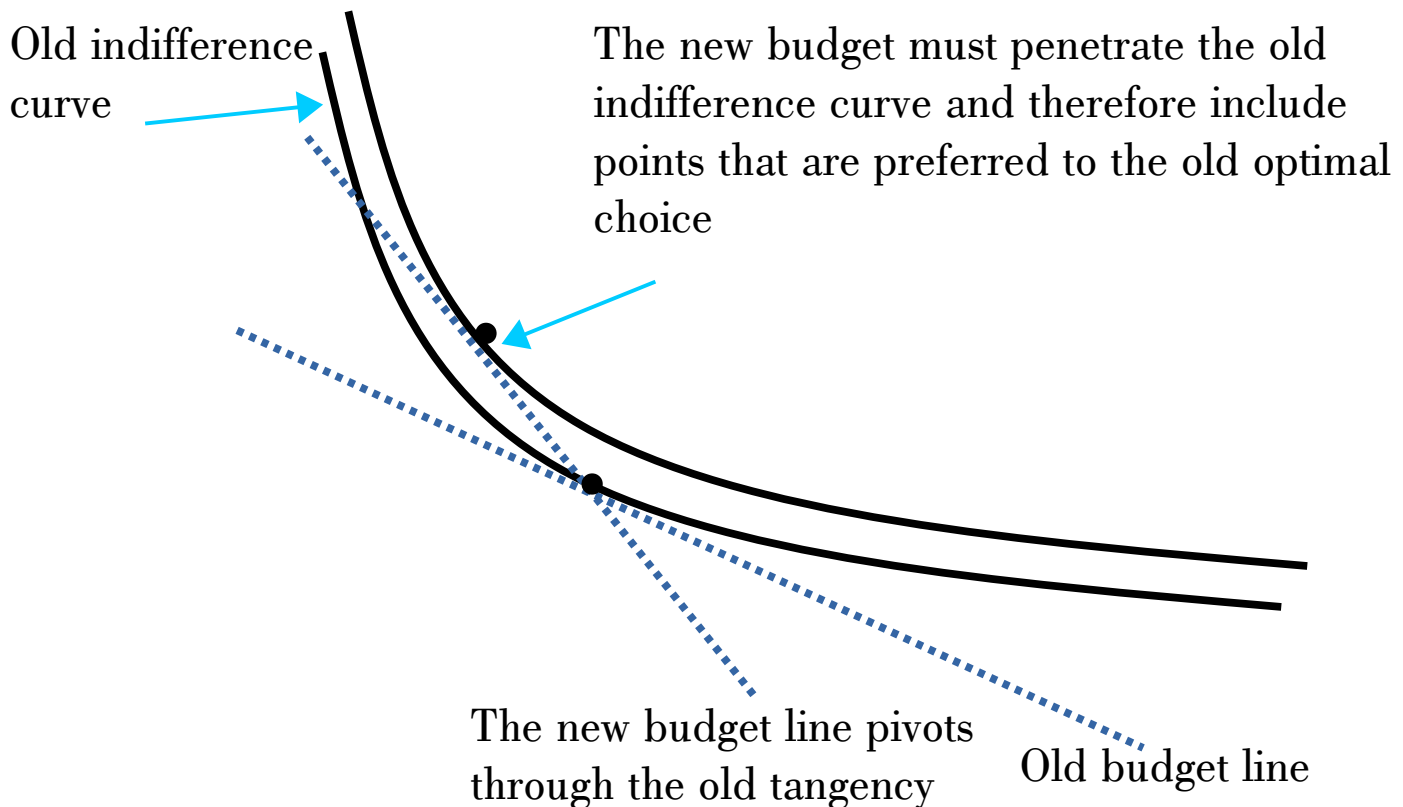


Figure 36: Agents are better off if a budget line penetrates their indifference curve

Unions members, government workers, retirees, and others whose income is partially determined by the CPI, benefit from this error of over-compensation. Each year, they are able to get to a higher indifference curve as a result. These agents are not shy about letting their congressmen and other government representatives know how much they prefer this method of calculating the COLAs to others. As a result, Laspeyres indices have solid political backing.

We agreed that finding the ideal price index is impossible. What about the CPI? The Bureau of Labor Statistics (BLS) is in charge of collecting this number. At first glance, it does not seem that complicated. We need only two things.

First, we need a list of what is in the “old bundle”. More precisely, we would need to survey consumers and ask them what they bought in a given week. What did they get at the grocery store, how many times did they ride the bus, eat out, stay in a hotel, buy a cup of coffee, watch a pay-per-view event, etc.?

Second, we need to know how much consumers paid for each and every purchase. Clearly, such data collection this is time-consuming and expensive to do. One also has to hope the consumers remember everything and report it honestly. It might seem that once we have the data, all we need to do is average the amount of each good purchased across consumers, and average the price paid for each of these the purchases. It turns out to be a bit more difficult than this. A short list of reasons follows:

What is in a consumer’s basket? Do we see two dozen eggs, or one dozen large cage-free eggs and one dozen conventional extra large eggs. Do we see two bottles of Chianti and three of Pinot Grigio, or five bottles of Italian wine? Do we count ten pound sacks of flour as a different good from five pound sacks, or do we just think of the good as pounds of flour? The point here is we have to make choices defining the set of goods that an agents purchase. If we group things together too coarsely (eggs, wine, flour) we ignore the important differences between the goods and also the wide variation in their prices. The data tend to lose their meaning as a result. If we distinguish goods too finely (each type of wine, each type of egg, or flour by brand and size of package), the data collection requirements can become overwhelming. How many consumers would we have to survey to find ten or twenty who bought a certain type of expensive car? Without a sufficiently large sample, the estimates of average price and quantity in a bundle lose their statistical significance.

What is the price of a good? A dozen wings at happy hour in a dive bar in East Paramus, New Jersey, costs much more per wing than at dinner time in a high-end gastropub in Manhattan. You can buy six wings as an appetizer, a dozen for dinner, or 100 to take home to your Super Bowl party. The price per wing will be different in each case. Thus, even for the exact same physical good, it is not clear what we should consider when figuring out the price. It would not be right to simply average the price per wing over all of these situations because we are not really talking about the same good. A wing at a dive bar, a fancy pub, or at a party at home are not really the same commodities. But again, how should we make these distinctions? Is every restaurant and bar different, are wings sold in different sized bundles different goods, or do we just average the per-unit price over all of these cases?

What do we do about technological changes? The first IBM personal computer had an Intel 8088 CPU with an 8 bit data bus running at 5 megahertz, with 16 kilobytes of RAM (memory) and sold for \$1565 in 1981. In 2015, you can get a Dell personal computer with an Intel Core i7-3770 Processor with a 64 bit data bus running at 3.4 gigahertz with 8 gigabytes of RAM for under \$1000. Do we say that a computer in 1981 went for 150% of what one goes for today? But today’s computer has 8 times the data bus, is 680 times faster, and has 500,000 times as much RAM. We get much more computer today for one third less money. We might argue that in 1981 as compared to 2015, the same computational capacity would have cost 12 times as much if data bandwidth is the measure, 1020 times as much if we used CPU speed and 750,000 times as much if we thought RAM was the right measure. So what is correct? Even if

we thought we should price a computer based on its functions or specifications, which function or combination of functions should we pay attention to?

What about new goods and obsolete goods? Some goods did not exist years ago. Others are no longer made today. How do we compare bundles across time in such cases? Twenty years ago, people bought cassette tapes and vinyl records. Now people buy or pirate MP3s. Do we look at the number of songs in the consumption baskets today as compared to yesterday? Do we count the fact that many songs today are pirated and so have a price of zero? Are MP3s not the same good as a record or cassette at all? If you have subscribed to a service that lets you access its whole library of MP3s, do we say that you own several million songs? Smartphones, in part, replace landlines, tablet computers, MP3 players (and Walkman), secretarial services, answering services, wristwatches, and so on. People still use all these component goods and services separately, however, so what do we do? Do we say that the cost per minute of a domestic call is some average of the cost of cell and landline calls? Are these different goods or just small variations of the same one? The same question could be asked about the other goods and services a smartphone replaces. Even if we knew the answers, what portion of a smartphone's cost should be apportioned to its wristwatch as opposed to the media player function?

How complete are our data? People work off the books, buy things illegally, pay cash for things, and barter. It is estimated that 25% or more of Greece's economy is underground, for example. How do we find prices and quantities there? Even in advanced economies, it is expensive and difficult to really see what people are consuming, and what prices they are paying for goods. Imagine how much more difficult it would be to collect cost of living data in developing countries or those with lots of hidden economic activity.

In short, theoretical level it is not at all clear what data we want and how we should be collecting it. If we could somehow figure this out, the CPI still would be systematically wrong in that it over-compensates for price changes. Finally, it is difficult and expensive to get the data we need so it is hard to estimate how much confidence we should put into it in any event. Still, it is better than nothing ... ?

Ignoring these difficult issues, we could also calculate other types of Laspeyres type price indices. The most prominent is the producer price index (PPI). This is exactly like the CPI except it uses an average bundle of inputs purchased by producers and tracks how the costs change over time. We could be more specific and create a price index for a state, for working mothers, for tourists, for students, etc. simply by choosing the reference bundle as the average consumption choice of people in the category of interest.

Finally, let's consider **inflation**. When all prices go up by the same proportion, the budget line shifts downward, but the slope does not change. If prices went up by 20% then obviously we must give the agent 20% more income to make the old bundle A affordable again. But this just brings the budget line back to where it used to be and makes A the optimal choice again. Thus, in the case of a pure inflation, the CPI correctly compensates for the change in prices. The Laspeyres indices really should be said to over-compensate in a weak sense. You are never worse off after having received Laspeyres compensation for prices changes, but will generally be better off. The figure below illustrates:

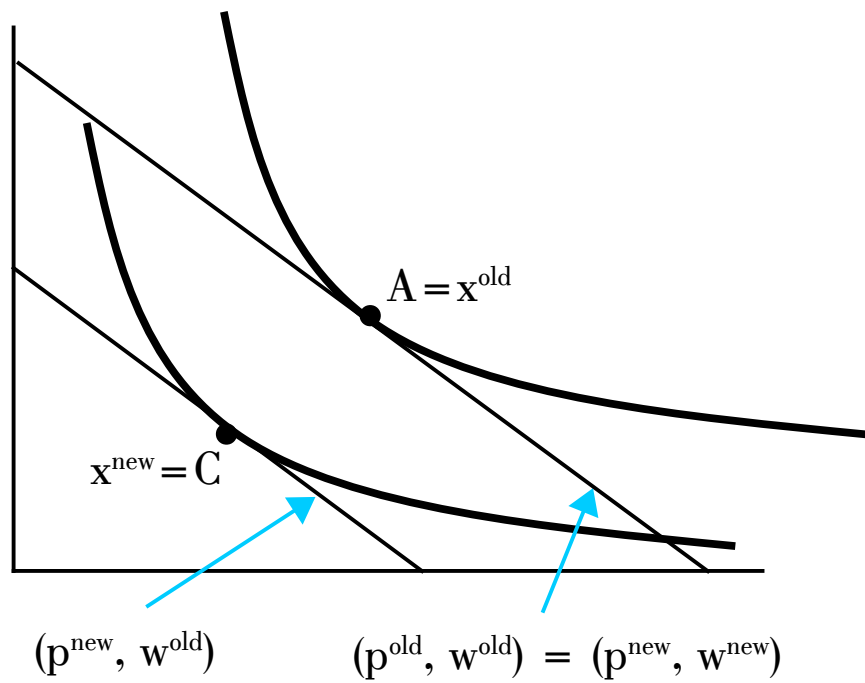


Figure 37: An example of a pure inflation

Section 3.8. Head Taxes and Income Taxes

We are all familiar with the income tax. In general, an **income tax** requires you to give a percentage of your gross income to the government. What remains is your net income, sometimes called your take-home pay. Such taxes can be **progressive**, meaning that the tax rate increases with an agent's income, **regressive**, meaning that the tax rate decreases with an agent's income, or **flat**, meaning that the tax rate stays the same at all income levels. Income taxes are distortionary in that they make choosing work over leisure less rewarding.

In contrast, **head taxes** are fixed payments that must be made regardless of any other choices made by an agent. (It may be that we call these head taxes because if you did not pay them, the king would cut off your head.) Head taxes are also called **lump sum taxes**. Income taxes affect an agent's decision of how much time to spend working both because they lower the net wage per hour which generates a substitution effect, and because they lower an agent's income, obviously generating a wealth effect. In contrast, head taxes have no effect on relative prices and so produce no substitution effects. Taking away a fixed amount of income produces only income effects.

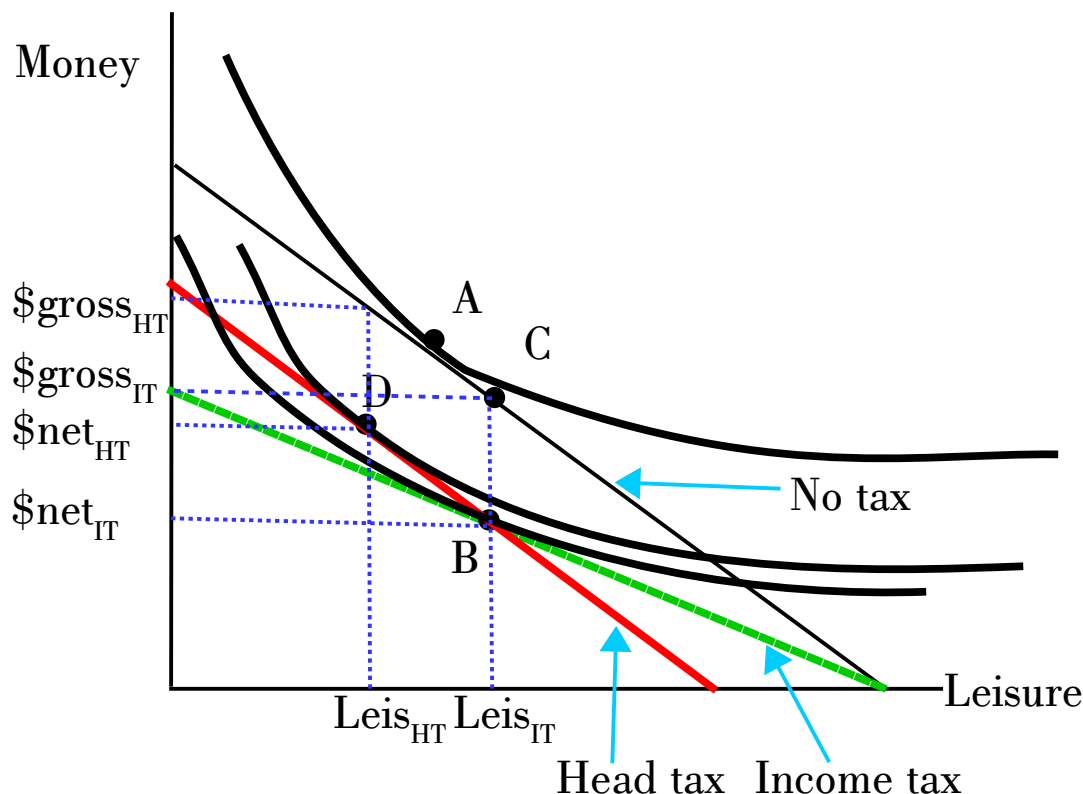


Figure 38: Head taxes vs. income taxes

Let's compare the effect of a flat income tax with a revenue equivalent head tax. The figure above illustrates.

With no taxes the agent faces the top budget line as his labor-leisure trade off. The slope of this budget line is the negative of his wage rate. We can see that the agent chooses bundle A as optimal.

If the government imposes a 40% flat income tax, the budget line pivots downward to the green dashed line and bundle B becomes the optimal choice. The net tax paid is equal to the gap between bundle C and bundle B, that is, $\text{Tax} = \$\text{gross}_{IT} - \net_{IT} .

Suppose we imposed a head tax of exactly the same amount. The budget line shifts downward everywhere by the amount of the tax. However, there is no change in the slope of the budget line since the wage rate is unchanged. The head tax is shown as the solid red line in the figure. We can see that the agent would choose a bundle like D and that this leaves him better off than he was under the income tax. Thus, it would seem that head taxes are a better way to collect revenue than income taxes.

What is going on here? Notice that bundle B is below bundle C by exactly the amount of the tax. Thus, bundle C is on both the income tax and the head tax budget lines. Bundle B, however, is a point of tangency between the income tax budget line and the agent's indifference curve. The head tax budget line goes through this point, but has a different (steeper) slope. Thus, the head tax budget line must penetrate the indifference curve through B, and therefore offers the agent some choices that are better than B, such as bundle D.

Notice that this is exactly the same geometric logic that gave us our conclusion about CPI overcompensating agents for changes in prices.

Intuitively, agents pay the same amount of tax under both systems, however, if the agent chooses to work more hours under the income tax, he only gets to keep 60% of his earning. Under the head tax, however, he gets to keep 100% of these extra earnings. This is why head taxes always make agents better off than revenue equivalent income taxes.

Why do governments almost universally use income taxes when it seems clear that head taxes are better? Most obviously, agents are different. Some are rich and some are poor. It seems unfair to make them all pay the same head tax. This problem is easily addressed, however. All we would have to do is to look at the income taxes agents paid in one year and then make this the agent's head tax for all future years.

Problem solved, right? Well, not quite. Income changes over the course of a lifetime. People graduate, retire, lose jobs, get promotions, become disabled, inherit wealth, go bankrupt, win the lottery, and so on. Thus, the lump sum taxes would soon fail to reflect an agent's economic standing.

We conclude that if a society wishes to connect taxes paid to an agent's "ability to pay", lump sum taxes will not do the job. We are stuck with some form of proportional income taxes.

Glossary

Arc-Price Elasticity of Demand: The percentage change in quantity associated with a one percent change in the price of a good over a range of prices (equivalently, over an arc of the demand curve) instead of at a single price (equivalently, at an infinitesimal point on the demand curve). Intuitively, this is a sort of average value of elasticity over an interval of prices. Formally, this is defined as:

$$\varepsilon = \frac{\% \Delta Q}{\% \Delta P} = \frac{\Delta Q/Q}{\Delta P/P}.$$

Compensated Demand Function/Curve: See Hicksian demand.

Complement: Complements are goods that are generally consumed together such as coffee and donuts or ham and cheese. More formally, goods m and n are called complements if and only if the cross price elasticity of demand is negative:

$$\varepsilon_{m,n} < 0.$$

Consumer Price Index (CPI): One of the family of Laspeyres price indices which tracks the money cost of a typical consumption bundle from year to year as commodity prices change. Formally:

$$\text{CPI} = \frac{p^{\text{new}} x^{\text{old}}}{p^{\text{old}} x^{\text{old}}}.$$

Corner Solution: A situation in which a consumer, firm, or other economic agent, finds that his optimal decision involves one or more of his choice variables taking its maximum or minimum possible values. In such a case, the agent does not make tradeoffs between the levels of all his choice variables, but instead goes to the extreme limits in some dimensions and while making tradeoffs among the remained variables.

Cost of Living Adjustment (COLA): A raise or increase in payments to agents based on the CPI. For example, each year, social security payments are increased automatically if the CPI is above a certain threshold.

Cost of Living Index: A cost of living index attempts to measure how much money income an individual needs to maintain the same standard of living as prices change from one period to the next. Changes in both absolute and relative prices may affect the cost of living.

Cross Price Elasticity of Demand: The ratio of the proportional change in quantity of one good in response to a proportional change in the price of another good given an agent's Marshallian demand function. Formally:

$$\varepsilon_{m,n} \equiv \frac{p_n}{x_m} \frac{\partial x_m}{\partial p_n} \equiv \frac{p_n}{x_m} \frac{\partial D_m^M(\bullet)}{\partial p_n}.$$

Dual Problem: In general, given a problem of maximizing an objective subject to a constraint, the dual problem is to minimize the constraint subject to the objective attaining some fixed value. In

consumer theory, the Hicksian demand function may be found by fixing an agent's utility level and then minimizing the expenditure needed to obtain this level:

$$\min \sum_{n \in \mathcal{N}} p_n x_n \quad \text{subject to} \quad u(x_1, \dots, x_N) = \bar{u}.$$

Engel Curve: A curve which shows an agent's optimal consumption level of a given good as a function of income (holding all prices constant). This is one of the two-dimensional "slices" of the Marshallian demand function.

Expenditure Function: A function which gives the minimum income required for an agent to obtain a given utility level under prevailing prices. Formally:

$$e(p, u) = \sum_{n \in \mathcal{N}} p_n D_n^H(p, u).$$

First Order Conditions (FOC): Let $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ and suppose that all the first derivatives exist. The first order conditions are defined as the following set of N equations for $n \in \mathcal{N}$:

$$\frac{\partial f(x^*)}{\partial x_n} = 0.$$

These equations are necessary conditions for a local maximum or a local minimum to be reached by the function. Thus, $x^* \in \mathbb{R}^N$ can be a local maximum or a local minimum of the function $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ **only if** x^* satisfies all the N FOC equations given above.

Flat Tax: A tax on an agent's income where the tax rate remains constant for all income levels.

Giffen Good: A good for which the demand curve slopes upwards.

Upper Half Space: Let $H_{p,k}$ be a hyperplane. Then the upper half space is defined as the set:

$$\{z \in \mathbb{R}^N \mid pz \geq k\}.$$

Lower Half Space: Let $H_{p,k}$ be a hyperplane. Then the lower half space is defined as the set.

$$\{z \in \mathbb{R}^N \mid pz \leq k\}.$$

Head Tax: A tax of a fixed amount that each agent must pay. The tax level is completely independent of consumption, production or any other choice made by the agent, and so does not distort an agent's economic decisions on the margin. Head taxes are also called a "lump sum taxes" and while they do not generate any substitution effects, they do have income effects that may affect an agents decisions.

Hicksian Demand Function/Curve: Hicksian demand functions give an agent's demand for a given good as a function of the prices of all goods and the utility level. This allows us to ask how much of a good an agent would consume if the price changed, but then we altered his income in such a way that he ended up on the same indifference curve after optimizing his consumption choice. In doing so, we isolate the substitution effect since we "compensate" him to remove any income effects from his consumption choice. The Hicksian demand curve is one of the two-dimensional "slices" of the Hicksian demand function and graphs the quantity of a given good as a function of price, holding all other prices and, in particular utility level, constant. Thus, along

a Hicksian demand curve, an agent has the same utility level, but different income levels, while along a Marshallian demand curve, an agent has the same income, but different utility levels. Hicksian demand is also called the “compensated demand”. Formally:

$$D_n^H(p_1, \dots, p_n, \dots, p_N, u).$$

Hyperplane: Let $p \in \mathbb{R}^N$ with $p \neq 0$ and $k \in \mathbb{R}$, then define the hyperplane generated by p and k as:

$$H_{p,k} \equiv \{z \in \mathbb{R}^N \mid pz = k\}.$$

In $\mathbb{R}^2$, a hyperplane is simply a line, and in $\mathbb{R}^3$, a hyperplane is an ordinary two-dimensional plane, for example.

Ideal Price Index (IPI): A fictional price index that is ratio of the “idealized income compensation level” (which is exact amount of income needed to obtain the same utility level at the current prices as an agent could under previous prices) to the income the agent had in the previous period. The IPI is fictional because finding this idealized compensation level would require knowing an agent’s indifference curves, which is, of course, impossible. Formally:

$$IPI = \frac{w^{IPI}}{w^{old}}.$$

Income Effect: The change in quantity demanded of a given good due purely to changes in the level of income and not to changes in relative prices.

Income Elasticity of Demand: The ratio of the proportional change in quantity of one good in response to a proportional change in the income level given an agent’s Marshallian demand function. Formally:

$$\eta_{m,n} \equiv \frac{w}{x_n} \frac{\partial x_n}{\partial w} \equiv \frac{w}{x_n} \frac{\partial D_n^M(\bullet)}{\partial w}.$$

Income Expansion Path (IEP): The locus of optimal choices in consumption space as income expands, but all prices stay the same.

Indirect Utility Function: A function which gives the maximum utility level achievable by an agent under given prices and income level. Formally:

$$v(p, w) = u(D_1^M(p, w), \dots, D_N^M(p, w)).$$

Inferior Good: A good is inferior if an income increase results in a decrease in quantity demanded (and an income decrease results in an increase in quantity demanded). Note that this a property of a good at each specific price and income combination, and not a general property of a good. Formally, good n is called inferior if and only if the income elasticity of demand is negative:

$$\eta_{n,w} < 0.$$

Laspeyres Price Indices: (LPI): A class of price indices that uses an average bundle consumed in the previous period as a reference to cost out under new prices. The ratio of the cost of this reference bundle under new as compared to old prices is the LPI. A seldom used alternative is called the “Paasche price index” which uses an average of this period’s consumption bundle as a reference bundle:

$$LPI = \frac{p^{new} x^{old}}{p^{old} x^{old}}.$$

Luxury Good: A good is a luxury if an income increase results in a more than proportional increase in quantity demanded. Note that this is a property of a good at each specific price and income combination, and not a general property of a good. Formally, good n is called a luxury if and only if the income elasticity of demand is greater than 1:

$$\eta_{n,w} > 1.$$

Marginal Utility: The rate at which “utility” changes when consumption of a given commodity changes. More formally, the partial derivative of a utility function with respect to any commodity is:

$$MU_n \equiv \frac{\partial u}{\partial x_n}.$$

Marshallian Demand Function/Curve: Marshallian demand functions give an agent’s demand for a given good as a function of the prices of all goods and the income level. The Marshallian demand curve is one of the two-dimensional “slices” of the Marshallian demand function and graphs the quantity of a given good as a function of its own price holding all other prices and the income level constant. This is also called the “uncompensated demand” because we do not compensate (positively or negatively) the agent for the fact that changes in price leave him better or worse off than he was initially. It is also called the “ordinary demand” since it is what we actually observe in real life. Formally:

$$D_n^M(p_1, \dots, p_N, w).$$

Normal Good: A good is normal if an income increase results in an increase in quantity demanded. Note that this is a property of a good at each specific price and income combination, and not a general property of a good. Formally, good n is called normal if and only if the income elasticity of demand is positive:

$$\eta_{n,w} > 0.$$

Ordinary Demand Function/Curve: See Marshallian demand.

Own Price Elasticity of Demand: The absolute value of the ratio of the proportional change in quantity of one good in response to a proportional change in its own price given an agent’s Marshallian demand function. Note that own price elasticity is always given as a positive number despite the fact that the derivative of the demand function is negative (except in the never observed case of a Giffen good). Formally:

$$\varepsilon_{n,n} \equiv \frac{p_n}{x_n} \left| \frac{\partial x_n}{\partial p_n} \right| \equiv \frac{p_n}{x_n} \left| \frac{\partial D_n^M(\bullet)}{\partial p_n} \right|.$$

Price Expansion Path (PEP): The locus of optimal choices in consumption space as the price of one good changes and all other prices and income stay the same.

Primal Problem: In general, a primal problem is to maximize (or minimize) an objective given a constraint. In consumer theory, the Marshallian demand function may be found by maximizing an agent's utility level subject to not spending more than his income under prevailing prices. Formally:

$$\max u(x_1, \dots, x_N) \text{ subject to } \sum_{n \in \mathcal{N}} p_n x_n \leq w$$

Progressive Income Tax: A tax on an agent's income where the tax rate increases as income increases.

Pure Inflation: If the absolute price of everything rises in the same proportion, so there is no change in any of the relative prices, then we have a pure inflation of the price level. Alternatively, a pure inflation can be seen as an increase in the relative price of all goods with respect to money, but with no other change in the relative prices.

Regressive Income Tax: A tax on an agent's income where the tax rate decreases as income increases.

Support: Let $p \in \mathbb{R}^N$ with $p \neq 0$ and $k \in \mathbb{R}$ generate the hyperplane:

$$H_{p,k} \equiv \{z \in \mathbb{R}^N \mid pz = k\}.$$

Then the set $S \in \mathbb{R}^N$ is said to be **supported** at point $x \in S$ by $H_{p,k}$ if:

$$px = k \text{ and } \forall z \in S, py \leq k.$$

That is, the set S is contained the lower half space of $H_{p,k}$ but touches the hyperplane at point $x \in S$. Another way to say this is that $H_{p,k}$ is tangent to S at x .

Separation: Let $p \in \mathbb{R}^N$ with $p \neq 0$ and $k \in \mathbb{R}$ generate the hyperplane"

$$H_{p,k} = \{z \in \mathbb{R}^N \mid pz = k\}.$$

Then two sets $S, T \subset \mathbb{R}^N$ are said to be **separated** by $H_{p,k}$ if

$$\forall x \in S, px > k \text{ and } \forall y \in T, py < k.$$

That is, the set S is contained in the upper half space and T is contained in the lower half space of $H_{p,k}$. In more than two dimensions, lines of separation are called separating hyperplanes.

Substitution Effect: The change in quantity demanded of a given good due purely to changes in its relative price and not to changes in income or utility levels. Substitution effects are seen as agents move along a given indifference curve to bundles with different slopes reflecting different relative prices of goods.

Substitutes: Substitutes are goods that are generally consumed in place of one another such as coffee and tea, or sandwiches and pizza. More formally, goods m and n are called substitutes if and only the cross price elasticity of demand is positive:

Tangency: A line is tangent to a curve at a point if the line just touches the curve at that point without penetrating the curve (at least in a sufficiently small neighborhood).

Uncompensated Demand Function/Curve: See Marshallian demand.

Walras' Law: If a consumer has monotonic preferences, then at an optimal consumption choice, $x \in X$, $p \cdot x = w$. That is, agents with monotonic preferences will always spend their entire income at an optimal consumption choice.

Problems

1. Joe Sixpack hates work, but loves money. For him, work is a bad, but money is a good. He can work as many as 40 hours per week at a wage rate of \$ 20 per hour, and also gets a small pension of \$100 per week from the Army.
 - a. Draw a picture with work on the x-axis and money on the y-axis. *Note that WORK, which is a bad, is on the x-axis, not leisure, which is a good.* Now draw in a budget constraint for a typical week that incorporates the conditions above. Draw convex preferences and show his optimal choice.
 - b. Suppose that Joe's Army pension increases to \$ 200 per month and that money is a normal good. True, false or uncertain: Joe will definitely work less. Show this in a picture.
 - c. Suppose that Joe has strictly monotonic preferences for these commodities in the sensible directions. Write a formal definition of Strong Monotonicity that accounts for this when the consumption set includes only these two commodities.
2. Suppose the agent has *strictly* convex indifference curves. Is it ever possible that there could be two optimal consumption points? If your answer is no, give your reasoning. If your answer is yes, draw an example and indicate what is special about this problem that makes this possible.
3. People who are eligible for the food stamp program have the right to buy up to a certain number of food stamps for a fraction of their face value. Chuck Roste is a recipient who is allowed to buy up to \$100 worth of food coupons for half the face value. In effect, the food stamp program gives him a \$100 subsidy on the first \$200 of food he buys. Thus, if he bought \$50 worth of food, only \$25 would come out of his pocket, and the rest would be paid for by the federal government. *Please remember that although Chuck is eligible to buy up to \$200 worth of stamps, he has the option of buying less if he prefers.*
 - a. Suppose that Chuck's income is \$400. Draw Chuck's budget constraint for dollars spent on food vs. dollars spent on all other goods before he signs up for the food stamp program. Be sure to label your axes!
 - b. Now draw Chuck's budget constraint *with* the food stamp program.
 - c. Suppose that under the food stamp program, Chuck chooses to eat only \$150 worth of food. Would he be made better off, worse off, or remain just as well off, if the government replaced the food stamp program with a direct \$100 cash grant?
4. You are a poor potato farmer who lives at the foot of the Ural Mountains. Potatoes sell for 100 rubles a kilo, and your crop consists of 1000 kilos. Consider the choice between consuming potatoes versus all other goods (AOG) and suppose that you have convex preferences. Using pictures to illustrate your argument, answer the following questions:
 - a. Graph the budget set between potatoes and AOG. Be sure to label the intersections and the slope. Show an optimal choice between these two goods.
 - b. Now suppose that you choose sell half you crop and keep half for you own consumption (and that this was an optimal choice given your preference). Suppose also that the price of potatoes

falls to 75 rubles per kilo after you sell half your crop. Are you better or worse off as result of this fall in price, or can you tell? Can you tell if potato consumption goes up or down, or does it depend on whether potatoes are a normal or inferior good? If so, how?

- c. Suppose the price fall takes place before you have a chance sell half your crop. Are you better or worse off as result of this fall in price, or can you tell? Can you tell if potato consumption goes up or down, or does it depend on whether potatoes are a normal or inferior good? If so, how?
5. A local bagel store sells bagels at \$1 each but will sell you a “bakers dozen” (13 bagels) for \$10 Suppose that you have \$12 to spend and have convex and monotone tastes.
 - a. Draw the budget constraint between bagels and all other goods. Show the optimal consumption point. Could it ever be optimal to consume exactly 11 bagels? Could it ever be optimal to consume exactly 6 bagels?
 - b. Suppose the store changes its pricing policy and starts to charge \$.90 per bagel with no quantity discount. Can you say for certain whether you are better or worse off under this pricing policy? Can you say for certain whether you will consume more or fewer bagels?
6. Suppose that there are only two goods in the world: bread and wine. Suppose that the price of wine increases, and that bread is an inferior good. Can you say for certain whether the quantity of wine will go up or down? What if bread is a normal good.
7. Felicity and Dolores Doublemint are twins. Felicity is a very happy-go-lucky person, while Dolores takes a much more serious-minded view of world. Dolores has a bank account with \$200 and also owns two ponies she enjoys riding. Felicity has a bank account with \$1000 dollars and has only one pony. It seems clear to their parents that Felicity enjoys life and riding ponies much more than Dolores. She is filled with joy when she rides, while Dolores just sits on the pony looking sad. They hire a psychic who figures out that:

$$U_F(m, p) = m + 30 p^2 \text{ and } U_D(m, p) = m + 100 p,$$

are the utility functions of Felicity and Dolores respectively, where m and p are the number of dollars and ponies each of the girls has. The parents propose giving one Dolores' ponies to Felicity on the grounds that she clearly would get more pleasure from it than Dolores. Answer and explain the following:

- a. Are both girls better off after this is done?
- b. Is there a way to make both girls better off when Dolores has to give on pony to Felicity?
- c. Is there any allocation that makes both girls better off?
- d. Can we say that Felicity is happier than Dolores or that Dolores is happier than Felicity at the allocation you identified in (c)?
8. Val Halla is a Norwegian farmer who owns 40 chicken and 20 sacks of wheat. As it happens, chicken and wheat are all he consumes as well. He goes to market one day and finds the price of chickens is \$4 and the price of wheat is \$2.

- a. At these prices he chooses to consume 20 chickens and 60 sacks of wheat. In a diagram, show Val's endowment and budget constraint. Also, sketch in his indifference curves. Assume his preferences are convex and monotonic, and of course be sure to draw them such that 20 chickens and 60 sacks of wheat is his optimal choice.
- b. Suppose Val finds that the price of chickens is \$2 and the price of wheat is \$2. Assuming wheat is a normal good, Can you say for certain whether Val will consume more wheat than if chickens cost \$4? Can you say for certain whether Val is better off or worse off with when chickens are cheaper like this.

9. Suppose that your utility function for Coke and Pepsi is given by the equation:

$$u(x_c, x_p) = 2(x_c + x_p).$$

- a. Sketch the indifference curves that come from this utility function.
 - b. Suppose that your income is \$10 and the price of Pepsi is \$1. Draw the price expansion path for Coke.
 - c. Using this price expansion path, draw the demand curve for Coke. Be sure to label any critical prices and quantities where demand behavior changes dramatically.
10. Suppose that $w=100$, and $p^a = p^i = 2$ and your demand for new Android phone as a function of the price of Androids, iPhones, and income is given by:

$$q^a = \frac{w}{4} - \frac{w^2}{2000} - 3p^a - \frac{(p^a)^2}{2} + 2p^i$$

- a. Solve for the own price elasticity of demand. What is value of the own price elasticity of demand?
 - b. Are Androids and iPhones complements or substitutes at these prices and income? Be sure to show how you know this.
 - c. Are Androids normal or inferior good at these prices and income? Be sure to show how you know this.
11. Many countries use the consumer price index (Laspeyres price index) to calculate the annual raise given to civil servants. The idea is to compensate the workers for increases in the cost of living.
- a. Define the Laspeyres price index.
 - b. Suppose that the only two goods in the world are food and clothing, and that the price of food goes up. Using a picture, show how giving raises based on the Laspeyres index actually leaves workers better off in general.
 - c. Will this always be true, or is there at least one case in which the Laspeyres index compensates workers correctly? In other words, can you think of combination of price changes such that the Laspeyres index correctly estimates the rise in the cost of maintaining the same standard of living?

Chapter 4. Producers

Section 4.1. Production Functions and Isoquants

To an economist, a consumer is really just a preference relation, a black-box that consumes goods and produces utility. Similarly, to an economist, a firm is just a technology, a black box that consumes certain goods as inputs, and produces other goods as outputs. In many ways, the problem of the firm and the problem of the consumer are very similar both conceptually and mathematically. There are also a few important differences, however, that we will outline below.

One way to describe a firm's technology is with a production function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}$,

$$f(x) = y.$$

In this form, x , is a bundle of inputs used by a firm to produce an output of a single good y . This is mathematically similar to a utility function. As such, we can also represent a firm's technology with something very similar to indifference curves. Just as an indifference curve is the set of all the consumption bundles that are equally good to an agent (that is, that give the same utility level), an **isoquant** consists of all the bundles of inputs that are equally good in the sense that they produce the same level of output. In fact, we could have called indifference curves "isoutility curves".

Formally $x, \bar{x} \in \mathbb{R}^N$ are on the same isoquant if and only if $f(x) = f(\bar{x})$. Graphically, suppose we used two inputs, capital and labor, denoted k and ℓ ,¹ and produced an output called yo-yos, denoted y :

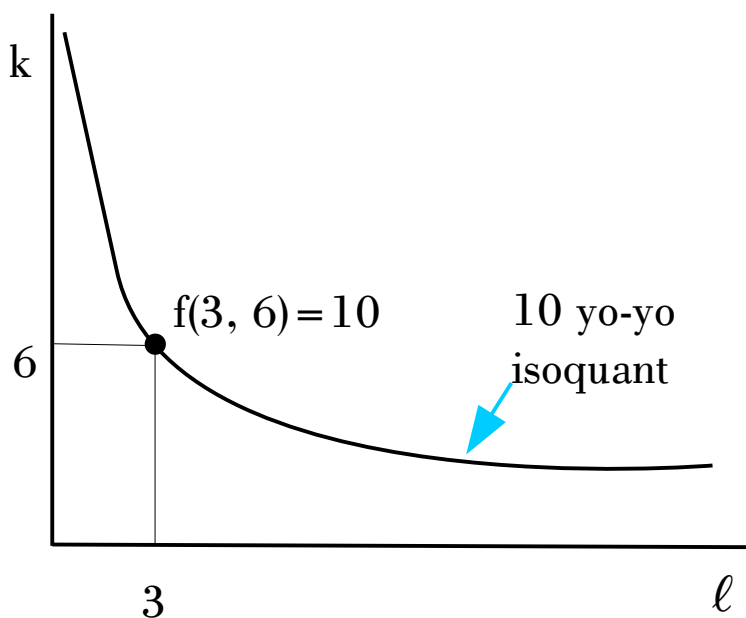


Figure 39: An isoquant

¹ Note that we use k as the subscript for capital to prevent confusion with c which might suggest cost. We also use the lower case script ℓ as the subscript for labor to prevent confusion with the Arabic number "1".

In general, we will assume that these isoquants satisfy the same assumptions we imposed on utility functions for essentially the same reasons:

Completeness: Since every input bundle is associated with some level of output, even if it is zero.

Transitivity: For obvious reasons.

Continuity: Because similar input bundles should produce similar output bundles.

Monotonicity: Since more inputs should result in more output, or at any rate, not less output.

Convexity: See the discussion below.

Although the reasons to assume the convexity of the isoquants are similar to the reasons we do so for indifference curves, it is worth elaborating a bit on this. Suppose we were trying to produce Cadillacs using machines and workers as inputs, and we could produce one Cadillac if we had 5 workers, and 5 machines.

Suppose that one worker quit. After reassigning the remaining workers to the jobs that humans are best at doing relative to machines, it might turn out that we only needed one half of a machine to replace a lost worker.

Now suppose we lost a second worker. Again, we would optimally reassign the remaining three workers, but now we would have to depend on machines to do jobs that they were less well suited to. It might therefore take an entire machine to do this work. The next worker is even harder to replace, and it might take two machines. In other words, the marginal product of workers compared to machines becomes higher as workers become relatively scarce (and lower as workers become relatively abundant). An illustration of this can be seen below:

Just as the slope of the indifference curve was called the MRS, the slope of the isoquant is called the **Marginal Rate of Technical Substitution**: $MRTS_{k,\ell}$. Thus, the MRTS is decreasing. Another way of saying this is that average bundles of inputs are more productive than extreme bundles. We conclude that isoquants are convex.

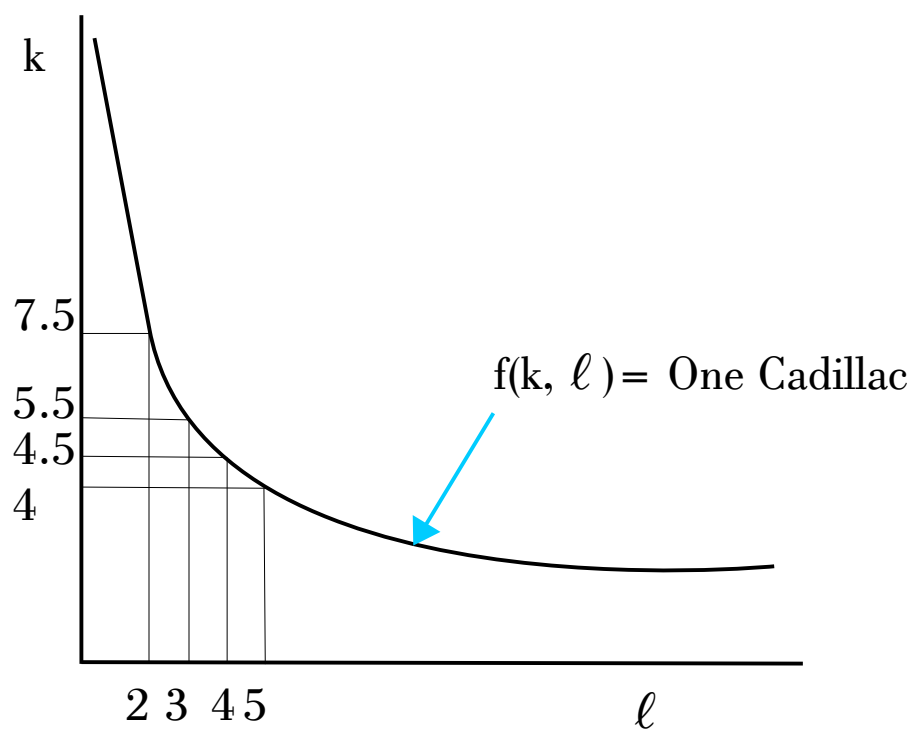


Figure 40: Convexity and DMRTS for isoquants

Section 4.2. Returns to Scale

A major difference between firms and consumers is that a firm produces a physical output. Thus, if $f(x) = 5$, then 5 actual physical units of a good are produced. If another input bundle $\bar{x}$ is used and the result is 20 units of y , then we can say that 15 more or four times as many units of output are produced. If two firms produce 20 units of output, then 40 are produced in total. In other words, the outputs of firms have cardinal meaning and so can be added and compared across firms. This contrasts with utility functions which have only ordinal meaning. A consumption bundle that gives you 20 utils rather than 5 does not make four times as happy or 20 units better off, it merely makes you happier than you were before.

The cardinality of production functions allows us to define three types of *returns to scale* that a production function might satisfy. Let $x \in \mathbb{R}^N$ be bundle of inputs, and $k > 1$ be a positive constant (confusingly, not to be confused with k , the subscript for capital).

Increasing Returns to Scale (IRS): $f(kx) > kf(x)$.

Constant Returns to Scale (CRS): $f(kx) = kf(x)$.

Decreasing Returns to Scale (DRS): $f(kx) < kf(x)$.

At first glance, CRS might seem to be an uninteresting knife-edge case. Suppose, however, that we had a giant, three-dimensional copy machine. If we used this to copy a firm and thereby exactly doubled all the firm's inputs, the result would have to be to exactly double output. Logically, nothing else could possibly happen. However, when we look at real world data, DRS is by far the most common thing we observe, with IRS turning up as well from time to time. Why is it that our real world observations differ from our theoretical predictions that CRS should be ubiquitous? The reasons are as follows:

- We typically do not really double all inputs. For example, when we double a firm's size we do not generally add a second CEO. Thus, managerial inputs are not doubled. This gives us DRS due to the implied measurement error. (Of course, having less management per worker might also increase productivity. You never know.)
- We may hire twice the labor, but these are the second-best workers not hired in the first round. We may double the amount of land we use, but the only land available is likely less desirable than what we bought the first time. Thus, measured inputs double, but the quality of the inputs is not constant. This also gives DRS due to the measurement error.
- We may have twice the capital and labor, but it is in a different form, perhaps using more efficient or specialized machines or employees. Firms change the organization of production when they increase the size of their operation to take advantage of such *economies of scale*. Thus, output may go up more than proportionally to inputs, and we will see IRS.
- Even with the same general production technology, employees might get better at doing their jobs as they get more practice. It might also be that the management or the engineering staff learn ways to streamline or otherwise improve the production process as they observe them in

action. This is called *learning by doing*. In our context, the more output a firm produces, the cheaper the costs per unit become. Thus, if you double inputs, you more than double output, and so again we see IRS.

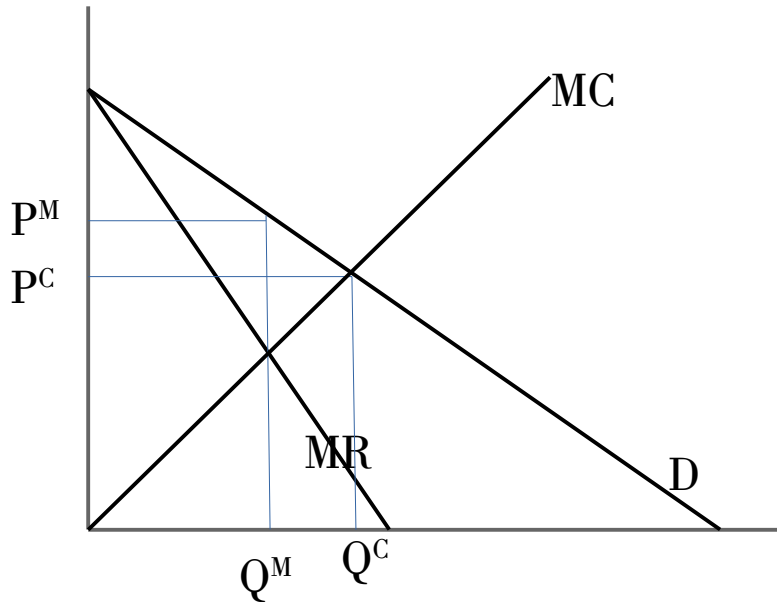


Figure 41: Monopoly optimal profit

Subsection 4.2.1. Natural Monopoly

Natural monopolies occur when an industry is characterized by high costs of initiating production, but lower incremental costs for producing subsequent units of output. Formally:

Natural monopoly: A firm is said to be a natural monopoly if the minimum of its AC curve is above the market demand curve.

For example, to build even a single new passenger aircraft, Boeing has to design it, test the model, get regulatory approval, build an assembly line, retool its machines, train its workers in the new process, and so on. Thus, if it made only one copy of the new design, it would be extremely expensive. The second unit, however, only requires the additional time and materials that go into actually producing the aircraft, so the incremental cost of the second unit is much smaller than the first. In other words, the AC of two aircraft is smaller than the AC of one aircraft, and that AC continues to decrease for a long time (maybe forever).

In the ICT space, businesses characterized by high first copy costs are often natural monopolies. Products like pharmaceuticals that involve significant research and development, or complex computer chips and devices that require extensive design and testing, also fall into this category. More generally, industries with significant economies of scale or scope, or with large “fixed costs” of beginning production are likely to be natural monopolies.

An example can be seen in the figure below. Notice that if the firm was forced to behave like a competitor, it would set choose price and quantity where marginal cost curve crossed the demand curve. This is unsustainable, however, because this “free market” price is below average cost. Thus, average cost of production less than average revenue (P^{FM}), implying the firm loses money on every unit of output. It would have to go bankrupt since profits would be negative.

On the other hand, allowing the firm to set the monopoly price results in consumers paying a higher price, for a smaller quantity. If the price and quantity were set by regulation at P^* , where $AC = D$, consumers would be better off with a lower price for a larger quantity, while the firm would just cover costs.

The problem is that a regulator such as the **Federal Trade Commission (FTC)**, or **Department of Justice (DOJ)**, has no way of knowing the shape of the AC curve, or even the whole shape of the demand curve. Indeed, the firm most likely does not know the shape of its AC curve, except locally. We are stuck in a second-best situation because of this incomplete information.

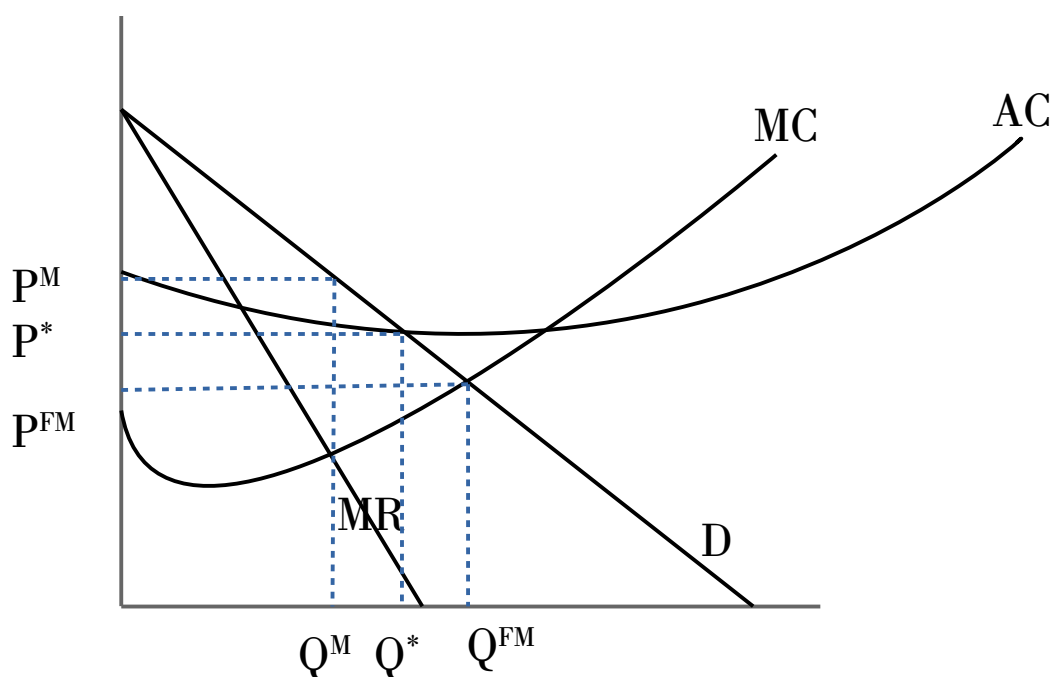


Figure 42: Natural monopoly

In the extreme, but common, case of large first copy costs and very low marginal costs, the AC may always be above the MC. The blue curves show the monopoly outcome. If we broke this natural monopoly up and insisted that two firms split production between them, then the average cost would approximately double (since twice the fixed first copy costs would have to be divided into the same total output).

In the illustration below, we have a situation where this causes the AC in the case with two firms to be above the demand curve at every quantity. This implies that there is no quantity where the competitive price could cover the average costs of production, and so the two firms would necessar-

ily make negative profits and have to shut now. Clearly the monopoly outcome is better than zero output. In all cases, it is unclear whether consumers or society benefit from breaking up natural monopolies.

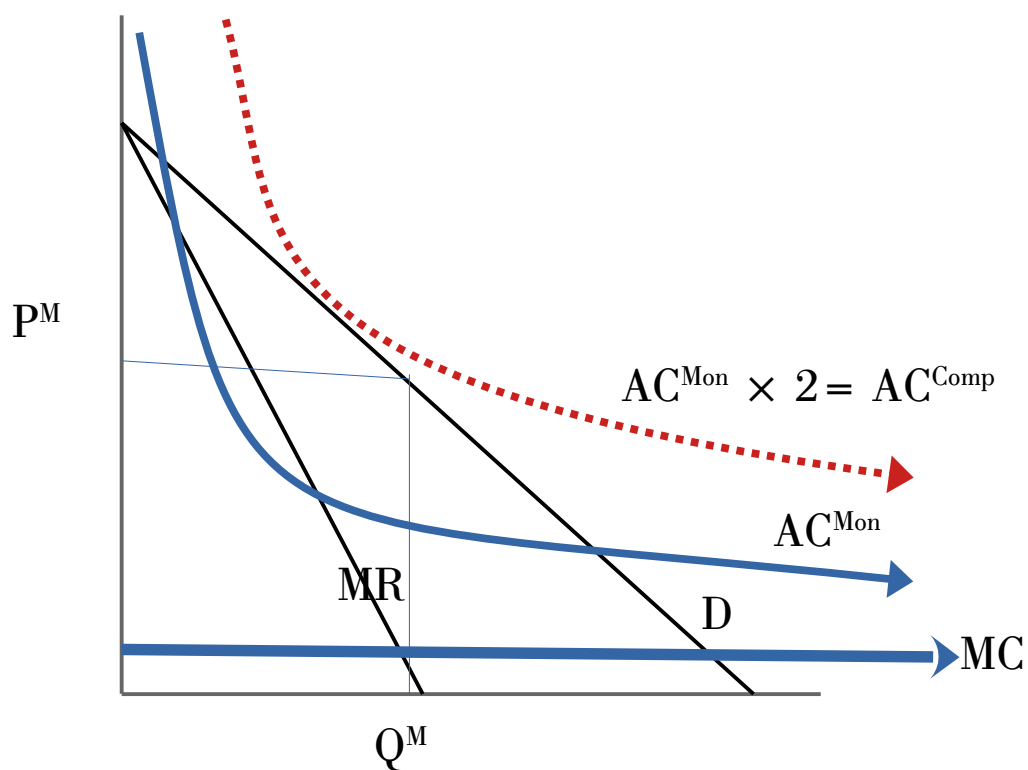


Figure 43: Natural monopoly when two firms cannot be profitable at the same time

Section 4.3. Cost Minimization

Another difference between firms and consumers is that consumers have fixed endowments which imply a fixed budget constraint. The consumer's problem is therefore to maximize utility given this fixed budget. Firms, on the other hand, have no endowments and therefore no fixed budgets. If they can profitably produce an output, they should be able to obtain financing to produce it (at least assuming perfect capital markets).

Ultimately, the objective of the firm is to maximize profits. Doing so requires two steps. First, the firm finds the least cost way of producing any specified level of output given the firm's technology and taking input prices as fixed. Second, the firm takes the minimal total cost function it just derived, and, taking output prices (or at least the demand conditions) as fixed, chooses a profit maximizing output level. These are logically separate and independent optimizations.

The cost minimization problem can be stated formally as:

$$\min p x \text{ subject to } f(x) = \bar{y}.$$

This is exactly like the “dual problem” for consumers. Recall that we can think of consumers trying to obtain a particular level of utility (that is, achieve a certain indifference curve) at the least possible cost (that is, while being on the lowest possible budget line).

Here, firms try to produce a certain level of output (achieve a certain isoquant) using the lowest cost set of inputs (here we say: while being on the lowest “isocost line”).

An **isocost line** is a set of bundles of inputs that each cost the same amount. This looks exactly like a budget line. In the figure below, the cheapest way to produce 20 units of output is to spend \$30 and use 20 units of labor and 10 units of capital, each costing \$1.²

² Note that we use k as the subscript for capital to prevent confusion with C , which might suggest cost. We also use ℓ the lower case script “l”, to prevent confusion with the Arabic number “1”.

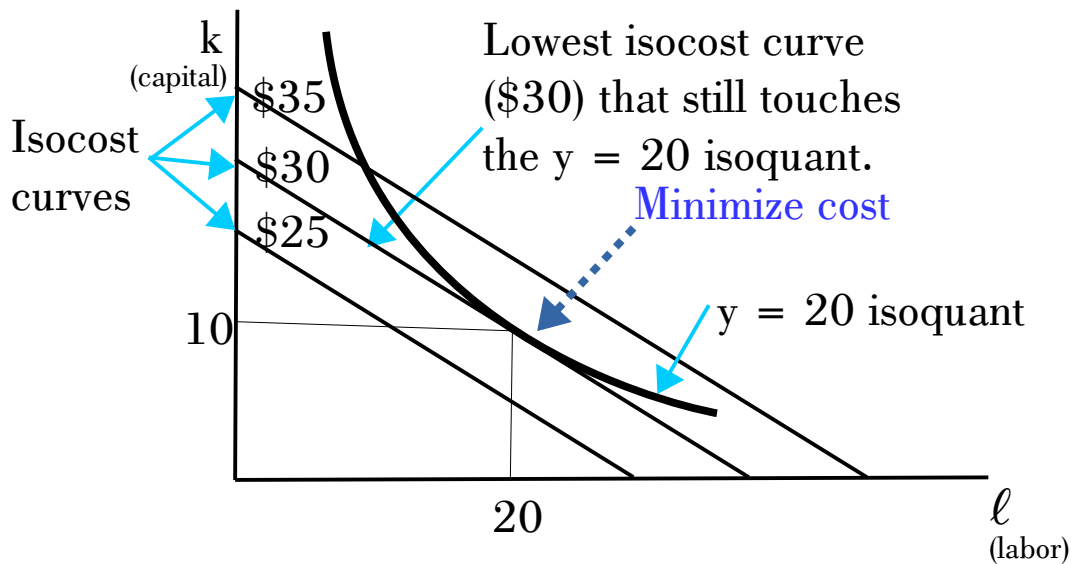


Figure 44: Isocost lines and cost minimization

Now suppose that the price ratio of inputs is represented by the solid isocost curves in the figure below. Given these input prices, the firm needs to know the least cost way of producing any given level of output. By finding a succession of tangencies to different isoquants, the firm can trace out the least cost ways of producing different levels of output. This is called an **expansion path**. If the price ratio of inputs changes (in the example below, the dashed lines show the case of lower relative labor prices), the firm finds a different expansion path that has both different minimal costs of production, and different cost minimizing levels of each input.

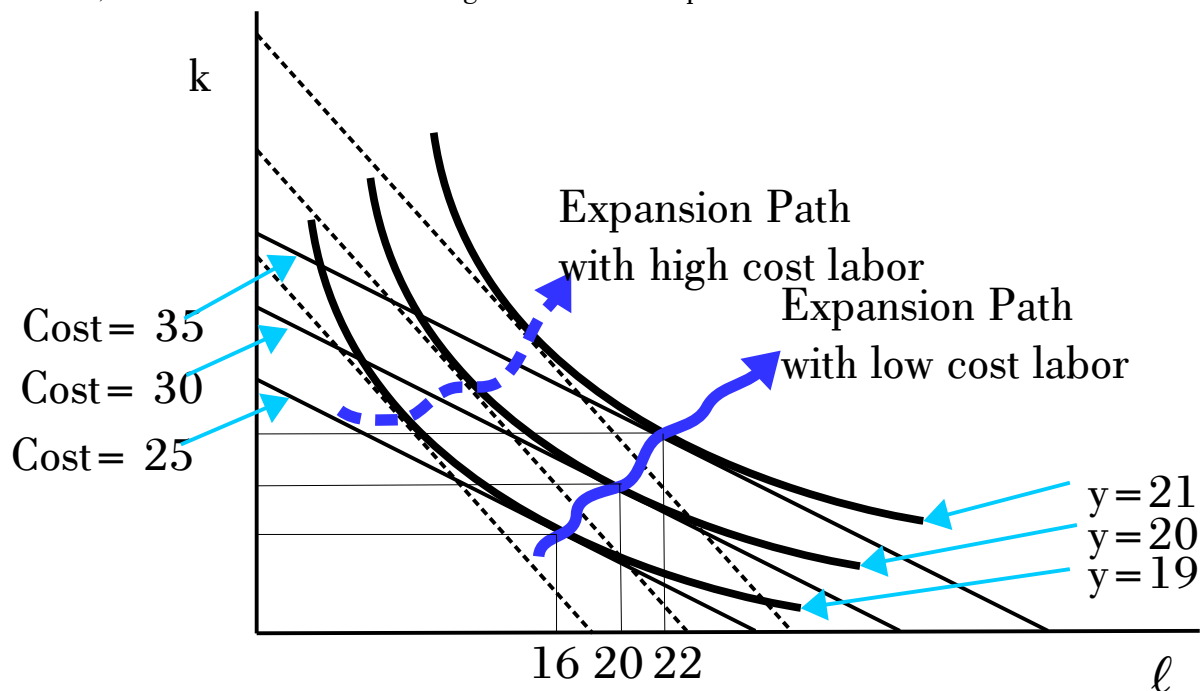


Figure 45: Cost minimization and the expansion path in production

Now consider only the solid isocost lines. From the graph, we can see that the least cost way to produce 19 units uses \$25 of inputs, and so on. This gives us the total cost function:

$$TC(19) = 25, \quad TC(20) = 30, \quad TC(21) = 35.$$

From this we can immediately derive the average cost and marginal cost functions:

$$AC(y) \equiv TC \frac{(y)}{y}, \quad MC \equiv \frac{\partial TC(y)}{\partial y}.$$

From a formal standpoint, the cost minimization part of the problem of the firm with many inputs is the following:

$$\min \sum_{n \in \mathcal{N}} p_n x_n \quad \text{subject to} \quad f(x) = \bar{y},$$

where x is the vector of inputs. To illustrate, suppose that we have two inputs and the production function has a Cobb-Douglas form. Then the firm's problem is this:

$$\min p_\ell \ell + p_k k = TC \quad \text{subject to} \quad f(\ell, k) = \ell^{1/2} k^{1/2} = \bar{y}.$$

Following our strategy for consumers, we can solve the production function for labor in terms of capital and substitute this into the cost function. This turns the constrained minimization into an unconstrained minimization:

$$\min p_\ell \ell + p_k \frac{\bar{y}^2}{\ell}.$$

Taking the derivative with respect to ℓ gives:

$$p_\ell - p_k \frac{\bar{y}^2}{\ell^2} = 0 \quad \Rightarrow \quad p_\ell = p_k \frac{\bar{y}^2}{\ell^2} \quad \Rightarrow \quad \ell^2 = p_k \frac{\bar{y}^2}{p_\ell},$$

and so:

$$\ell = \sqrt{\frac{p_k \bar{y}^2}{p_\ell}}, \quad \text{and by symmetry,} \quad k = \sqrt{\frac{p_\ell \bar{y}^2}{p_k}}.$$

These are called **factor demand curves**, and they give the cost minimizing level of each input that the firm should use in order to produce any given quantity of output when the prices of inputs are p^ℓ and p^k . These are the exact analogues of the Hicksian demand curves for consumers.

How do we get a total cost function? This is just the cost of the least cost way of making a given quantity! Thus:

$$TC(p_\ell, p_k, y) = p_\ell \sqrt{\frac{p_k y^2}{p_\ell}} + p_k \sqrt{\frac{p_\ell y^2}{p_k}} = \sqrt{p_\ell p_k y^2} + \sqrt{p_k p_\ell y^2} = 2\sqrt{p_\ell p_k y^2}$$

which would have been the expenditure function had we done the same thing in consumer space. More generally, solving the cost minimization problem while taking all inputs as variable, as we just did, really gives us what are called the **long run factor demand functions**:

$$FD_n^{LR}(p, y).$$

We will elaborate on this in the next section.

Section 4.4. The Long Run, the Short Run, and Cost Functions

It takes time for firms to respond to changes in demand or factor costs. Leasing new space, building and installing new machines and equipment, even hiring new workers, cannot be done instantaneously. Of even greater concern is that if costs go up or demand goes down, firms might be locked into leases or long term contracts with suppliers and workers. It might also take time to scrap or sell under utilized machinery. Thus, firms take time to reduce factor consumption to the optimal levels.

Because of these constraints, the short run choices of a firm differ from the long run choices. Formally, the short run and long run are defined as follows:

Short Run: A period of time over which at least some inputs are fixed and, therefore, not avoidable.

Long Run: The period of time over which all inputs are avoidable.

Notice that there are many potential short runs with more or fewer of the factors being fixed, but only one long run. We can now distinguish between two kinds of costs:

Fixed Cost: The cost of hiring fixed factors.

Variable Costs: The cost of hiring avoidable factors.

The key difference between variable and fixed factors are that variable factors are avoidable in the short run while the fixed factors are unavoidable. **Avoidability** means that a firm can change the levels of a factor that it uses, and can even reduce the level to zero. Variability in this context *does not imply* that an input can be varied in small or infinitesimal increments, simply that the input can be freely varied while respecting any physical aspects of the input and the effect it has on production.

For example, consider a hot dog vendor. To sell even one hot dog, he has to buy a whole hot dog cart. It does him no good to try to sell hot dogs out of half of a cart. The hot dog water would just spill out onto the sidewalk. This does not mean that hot dog carts are a fixed factor! It might be that there is an active, used market so that a vendor can buy one today for \$10,000 and sell it tomorrow for the same price. This means that hot dog carts are a variable factor since it is an avoidable cost in the short run period of one day. It is, however, a **lumpy variable cost**, that is, an *avoidable cost*, but one that must be *paid all at once* if output is positive. In cases like this, the marginal costs of even infinitesimally small positive output levels can be very large.

Of course, we could have considered the short run demand decision for consumers. Housing, for example, might be leased and so it may be difficult for agents to quickly increase or decrease consumption of this good. We do not seem to gain a great deal of additional insight regarding consumer behavior by considering this, however, and so not much has been done in this direction.

What happens to fixed costs in the long run? There are no fixed factors in the long run! All factors are variable, and no costs are fixed.

We will find it useful to talk about total, marginal, average, average variable costs below. In general, these are defined as follows:

Total Cost: $TC = (FC + VC)$

Average Cost: $AC = \frac{(FC + VC)}{y} = \frac{TC}{y}$

Average Variable Cost: $AVC = \frac{VC}{y}$

Marginal Cost: $MC = \frac{\partial TC}{\partial y}$

Note that the marginal cost is the incremental cost of increasing the output level. If we considered discrete changes, for example, increasing output from 12 to 13 units, the marginal cost would be the change in total cost: $TC(13) - TC(12)$. On the other hand, the effect on the TC of an infinitesimal change in output is the derivative of total cost. To put this another way, the marginal cost is the rate at which total cost changes as output goes up. We could also turn this around. If we integrated the MC function over the interval from $Q = 0$ to $Q = 13$, for example, we would get the VC of producing 13 units of output. That is, the VC is the “sum” of the MC s in a sense.

Now let’s focus on the relationship between average cost and both marginal and variable cost. Suppose we solved for the output level, y , where the average cost curve achieved its minimum:

$$\min_y \frac{1}{y} TC(y) \Rightarrow -\frac{1}{y^2} TC(y) + \frac{1}{y} MC(y) = 0 \Rightarrow \frac{1}{y} TC(y) = MC(y) \Rightarrow AC(y) = MC(y).$$

This implies that the at the output level where the firm is producing at the smallest possible average cost, it must be that $AC = MC$. This is equivalent to saying that the MC penetrates the AC curve at exactly the minimum of the AC . A similar argument shows that the MC penetrates the AVC curve at the minimum point of the AVC . To see this intuitively, suppose that the marginal cost is increasing the in quantity of output.

- If $MC < AC$, as output increases, the incremental cost of the last unit produced is less than the average cost of all produced units. Thus, when we average in this last, cheaper unit, the average cost of all units goes down. The AC is therefore downward sloping if $MC < AC$.
- If $MC > AC$, as output increases, the incremental cost of the last unit produced is greater than the average cost of all produced units. Thus, when we average in this last, more expensive unit, the average cost of all units goes up. The AC is therefore upward sloping if $MC > AC$.
- If $MC = AC$, as output increases, the incremental cost of the last unit produced is identical to the average cost of all produced units. Thus, when we average in this last, identically expensive unit, the average cost of all units stays the same. The AC is therefore flat if $MC = AC$. In other words, the AC has stopped sloping downwards and has not yet started sloping upwards when $MC = AC$, and so is at its minimal level.

The figure below illustrates this:

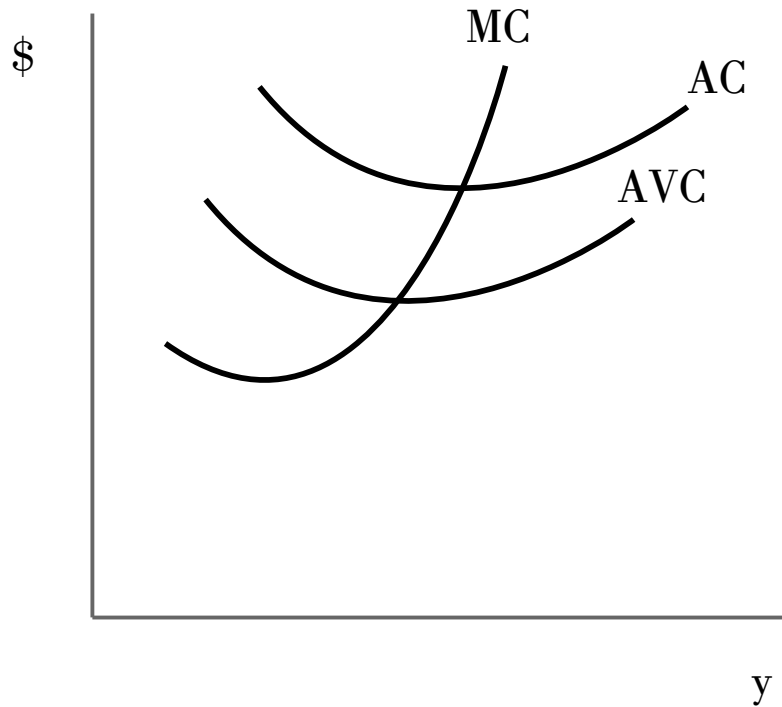


Figure 46: The relationship between MC, AC, and AVC

What about the short run? Recall that the cost minimization problem for the long run is:

$$\min \sum_{n \in \mathcal{N}} p_n x_n \quad \text{subject to} \quad f(x) = \bar{y}.$$

However, in the short run, some factors are fixed. They are no longer choice variables. This gives a short run cost minimization problem that is of reduced dimension. As the time span becomes shorter, more factors become fixed, and the cost minimization problem becomes one of smaller dimension.

In general, suppose that $\mathcal{N}^F$ is the set of fixed factors in some short run and $\mathcal{N}^V$ is the set of variable factors. By construction, this is a nonoverlapping partition of the inputs and so $\mathcal{N}^F \cap \mathcal{N}^V = \emptyset$ and $\mathcal{N}^F \cup \mathcal{N}^V = \mathcal{N}$. Formally, the **short run cost minimization problem** is the following:

$$\min \sum_{m \in \mathcal{N}^V} p_m x_m \quad \text{subject to} \quad f(x) = \bar{y}.$$

Solving this gives us **short run factor demand functions**. Although these factor demands include the price of fixed factors as arguments, they have no effect at all on the optimization.

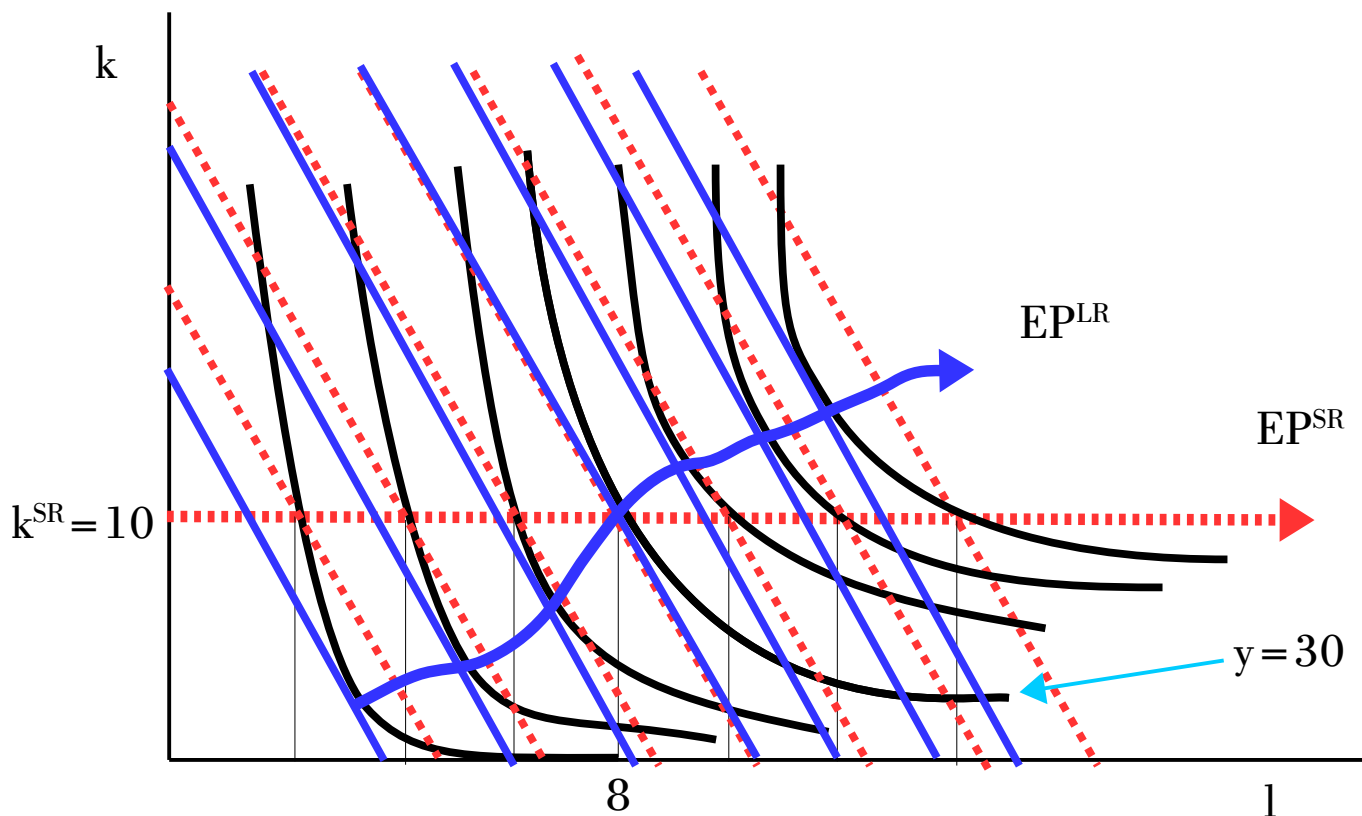


Figure 47: Cost minimization and the expansion path in the long and short run

The figure above illustrates the short run and long run cost minimization problem if there are only two inputs. In the short run, capital is a fixed factor. Thus, the cost minimization problem is degenerate. The short run expansion path is just the horizontal dashed red line at k^{SR} . For example, to produce 30 units of output takes k_{SR} units of capital and 8 units of labor in the short run. The cost is given by the red dashed isocost curve through this input bundle.

In the long run, capital is just another variable factor, so we choose the least cost input bundles where the blue solid isocost curves are tangent to each of the isoquants. This gives the blue solid long run expansion path.

Observe the following key point: It just so happens that given the relative factor prices shown by the slopes of the red and blue isocost curves, the long run, least cost, level of capital is the same as the short run fixed level of capital when output level is 30. You can see that the short run and long run minimum isocost curves happen to be the same here. However, when we reduce production, the long run cost minimizing input bundle has less capital. The long run minimal isocost curve is also below the short run isocost curve. When we increase production on the other hand, the long run cost minimizing input bundle has more capital. The long run minimal isocost curve is also below the short run isocost curve.

The long run minimal total costs are always less than short run minimal total cost except when the short run and long run cost minimizing level of fixed factors happen to agree, in which case short run and long run minimal total costs are equal.

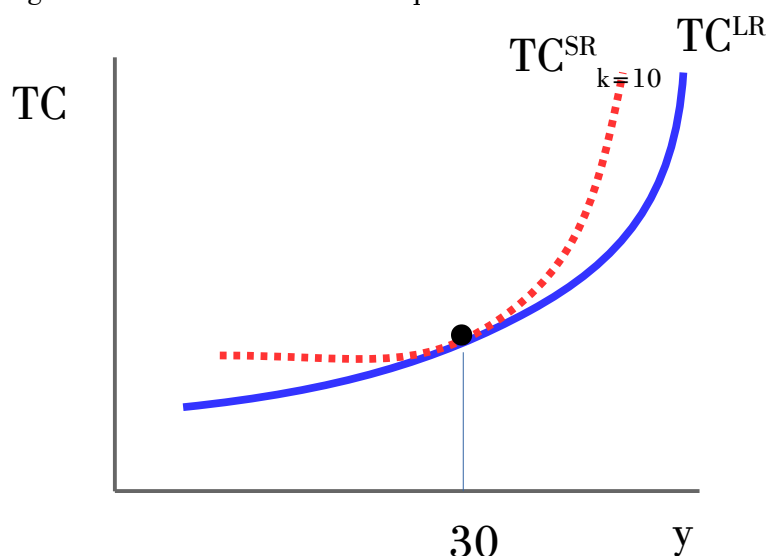


Figure 48: The relationship between the long and short run TC

This gives the following result. We can see that the $TC_{k=10}^{SR}$ is always (weakly) above TC^{LR} . Also note that the $TC_{SR}^{k=10}$ is only one of many possible TC^{SR} .

In the next figure, we draw in several other TC^{SR} for other fixed levels of capital. Notice that the TC^{LR} is an *envelope* that contains them all with each TC_{SR} tangent to the TC_{LR} at only one point.

You can see that the slope of the TC_{SR} is shallower when $y < 30$, the same when $y = 30$, and steeper when $y > 30$. The intuitive explanation is the following:

- When $y > 30$, costs increase faster in the short run since the firm has to produce output with a suboptimal level of capital. It has to replace this missing capital with labor, and this costs more on the margin.
- When $y < 30$, costs decrease more slowly in the short run since the firm has to produce output with a superoptimal level of capital. It would prefer to get rid of this extra capital, but it is a fixed factor in the short run. While having this extra capital allows it to produce with less labor than it would otherwise, costs would go down even further if it could stop paying for the inefficiently used capital and replace it with more cost-efficient labor.

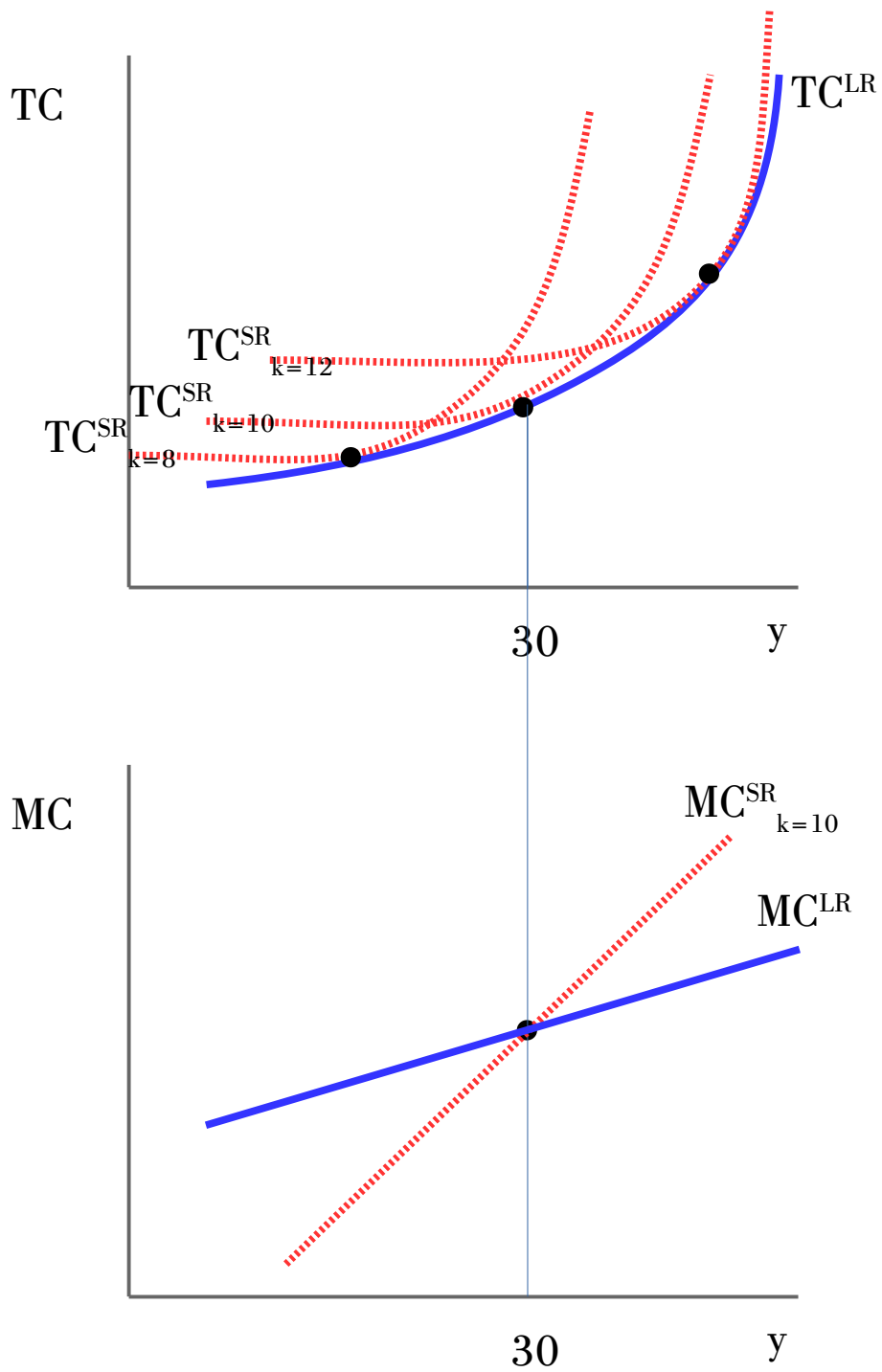


Figure 49: The relationship between the long and short run TC and MC

Putting this together gives us the following for the short run and long run:

Short Run Costs

Fixed Cost:	$\sum_{m \in \mathcal{N}^F} p_m \bar{x}_m$
Variable cost:	$\sum_{m \in \mathcal{N}^V} p_m FD_m^{SR}(p, y)$
Total Cost:	$\sum_{m \in \mathcal{N}^F} p_m \bar{x}_m + \sum_{m \in \mathcal{N}^V} p_m FD_m^{SR}(p, y)$
Marginal Cost:	$\sum_{m \in \mathcal{N}^V} p_m \frac{\partial FD_m^{SR}(p, y)}{\partial y}$
Average Cost:	$\frac{\sum_{m \in \mathcal{N}^F} p_m \bar{x}_m + \sum_{m \in \mathcal{N}^V} p_m FD_m^{SR}(p, y)}{y}$
Average Variable Cost:	$\frac{\sum_{m \in \mathcal{N}^V} p_m FD_m^{SR}(p, y)}{y}$

Long Run Costs

Fixed Cost:	0
Variable Cost:	$\sum_{n \in \mathcal{N}} p_n FD_n^{LR}(p, y)$
Total Cost:	$\sum_{n \in \mathcal{N}} p_n FD_n^{LR}(p, y)$
Marginal Cost:	$\sum_{n \in \mathcal{N}} p_n \frac{\partial FD_n^{LR}(p, y)}{\partial y}$
Average Cost:	$\frac{\sum_{n \in \mathcal{N}} p_n FD_n^{LR}(p, y)}{y}$
Average Variable Cost:	$\frac{\sum_{n \in \mathcal{N}} p_n FD_n^{LR}(p, y)}{y}$

Glossary

Average Cost Function: The total cost of production divided by the quantity produced. In general:

$$\frac{\sum_{n \in \mathcal{N}} p^n FD_{LR}^n(p, y)}{y}.$$

Average Variable Cost: The variable cost of production over some short run divided by the quantity produced. Variable costs are avoidable, but may or may not be continuously divisible (see lumpy variable costs). In general:

$$\frac{\sum_{n \in \mathcal{N}} p_n FD_n^{LR}(p, y)}{y}.$$

Avoidable Cost: The cost of avoidable or variable factors. Avoidability means that the firm can choose to reduce the quantity of a factor (an input to production) to zero in some short run. More factors become available as the time frame increases. Identical to variable cost.

Constant Returns to Scale (CRS): A property of the production function that requires that increasing inputs by some multiple $k > 1$ results an increase in output level of exactly k :

$$f(kx) = kf(x).$$

Decreasing Returns to Scale (DRS): A property of the production function that requires that increasing inputs by some multiple $k > 1$ results an increase in output level of exactly k :

$$f(kx) < kf(x).$$

Expansion Path: The locus in the factor space of the least cost bundles of inputs required to produce different levels of outputs under given factor prices.

Factor Demand Function/Curve: A function that gives the cost minimizing level of each input that the firm should use in order to produce any given quantity of output for any combination of factor prices. These are the exact analogues of the Hicksian demand curves for consumers. Factor demand curves are two-dimensional slices of the factor demand functions that graph the cost minimizing choice of a given factor required to produce a fixed level of output as a function of the price of the factor holding all other factor prices constant. Formally, the factor demand function is:

$$FD_n(p, y).$$

Fixed Cost: The cost of hiring fixed factors, that is, inputs that cannot be avoided or altered in some short run.

Increasing Returns to Scale (IRS): A property of production functions that requires that increasing inputs by some multiple $k > 1$ results an increase in output level of exactly k :

$$f(kx) > kf(x).$$

Isocost Line: A set of bundles of inputs that each cost the same amount under given factor prices.

Isoquant: A set of bundles of inputs that can be used to produce some specific level output given a firm's production function. The set of all isoquants contains the same information as the production function of a firm.

Long Run: The period of time over which all inputs are avoidable.

Lumpy Variable Cost: The cost of an input that is on the one hand avoidable, but on the other hand, not divisible. That is, if the firm leaves the industry, it is able to avoid using this variable factor. However, if a firm chooses to produce any positive level of output, it must purchase a discrete level of the factor instead of an amount of the input that is continuously proportional to the output level.

Marginal Cost Function: The incremental cost of increasing output level. Formally, this is the derivative of the total cost function with respect to output quantity:

$$\sum_{n \in \mathcal{N}} p_n \frac{\partial FD_n^{LR}(p, y)}{\partial y}.$$

Marginal Rate of Technical Substitution: Just as the slope of the indifference curve was called the MRS, the slope of the isoquant is called the Marginal Rate of Technical Substitution: $MRTS^{k,l}$. This is also equal to the ratio of the marginal product functions of two inputs.

Problem of the Firm: All firms minimize the cost of producing any level of output regardless of the type of market they sell goods in (competitive, monopoly, oligopoly, etc.). There are short run versions of this problem where fixed factor input levels are not choice variables and long run versions in which all input levels may be freely chosen. Formally:

$$\min \sum_{n \in \mathcal{N}} p_n x_n \quad \text{subject to} \quad f(x) = \bar{y}.$$

Production Function: A production function gives the level of output that results from using any given bundles of inputs. This way of describing a firm's technology is a special case of the more general "production set" approach since it assumes each firm produces one and only output, and that outputs are a separate set of goods that are distinguished from inputs. Formally:

$$f: \mathbb{R}^N \Rightarrow \mathbb{R}.$$

Short Run: A period of time over which at least some inputs are fixed and, therefore, not avoidable.

Total Cost Function: The cost of producing any given level of output under given factor prices:

$$\sum_{n \in \mathcal{N}} p_n FD_n^{LR}(p, y).$$

Variable Cost: The cost of hiring avoidable factors in some short run.

Problems

1. McDonald's is experimenting with androids to take over burger flipping jobs from humans. Suppose the production function for Big Macs produced by androids and humans, or a combination of the both is the following:

$$b = f(a, h) < f(a, h) = 10\sqrt{a+4h}.$$

- a. Graph the $h = 50$ isoquant. Using this, draw in a few more isoquants to show what the whole map looks like.
 - b. Is this production functions DRS, CRS, IRS, or none of these?
 - c. Suppose that androids cost \$1 per unit. What is the cost minimizing bundle of inputs to produce 50 Big Macs when humans cost \$1, \$2, \$4, \$6, and \$10, per unit. (Thus, give five cost minimizing bundles.)
 - d. Given what you learned above, write the total cost function for output when the input prices are fixed as $p_a = 1$ and $p_h = 2$. (That is $TC(b) = ?$)
2. As everyone knows, the perfect Martini consists of six parts gin and one part vermouth. No other combination is acceptable.
 - a. Using this information, write the production function for Martinis. Assume units of inputs are "shots" and it takes one shot of gin to make a Martini. (You should give an equation.)
 - b. Using the production function above, draw a picture of the isoquants where gin and vermouth are the inputs.
 - c. Suppose that gin costs \$5 per shot and vermouth costs \$7 per shot. Show the expansion path in the picture. Be sure the label at least three of the isoquant and isocost curves as well the quantities of gin and vermouth at the least cost tangencies.
 - d. Using this, write the total cost function for Martinis as function of quantity in liters assuming the input prices above are fixed.
 - e. Can you give the exact functional forms for the demands for any of these factors?
 3. Suppose that if a firm produces goods at all, its long run total cost, taking factor prices as fixed, is given by:

$$TC^{LR}(Q) = 300 + \frac{(Q^2+1)(Q+2)}{Q^2}$$

- a. What is the long run average cost, average variable cost and marginal cost?
- b. Suppose that there are only two inputs, capital and labor, and that in the TC^{LR} function given above, the price of capital is 100 and the price of labor is 10. Now suppose that the price of capital drops to 50 and the price of labor drops to 5. Either give the new AC^{LR} , VC^{LR} , and MC^{LR} or argue why it is impossible to determine without knowing the production function.

- c. Suppose instead that the price of capital drops to 50 while the price of labor goes up to 20. Either give the new AC^{LR} , VC^{LR} , and MC^{LR} or argue why it is impossible to determine without knowing the production function.
4. Tell whether the following statements are true, false, or uncertain? Be sure to explain your answers.
 - a. The long run total cost is *always* above the short run total cost.
 - b. The long run marginal cost is *never* above the short run marginal cost.
 5. Suppose that Oreo cookies oc are made using four inputs: cookie dough (c) white-stuff (w), labor (ℓ) and capital (k).
 - a. Write an example of a possible production function for Oreos (as in: $f(c, w, \ell, k) = ?$). For example, you may wish to write a Cobb-Douglas production function.
 - b. Suppose that the price of input x is p_x per unit. Assume in the short run, capital and labor are fixed while cookie dough and white-stuff are variable. Write the short run cost minimization problem for the firm.
 - c. Solving (b) gives four short run factor demand functions. What will these factor demands be functions of? For example, factor demand for white-stuff will depend on what variables?

Chapter 5. Supply and Demand in Competitive Markets

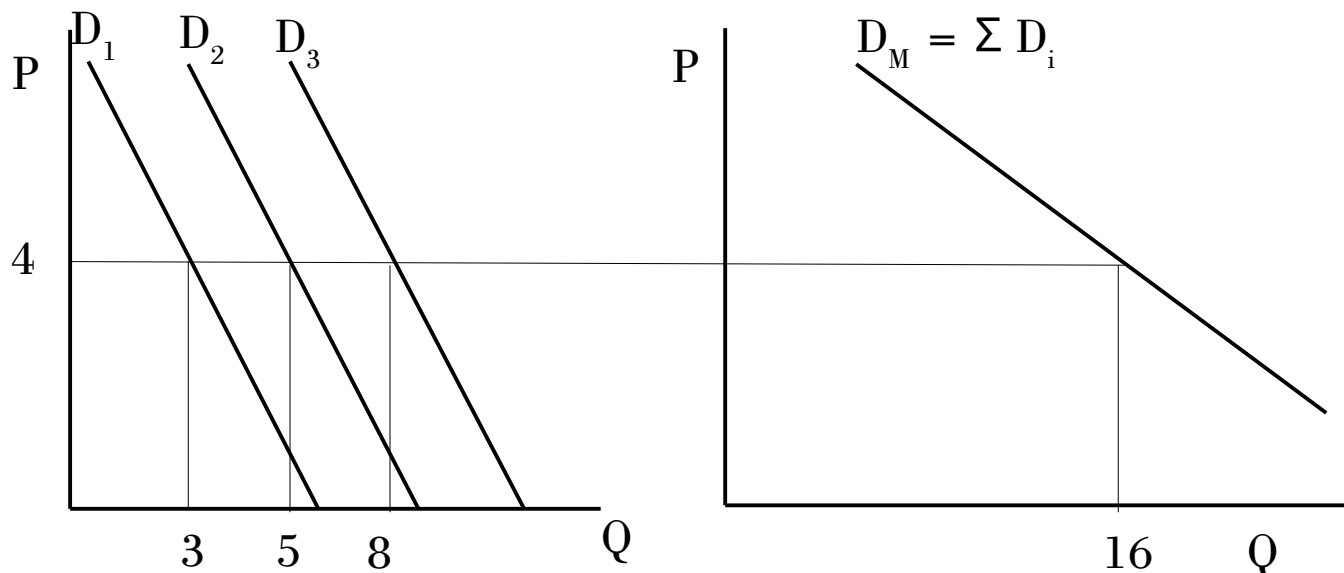
In this section, we consider the behavior of firms and consumers in a perfectly competitive market both as individuals, and in aggregate.

In general, markets tend to be competitive if there are many suppliers and many consumers. This implies that no agent has any power in the market and therefore simply accepts the conditions of the market and the actions of others as given and as something he cannot affect. Think of wheat farmers and bread consumers as examples on each side of a market. No wheat farmer thinks that by planting more wheat, he risks flooding the market and lowering the market price. Similarly, no consumer of bread worries that if he buys all the bread he wishes, the price might increase. Both agents contribute too little to the total supply or demand for the market as whole to notice their behavior.

A more precise way to say this is that, a market is **competitive** if all firms and all consumers are **price takers**. In contrast, if individual firms or consumers can affect the price in their markets, they have **market power**. They are **price makers** and the markets they participate in are not competitive. In the extreme case where only one firm supplies all the demand, we call the firm a **monopolist**. If there is only a single consumer buying all the supply, we call him a **monopsonist**. We consider noncompetitive markets in a subsequent chapter.

Section 5.1. Aggregate Demand of Competitive Consumers

In the previous section, we derived the demand behavior of an individual price taking consumer.



Individual demand curves

Market demand curve

Figure 50: Adding up individual demand curves to get the market demand

We can use this to directly derive the **aggregate** or **market demand**. The figure below illustrates this. Each consumer takes the price of the good as given, and buys the utility maximizing amount. For example, at a price of \$4, agent 1 demands a quantity of 3, agent 2 demands a quantity of 5, and agent 3 demands a quantity of 8. The market demand at a price of \$4 is therefore 16 units of the good. You can see that all we did was to horizontally add up the demand curves of each agent to find the market demand.

Mathematically, this is even more straightforward. The aggregate or market demand is simply:

$$D_n(p, w_1, \dots, w_I) = \sum_{i \in \mathcal{I}} D_{i,n}(p, w_i)$$

A Note on Notation

Recall that when there are two subscripts, the first refers of the agent and the second to the good. When there is only one subscript, it refers to the good or to the agent depending on context. Thus,

$$D_{i,n}(p_1, \dots, p_I, w_i) \equiv D_{i,n}(p, w_i) \Rightarrow \mathbb{R}$$

is the demand function of agent i for good n , while

$$D_i(p, w_i) \Rightarrow \mathbb{R}^N$$

is the demand function of agent i for all goods. That is, a multivalued function that maps prices the agent's endowment to an N – dimensional consumption vector. On the other hand,

$$D_n(p, w_1, \dots, w_I) \Rightarrow \mathbb{R}$$

is the aggregate, or market, demand for good n . This difference can be inferred from the fact that the aggregate demand depends on the income of all agents, not just that of a single agent i . Fortunately, we will not need to use these potentially confusing constructions to discuss either partial or general equilibrium in the coming chapters.

With this in mind, suppose we had two agents and two goods with the following Marshallian demand functions for good 1.

$$D_{1,1}(p_1, p_2, w_1) = \frac{w_1}{3} - 2p_1 + p_2,$$

$$D_{2,1}(p_1, p_2, w_2) = w_2 - (p_1)^{\frac{1}{2}} + \frac{p_2}{5}.$$

Then the aggregate demand would be given by the following formula:

$$D_1(p_1, p_2, w_1, w_2) = D_{1,1}(p_1, p_2, w_1) + D_{2,1}(p_1, p_2, w_2) = \frac{w_1}{3} + w_2 - 2p_1 - (p_1)^{\frac{1}{2}} + 6\frac{p_2}{5}.$$

Note several points here:

- To find the aggregate market demand, we add up all the individuals' Marshallian demands.
- We can see immediately that while the notions of own price and cross price elasticity are well-defined for market demand, income elasticity is not. The effect on market demand from an increase in consumer income, however, depends upon exactly *which* consumers are getting the extra income. (An exception to this is if all consumers happen to have homothetic preferences as where described in a previous section.)

- When doing **partial equilibrium analysis** (that is, considering each market in isolation and ignoring the effects of other markets), we simply take income and all other prices as fixed. This allows us to state the quantity demanded as a function of the price of good alone.

For the remainder of this section, we will focus on the aggregate demand for a single good. We will therefore drop the subscripts for agents, and goods as well as the superscripted A. Instead we will use Q for the quantity and P for the price of the single good we are considering. We will take all the other variables in the Marshallian demand function as fixed. The aggregate demand equation above therefore becomes:

$$Q(P) = \left(\frac{w_1}{3} + w_2 + 6\frac{P^2}{5}\right) - 2P - (P)^{\frac{1}{2}} = K - 2P - P^{\frac{1}{2}}$$

where

$$\frac{w_1}{3} + w_2 + 6\frac{P_2}{5} \Rightarrow \frac{\bar{w}_1}{3} + \bar{w}_2 + 6\frac{\bar{P}_2}{5} \equiv K$$

Section 5.2. The Supply Curve for Competitive Firms

Now consider the aggregate behavior of price taking firms in a competitive market. So far, we have solved the cost minimization problem for firms. This allowed us to derive short run and long run cost functions of various kinds. We stopped short of describing the individual supply curves. It turns out that this is a little more complicated than deriving the demand curves for consumers.

We begin in the same way. Since competitive firms are price takers, each firm believes it can supply as much or as little as it pleases at the market price. This means that from the perspective of each individual firm, the market demand curve is completely *flat, or horizontal, or elastic*. Of course, the market demand curve does slope downward in actual fact, however, over the range of output possible for any single firm, the price changes only trivially.

Taking a step back, notice that for *all firms*, competitive or not, the total revenue is equal to *price times quantity*, while profits, denoted π , are equal to *total revenue minus total cost*:

$$\pi(Q) = TR(Q) - TC(Q) = P(Q)Q - TC(Q),$$

where $P(Q)$ is called the “inverse demand curve” since it expresses price as a function of quantity instead of the reverse.

To maximize profits, the firm takes the derivative of the profit function and finds the quantity that causes the derivative to equal zero. Thus, profit maximization requires:

$$\frac{\partial \pi(Q)}{\partial Q} = MR(Q) - MC(Q) = 0 \Rightarrow MR(Q) = MC(Q).$$

The equation above is a general necessary condition for profit maximization. We will see below that it is not sufficient, however.

For the special case of competitive firms, recall that price is a fixed constant and is not affected the quantity the firm produces. In other words, the inverse demand curve takes this form: $P(Q) = P^{FM}$ where P^{FM} is the **free market price**, that is, the price that equates quantity supplied and quantity demanded in an ideal perfectly competitive market. This means that $TR(Q) = P^{FM}Q$ and so, $MR(Q) = P^{FM}$.

It follows that the necessary condition for profit maximization for competitive firms to choose a quantity such that:

$$MC(Q) = MR(Q) = P^{FM}.$$

In other words, in order to maximize profits, competitive firms should choose a quantity that equates the MC with the free market price. Another way to say this is that firms supply along their MC curves. *This seems to suggest that MC is exactly the same as the supply curve for competitive firms.* This turns out to be only partly true.

Section 5.3. Short Run and Long Run Shutdown

The presence of fixed costs does not seem to directly affect the argument that firms should equate MC and price. (Recall that by fixed costs, we mean unavoidable costs, not the costs associated with zero production.) Again, we will see that this is only partly true. Recall that $MC = P^{FM}$ is only a necessary condition for maximal profits. It may be, however, that these maximal profits are negative! In this case, the firm should shut down in the long run and leave the industry. Thus, that profits be positive is also a necessary condition for an output choice to be profit maximizing. Formally:

$$\pi(Q) > 0 \Leftrightarrow TR(Q) > TC(Q) \Leftrightarrow AR(Q) > AC(Q) \Leftrightarrow P^{FM} > AC(Q).$$

That is, the one and only way that profits can be positive is when the competitive price is higher than the average cost at the (profit maximizing) quantity produced. From this we can conclude:

THE PART OF THE MC CURVE THAT IS ABOVE THE MINIMUM OF THE AC CURVE IS THE LONG RUN SUPPLY CURVE FOR COMPETITIVE FIRMS.

This also is an example of a very important general principle that holds for firms, consumers and all other types of agents:

EQUIMARGINAL PRINCIPLE: IF A THING IS WORTH DOING AT ALL, IT IS WORTH DOING UNTIL THE MARGINAL BENEFIT EQUALS THE MARGINAL COST.

Suppose, for example, that the property taxes go up on the local McDonald's restaurant. This would lower profits. Will these increased costs be passed on to the consumer? Will the price of happy meals go up? No! This is an increase in fixed cost. It does not affect either marginal costs or marginal revenues. Thus, there is nothing that McDonald's can do to increase profits.

The point is that McDonald's absorbs these costs, but not because it wants to. We all know that clowns are evil, and we certainly would not expect Ronald McDonald to be an exception. The price that Ronald set before property taxes went up is already the most exploitative possible in the sense that it extracts the most profits possible. Any change in price would only decrease profits. To put this another way, if raising prices would increase profits after the tax increase, the price must have been set too low to begin with. It should have already been at this higher, more profitable level.

We should note that if we assume that the production function is concave, these two necessary conditions are also sufficient. That is, in the most typical case of firms operating under DRS tech-

nology, if there is output level Q such that $MC(Q) = P^{FM}$ and $\pi(Q) > 0$, then Q is the one and only profit maximizing choice.

With the positive profit condition in mind, let's turn to the short run problem for a competitive firm. In the short run, the firm can avoid *only the variable costs* if it shuts down. It cannot avoid the fixed costs by definition. Thus, if it shuts down, it must still pay the fixed costs and its profit is therefore $-FC$.

It follows that if the profits are negative but not as negative as the fixed cost, the firm is better off staying open. Even though the firm is making losses, the losses would be greater if it shut down right away. Staying in business allows the firm to partially offset the fixed cost and so it loses less money. In the long run when the fixed costs become avoidable, then the firm should shut down. Of course, if the losses in the short run are greater than FC , then the firm should shut down in the short as well. Formally, the firm should shut down in the short run if

$$0 > \pi \Leftrightarrow TR - TC < -FC \Leftrightarrow TR \leq VC \Leftrightarrow P^{FM} \leq AVC.$$

Putting this together in a picture:

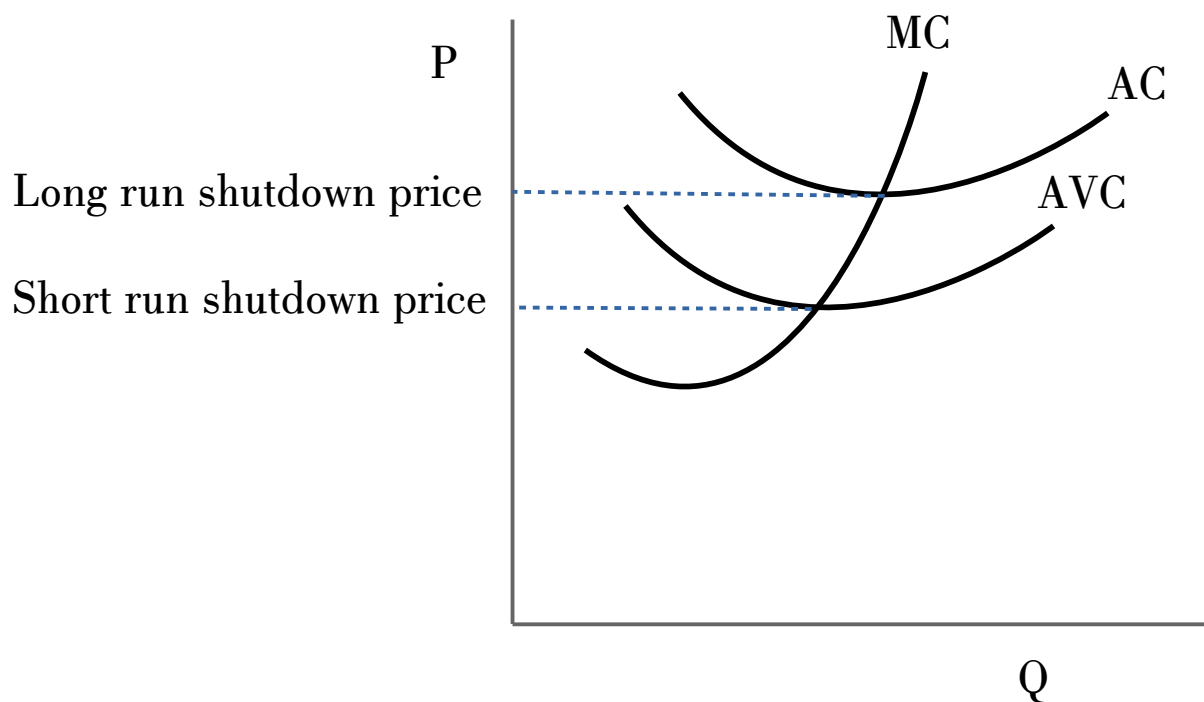


Figure 51: The zero profit prices in the long and short run

Fixed costs are avoidable in the long run. **Sunk costs**, on the other hand, are costs that can never be avoided. In a sense, sunk costs are fixed forever. It should be easy to see that sunk costs should never affect a firm's decision. If a cost is paid and cannot be recovered or reversed, then it has no effect on the optimality of any future decision.

For example, a firm might invest ten million dollars to develop a new product that it expects to monopolize and to sell for a high price only to find that a competing firm has developed a very similar product. If the first firm kept to its plan and set a high retail price, the competing firm could undercut it and take the entire market. On the other hand, if the first firm sets a lower price, it would never recover its research and development costs. What should the first firm do?

Research and development is a perfect example of a sunk cost. The firm can never unspend this money no matter what it does. If it sets a high price, it has spent ten million and then sells little or none of the resulting product. This only compounds the loss.

If it competes and sets a lower price, it has still spent the ten million, but it might make a small profit to partly offset this. If it could go back in time, it would choose not to spend ten million developing the product, but having spent it, the best thing to do is set the profit maximizing price based on current market and cost conditions. The cost of developing the product has nothing at all to do with determining this profit maximizing price. We summarize this idea as a general principle as follows:

SUNK COSTS ARE SUNK.

Section 5.4. The Short Run Aggregate Supply Curve

Recall that we found the market demand curve by horizontally adding together the individual demand curves. We can see from the argument above that each firm's supply curve is upward sloping, so if we followed a similar procedure, we would get an upward sloping market supply curve. In the short run, this is correct. We simply horizontally add up the part of the MC that is above the AVC of each firm that is currently in the market. The figure below illustrates this:

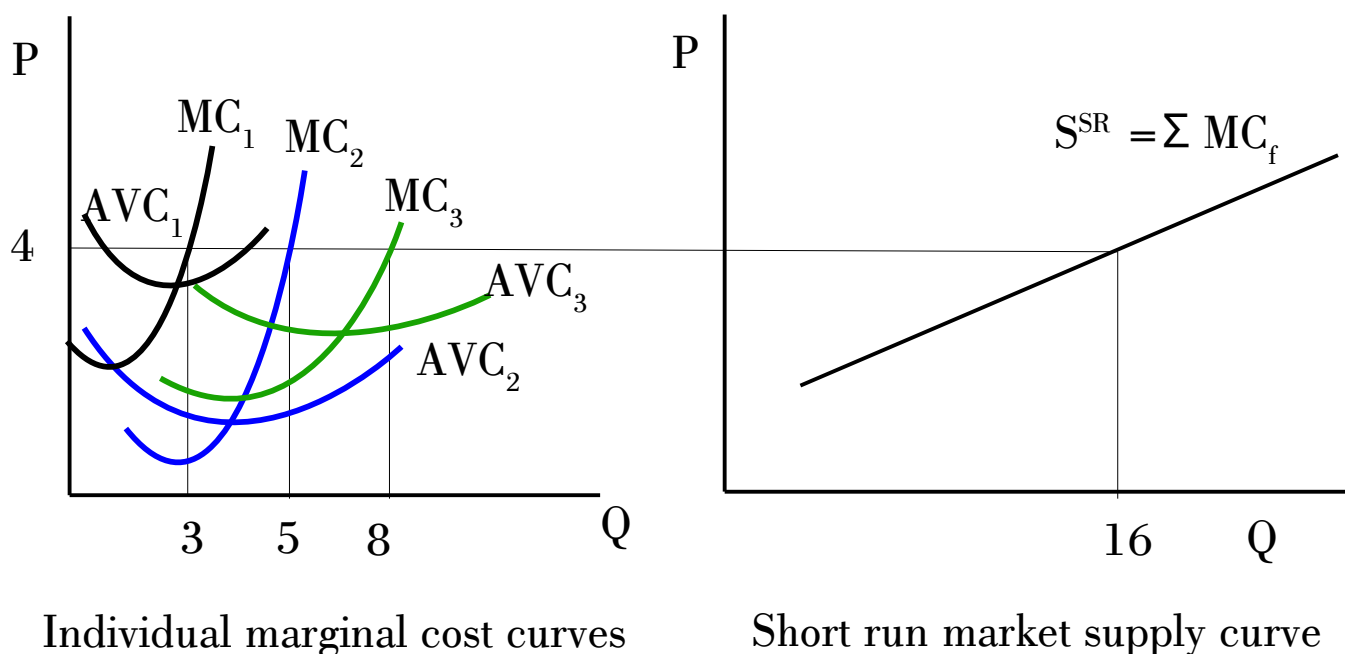


Figure 52: The short run aggregate supply

Note that in the short run, the number of firms in the market cannot increase. Firms may exit if they cannot cover variable costs, but by assumption, firms do not have time to enter.

Section 5.5. The Long Run Aggregate Supply Curve

Things are a little more complicated in the long run than the short run. In all of our discussion about consumer behavior, we have exogenously fixed the total number of consumers at some number I . This is probably reasonable. In the abstract, population might indeed be affected by market conditions, but it would take a very long time for fertility to respond to good economic times.

This is not the case for firms. In the long run, we would expect to see firms decide to enter profitable sectors and to exit unprofitable ones. Thus, in the long run, we need to account for how the entry and exit of firms affects the quantity supplied at any price. It is therefore unreasonable to exogenously fix the number of firms in a competitive market at some number, F .

To understand this, we need to be clearer about what we mean by “profits”. When most people think about profits, they have in mind the accountant’s definition, that is, the difference between total revenue and the total direct out-of-pocket costs of a firm. The problem with this is that it does not include indirect costs or other types of opportunity costs.

For example, if you start a restaurant in the first floor of your house, the space you give up is not a direct cost, but it does prevent you from using the first floor for any other purpose. Similarly, the time you spend is not a direct cost of the restaurant, but to work there, you must give up the opportunity to spend those hours working for someone else. An economist would subtract all these costs as well as the direct costs to get a more accurate measure of profits. Thus, there are two types of profits.

Accounting Profits: $\pi^A = (TR - \text{Direct Costs})$

Economic Profits: $\pi^E = (TR - \text{Opportunity Costs})$

When we say that there will be entry if profits are positive, we mean if *economic profits* are positive. That is, if something remains once all factors of production used by a firm have been compensated at the market rate (their opportunity cost), then the industry is more attractive than others. New firms would enter and drive down the price until economic profits return to zero. If economic profits in an industry are negative, then firms exit. This moves the market supply curve to the left since fewer firms exist to add their output to the total. Exit continues until the price rises enough to return industry profits to zero. To summarize:

**IN THE LONG RUN, ECONOMIC PROFITS IN ALL COMPETITIVE INDUSTRIES
MUST EQUAL ZERO.**

It may seem counter-intuitive that firms would participate in a market where profits are zero. Remember, however, that all factors are being compensated at their normal rate of return. For example, it might be that the ordinary rate of return to capital is 7%. If the hamburger industry gives a

return on capital of 7% after other factors are paid, then the hamburger industry is just as attractive a place to invest capital as any other.

If the hamburger industry gives a return of 10% however, it is a more attractive place to invest than others, and so investors will start new hamburger firms as long as the return stays above 7%.

If the return is only 5% then current firms in the industry will liquidate their holdings as quickly as possible in order to shift their capital into more profitable sectors. This results in lower quantities being supplied, which in turn results in an increase in the price of hamburgers. This continues until the return to capital goes back to 7%.

Now that we have established that the (economic) profits must be zero in the long run, we show that this implies something surprising about the long run price in an industry.

$$\pi^E = TR - TC = 0 \Leftrightarrow \frac{TR}{Q} = \frac{TC}{Q}.$$

Since $TR = P^{FM} Q$, we conclude $P^{FM} = AC$.

We now have two facts:

- Profit maximization requires for each firm that $P^{FM} = MC$.
- Long run competitive entry and exit of firms implies profits are zero and so $P^{FM} = AC$.

This implies that in the long run, firms in a competitive industry must operate where:

$$P^{FM} = MC = AC.$$

But there is only one place where $MC = AC$: *at the minimum of the AC function where it is penetrated by the MC!* Therefore, in the long run, this must determine the competitive price, P^{FM} , as well. Graphically:

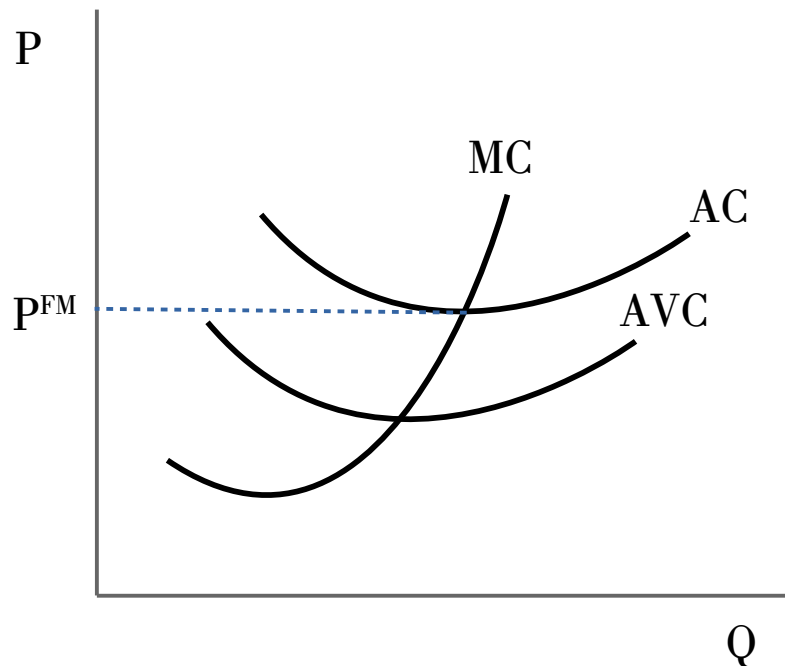


Figure 53: The zero profit prices in the long and short run

The next question is: what does the zero long run profit condition imply about the long run supply curve? There are possible cases:

Long Run Constant Cost Industry (LRCC): The long run supply curve is flat.

We find LRCC industries when all firms, and potential firms, have access to *identical technologies* and the *cost of inputs remain the same* even if there are industry-wide increases and decreases in the total output (which would increase the demand for factors of production in the industry).

In this simple case, the only zero profit price is the minimum of the average cost. Entry or exit takes place if and only if the price ever deviates from this. The long run supply curve is therefore flat, that is, completely elastic. For every level of industry output, we can imagine that just enough identical firms enter or exit to maintain exactly this price in the long run.

Long Run Increasing Cost Industry (LRIC): The long run supply curve is upward sloping.

What if firms are not all the same? Suppose that their technologies are different, perhaps because of differences in business plans or managerial skill. This would result in firms having different costs, and in particular, different average costs.

In this case, firms have different minimum average cost points, and so they become profitable at different prices. The firms with the lowest minimum average cost enter first, and as the prices go up, more firms enter as the price rises above their own minimum average cost point. Thus, higher

prices induce more firms to enter and add their own supply curve to the market supply curve. Similarly, lower prices cause the least efficient firms to exit as they can no longer cover their average costs, which in turn reduces industry supply curve.

In other words, the industry long run supply curve is upward sloping for two reasons: (1) each firm currently in the market has an upward sloping supply curve, so the sum of these is upward sloping and (2) higher prices cause new firms to enter and this increases the market output even further.

Now consider something new. While it is perfectly reasonable to imagine that input prices (also called “factor prices”) do not change if one firm changes its output, it is not as clear if all firms in an industry decide to increase output at the same time.

For example, if the nuclear power industry increases electricity production, we suspect that the price of uranium might go up. We call this phenomenon the factor price effect:

Factor Price Effect (FPE): The positive or negative effects seen on factor, or input, prices that result from changes in demand stemming from changes in output levels of industries that require them.

We only expect to see an FPE when an industry makes up a significant fraction of the total demand for a factor. For example, the nuclear power industry probably consumes almost all the Uranium produced each year (at least we devotedly hope so. Similarly, the airline industry consumes a large fraction of all aviation fuel produced. On the other hand, while the fast food industry requires a lot of beef to make hamburgers (at least we devotedly hope it is beef), this is nevertheless not a very high fraction of total annual beef production. We would expect to see FPE in the first two cases, by not the latter.

It is most common to see **positive FPEs**, where an increase in industry output causes an increase in factor prices. This is through the ordinary channel of upward sloping supply curves, and increasing marginal costs, in the industries producing the factors. However, we sometimes see **negative FPEs**, where an increase in industry output, causes a *decrease* in factor prices. This might happen in high technology sectors such as computer chips and aviation components. Negative FPEs are driven by the suppliers of these inputs being able to take advantage of economies of scale as they supply more of their factors to the industry that uses them.

Given this, Now consider a competitive industry in which all firms have identical technologies, but is subject to a positive FPE for some of its inputs. If the whole industry increases output, more inputs are required. The costs of inputs subject to positive FPEs increase as a result. In turn, this increases the total cost function for all firms (identically, in this case). Thus, the average cost, and so the minimum of the average cost curve, for all firms go up. This new minimum is then the zero profit price at the new, higher, level of aggregate output.

Since the zero profit output price goes up when output increases, and down, when output decreases, the long run aggregate supply curve is upward sloping. In the case of a positive factor price

the argument is reversed. The zero profit output price goes down when output increases, and so the long run aggregate supply curve is downward sloping.

Note that the change in factor prices in these cases is *endogenous* in the sense that it is systematically driven by the factor demand of a specific industry using the input. The change in factor prices is not a result of *exogenous* forces such as a change in the cost of inputs needed to produce the factor, or increases in demand for the factor by other industries.

We conclude that long run decreasing cost industries can result for firms having different technologies, and therefore costs, or positive factor price effects, or both at the same time.

Long Run Decreasing Cost Industry (LRDC): The industry long run supply curve is downward sloping.

If all firms, and potential firms, have access to identical technologies, but the industry is subject to negative FPEs, as the industry output goes up, factor prices decrease. This causes the minimum of the average cost point for each firm to go down. Thus, the zero profit price for the industry goes down when output goes up, and goes up when output goes down. In other words, the industry long run supply curve is downward sloping. If firms had access to differing technologies, it would partly, or wholly offset the negative FPE.

Putting this together we get the following conclusions:

Long vs. short run supply curve:

- The short run industry supply curve is the sum of supply curves of the individual firms currently in the market.
- In the long run, both the total number of firms and the shape of each of their supply curves may change if demand or factor prices change.
- In a long run equilibrium, the following are all true:
 - The industry price and quantity is determined by the intersection of the industry demand curve and the long run industry supply curve.
 - The short run industry supply curve will intersect the industry demand curve at the same price and quantity as above, even though the slope of the short run industry supply curve will generally be steeper than the long run industry supply curve.
 - The long run equilibrium price determined above will give all firms currently in the industry (or at least the marginal firm) zero economic profits, and so is also equal to the minimum of the firms' AC curve.

If market demand changes, but there is no exogenous change in factor prices:

- In the short run, the equilibrium price and quantity is determined by the intersection of the short run supply curve and the new demand curve.
- In the long run, the equilibrium price and quantity are determined by the intersection of the long run supply curve and the new demand curve. Note that the long run industry supply curve does not change when demand changes.

If some factor price changes exogenously, but there is no change in market demand:

- In the short run this might affect either variable or fixed costs of firms. If variable costs change, then marginal costs of all firms change, which means that their sum, which is the short run industry supply curve, changes as well. The short run equilibrium price and quantity is determined by the intersection of the new short run supply curve and the original demand curve. On the other hand, if the cost of a fixed factor changes, only the average cost curves of firms change. Their marginal cost curves are unaffected. As a result, the industry short run supply curve also stays fixed, and so neither the aggregate price or quantity change.
- In the long run, factor price changes affect total, average, and marginal cost of all firms. Recall that the long run (zero profit) industry supply curve is based on some initial set of factor prices, taking into account predictably, endogenous FPEs. If factor prices change exogenously, however, then the zero profit prices also changes. If factor costs exogenously increase, the long run industry supply curve shifts upward, and the opposite if factor prices decrease. The shape of the long run supply might also change, however, it would retain its direction of slope. LRDC industries would remain LRDC industries, for example, although with higher zero profit prices for any aggregate level of output. The long run industry price and quantity then determined by the intersection of this *new long run supply curve* and the initial demand curve.

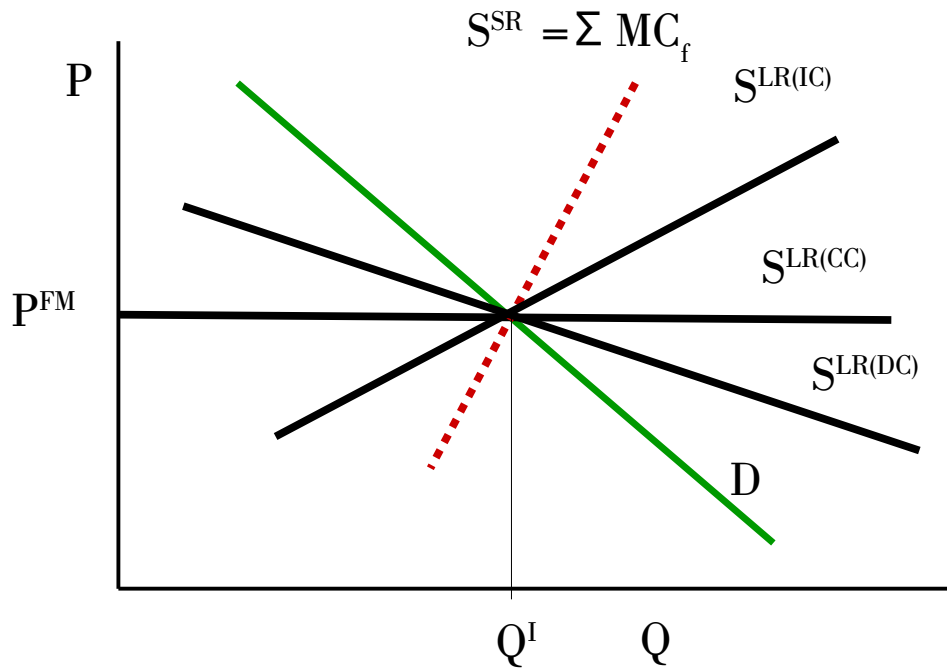


Figure 54: The long run demand for IC, DC and CC industries

The figure above shows the industry demand curve (the green downward sloping line) and the short run supply curve (the red dashed line). It also shows all three possible types of long run supply curves as heavy lines. Of course, only one of these can be correct at a time. The figure shows what the industry looks like when it is in equilibrium. This requires that the demand, short run supply curve and long run supply intersect at the same price/quantity point. At a long run equilibrium, P^{FM} is the free market price, and Q^I is the quantity supplied collectively by all the firms currently in the industry.

Zero profits?

Astute students might have noticed something fishy in the discussion of LRIC industries. You (I assume you are an astute student) might object if firms have different technologies and enter the market as they cover their minimum average costs, then more efficient firms must be earning more than average cost. That is price just covers average cost when a firm enters, but then price must go above his minimum average cost in order to induce the next, somewhat less efficient firm to enter.

Looking at figure below, the free market price is well above the average cost curve of firm 3 (green). It seems that profit should equal $\pi_3 = (P^{LR} - AC_3)Q_3$, the green hatched area. There are two answers: an easy wrong one, and a somewhat more complex correct one.

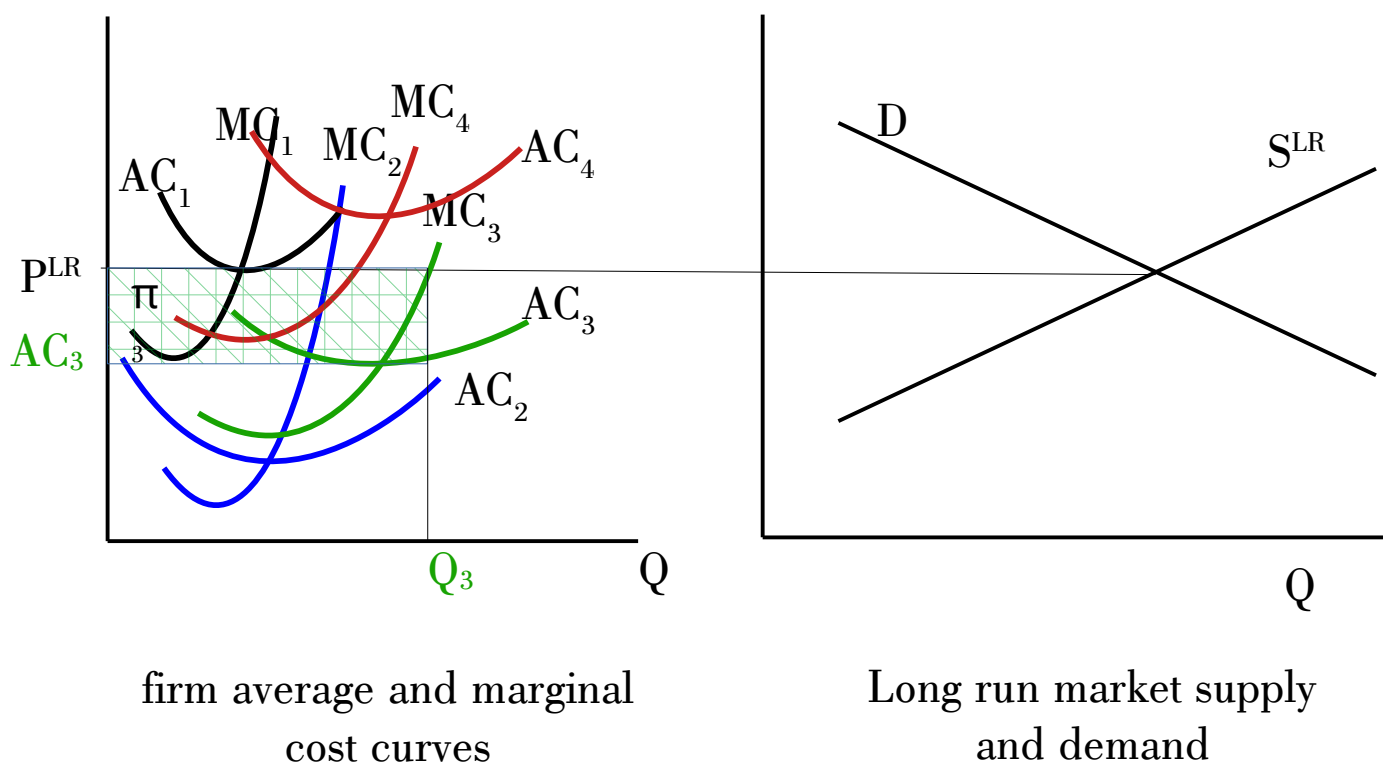


Figure 55: The long run in a LRIC industry

Wrong but workable answer: All that matters for the entry exit dynamic that allows us to conclude that the long run supply curve will be upward sloping is that the last firm to enter makes non-negative profits, and the most efficient potential cannot cover its minimum average cost at the market price. If so, we are in equilibrium, and all is well.

Correct but more complicated answer: All firms do, in fact, make zero economic profits in the long run, even firms that entered early. Suppose, for example, these firms are 7-11s and Adam, Bob, and Caroline are the managers of store 1, 2 and 3, respectively. All are paid \$30,000, the go-

ing rate for managers. Store 1 makes no economics profit after Charlie is paid, but its owner makes the ordinary rate of return on his invested capital. The store Bob manages makes a profit of \$75,000, while Caroline's makes \$100,000, even after both are paid.

What if Caroline offered her services to other store owners, and potential investors in 7-11s? If the owner of store 1 fired poor Adam, and hired Caroline at a wage of \$130,000, he would just as well off as he is now, and still make the ordinary rate of return on his capital. Any other investor could start a 7-11 and do the same. As a result, Caroline must make \$130,000 in equilibrium. If she is paid any less, there is an arbitrage opportunity to hire her, star a new 7-11, and make economic profits equal to \$130,000 minus whatever she is paid. Similarly, Bob must make \$105,000 in equilibrium (all this assumes full information, and efficient markets, of course).

Assume store 3's owner understands all this, and pays Caroline what she is worth. The effect is to raise store 3's average cost curve so that its minimum equals the market price. In other words, the value of factors (managerial skill of individuals) is determined endogenously by the market. In an efficient market the cost of these factors adjusts through compensation so that there is economic profit or advantage to hiring one over the other. The wage equals productivity, and so economic profits are zero regardless who is hired.

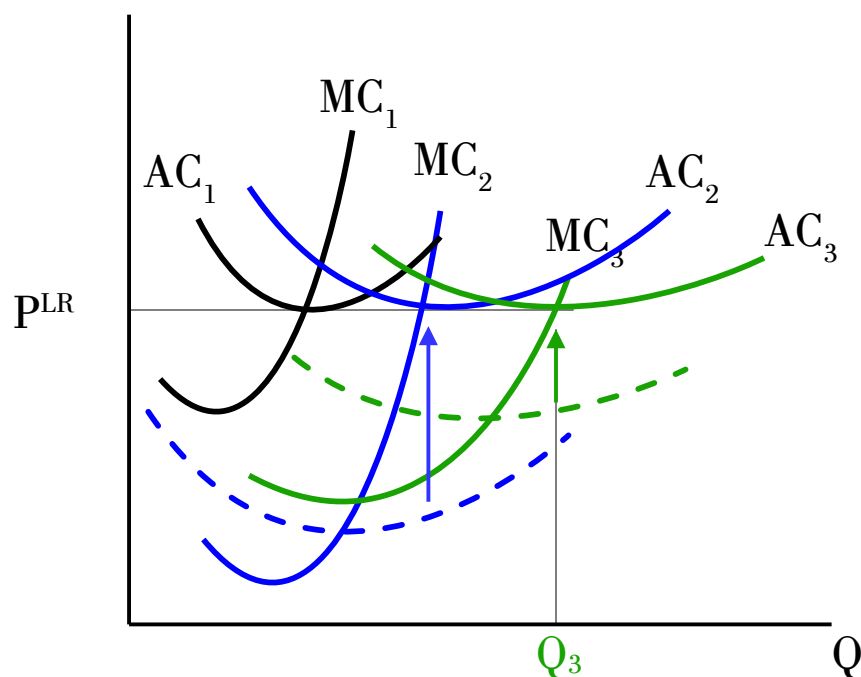


Figure 56: Endogenous valuation of factors leads to zero economic profit

Real world situations like this are everywhere, but consider sports as an example. There is nothing inherently more valuable about being good at basketball relative to being good at water polo. Both are difficult to sports to master, and require a high degree of athleticism and physical strength to excel. Nevertheless, basketball players make vastly more than polo players. Why are polo players treated so poorly?

You know the answer. Many millions of people love basketball and are willing to pay high prices to watch it. Very few want to watch water polo. Even if basketball players had no other skills, the opportunity cost of working for the Lakers is not working for the Bulls. Thus, basketball players make big salaries since market conditions make the next best opportunity for them very valuable.

Sadly, the next most lucrative opportunity for a polo player might be selling insurance, or running a business. Even the best player can only expect to be paid his marginal contribution to ticket sales, and other revenue.

A perhaps unpleasant lesson here is that markets determine the value of both our endowments and our efforts. They have no innate value, at least to world at large. Since we need other people to survive and enjoy life, we must, at least to some extent, figure out what the world wants and provide it. Of course, we can choose to ignore this, but the consequence is that we will find few opportunities for beneficial exchange with rest of world.

Section 5.6. Example: Short and Long Run in a Long Run Constant Cost Industry

To see the dynamic described above at work, consider what happens in a LRCC industry when demand, fixed cost, or variable costs change.

To begin with, the figure below shows such an industry in an equilibrium state. On the left side we show the cost curves for a representative firm in the industry. Notice that the price is equal to the minimum of the average cost curve. If it were anything else, we would have entry or exit of firms until this price prevailed.

On the right side, we show the demand and supply curves for the industry as a whole. The short run supply curve is equal to sum of the marginal cost curves of the firms currently in the industry. It is upward sloping as a result. The long run supply curve, however, is flat. This is because in the long run, firms cannot make economic profits or losses. As a result, enough identical firms will enter or exit this industry to supply any amount consumers might demand at the long run zero profit price P^{LR} . Finally, note that the long run supply, short run supply and demand curves all intersect at the same price and quantity.

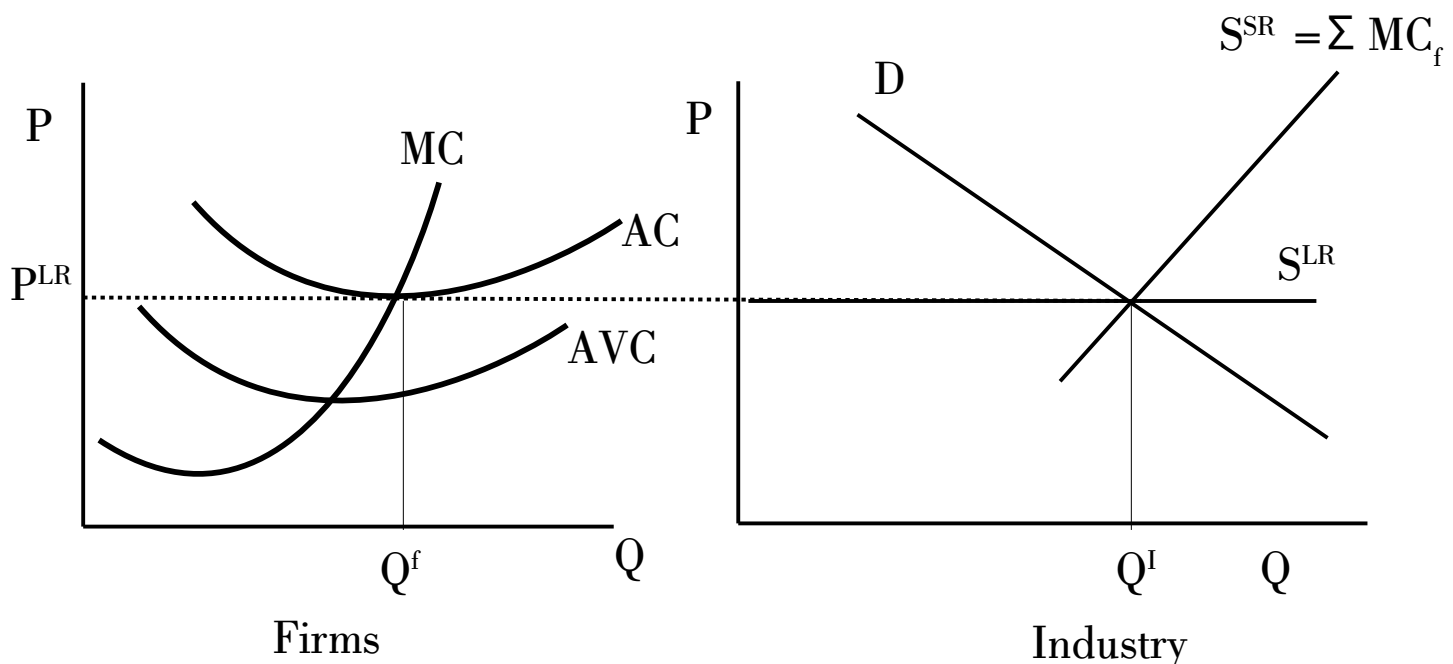


Figure 57: LRCC industry in equilibrium

Note that Q^f is the quantity supplied by a representative firm in the industry. Since by assumption all firms have identical technologies and costs in a LRCC industry, each firm supplies exactly Q^f .

Subsection 5.6.1. Change in Demand

Suppose that there is an exogenous *increase in demand*. Since there are no exogenous changes in any of the factor costs, the MC, AC, and AVC, of the firms in the industry are all unchanged. This implies that the short run supply curve, which is the sum of the MC curves, is also unchanged. Therefore, the short run equilibrium price and quantity for the industry are determined at the intersection of the new demand curve and the short run supply curve.

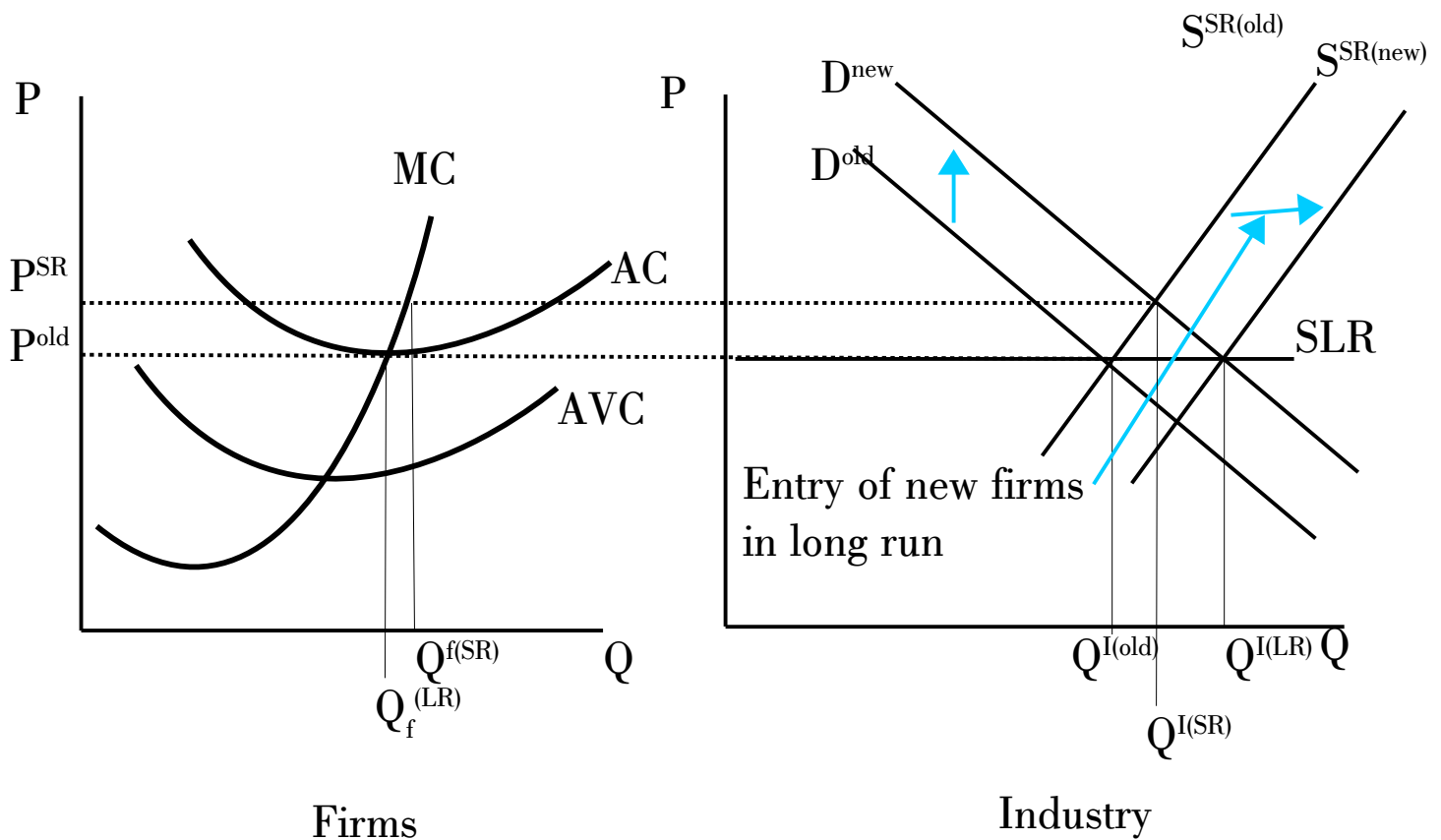


Figure 58: LRCC industry going to a new equilibrium when demand changes

Since the demand has gone up, we move up the unchanged old short run supply curve and the new short run equilibrium is therefore the following:

Short run

- Industry price goes up: $P^{old} \rightarrow P^{SR}$
- Industry quantity goes up: $Q^{I(old)} \rightarrow Q^{I(SR)}$
- Firm quantity goes up: $Q^{f(LR)} \rightarrow Q^{f(SR)}$

You can see that the new short run price is above the firm's AC. Profits are therefore positive in the short run. This generates entry by new firms. As firms enter, their MC curves are added the current industry short run supply curve. This causes the industry supply curve to start shifting to the right. As a result, the intersection of the short run supply and demand shifts down and right toward higher quantities and lower prices. Entry stops when the price finally reaches the old zero profit price. Thus, the price returns to where it started, firms produce what they used to, but total industry output increases due to the presence of more firms in this LRCC industry.

Long run

- Industry price goes down: $P^{SR} \rightarrow P^{old}$
- Industry quantity goes up: $Q^{I(SR)} \rightarrow Q^{I(LR)}$
- Firm quantity goes down: $Q^{f(SR)} \rightarrow Q^{f(LR)}$

Subsection 5.6.2. Change in the Cost of a Fixed Factor

Now suppose there is an exogenous *increase in the cost of a fixed factor*. In the short run, this increases TC and therefore AC. However, the increase in FC has no effect on VC or AVC, and in particular, no effect on MC. Thus, firms' AC shift upward, but MC is unchanged. This means that there is no change in the short run supply curve for the industry. Since demand did not change either, there is no change in the price or quantity in the industry. In turn, since each firm's MC is unchanged, there is no change in quantity for the firm either.

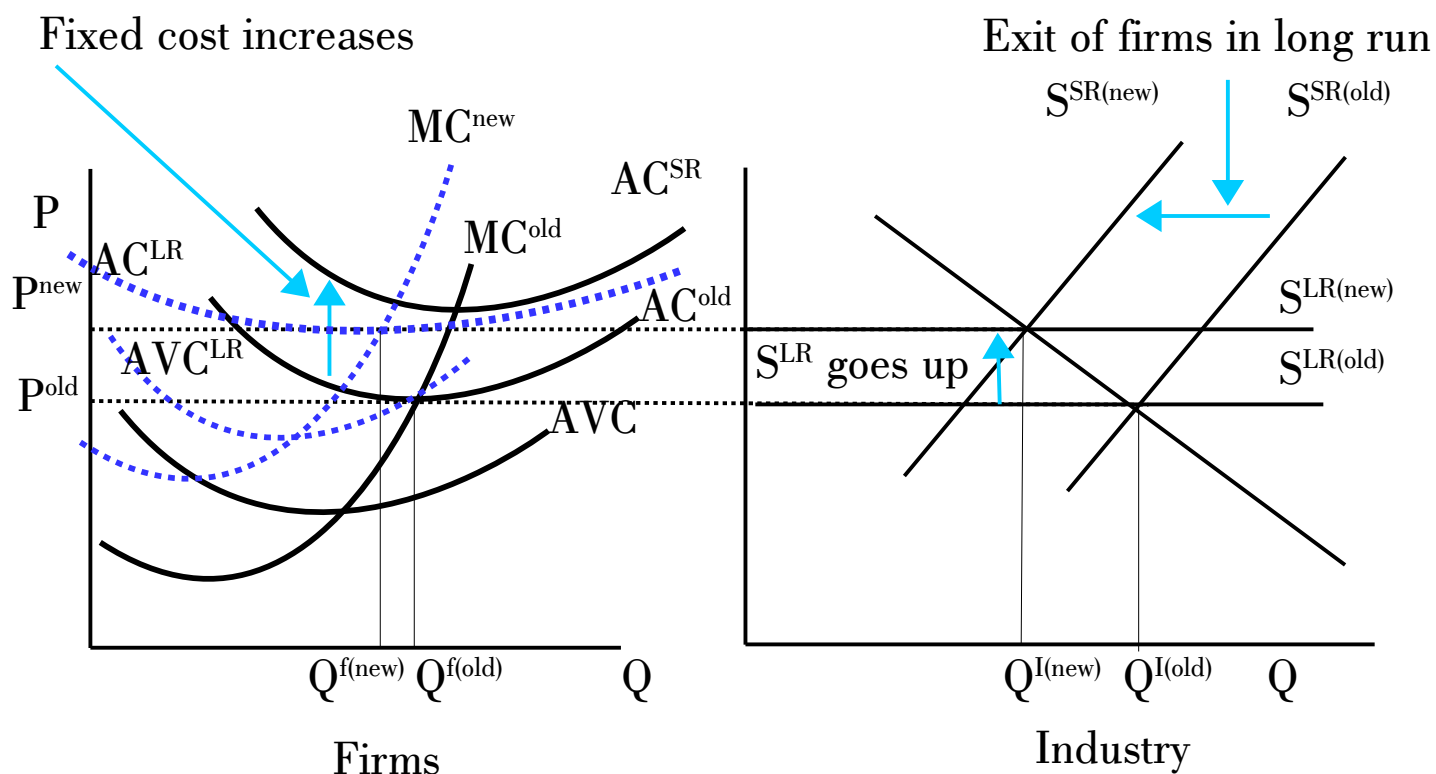


Figure 59: LRCC industry going to a new equilibrium when FC change

Short run

- Industry price is unchanged: $P^{old} \rightarrow P^{old}$
- Industry quantity is unchanged: $Q^{I(old)} \rightarrow Q^{I(old)}$
- Firm quantity is unchanged: $Q^{f(old)} \rightarrow Q^{f(old)}$

The short run price is below firms' AC. As a result, profits are negative in the short run. This generates an exit of firms in the long run. As firms exit, fewer MC curves are added to the short run

supply. This causes short run supply curve to start shifting to the left. As a result, the intersection of the short run supply and demand shifts up and left toward lower quantities and higher prices.

Exit stops when the price finally reaches the new zero profit price. *Note that since factor prices increased exogenously, this is a new, higher, long run industry supply curve.* This supply curve is still flat, but reflects the new, higher, zero profit price. In addition, since the relative price of factors has changed, the slope of the isocost curves has changed. This means that the (long run) minimal cost input bundles are different. Without knowing the exact production function of the firms, we cannot tell exactly how the MC, AC, and AVC, move.

We can tell that the minimum of the AC curve will be higher than it was before the factor price increase, but will be lower than it was in the short run before the firms could adjust to optimal usage of the fixed factors. (Note that the long run adjustment will most likely involve shifting out of fixed and into variable factors since a higher fixed factor price implies lower relative variable factor prices.) The dashed blue cost curves show one possible example of what might happen that is consistent with this. Thus:

Long run

- Industry price goes up: $P^{old} \rightarrow P^{new}$
- Industry quantity goes down: $Q^{I(old)} \rightarrow Q^{I(new)}$
- Firm quantity is uncertain: $Q^{f(old)} \rightarrow Q^{f(new)}$

The example shows firm quantity going down, but it is also possible that the firm quantity could go up.

Subsection 5.6.3. Change in the Cost of a Variable Factor

Now suppose there is an exogenous *increase in the cost of a variable factor*. In the short run, this increases all the cost curves. Thus, the MC, AC, and AVC, all shift upward and may change shape. This is shown with the lighter red dashed curves in the figure below. We see that the short run supply curve shifts upwards since MC of all firms shifts upwards.

Thus, in the short run, industry quantity and price are determined by the intersection of the demand and the (middle) short run supply curve. The price goes up, but at the same time MC goes up for each firm. Thus, each firm produces an output that is on the upward shifted MC, but at a higher price. It might look like these effects are offsetting, so we can not tell if firms produce more or less. However, in this example, the number of firms is unchanged in the short run. Since all firms are identical and the industry output decreases, each firm must be producing less.

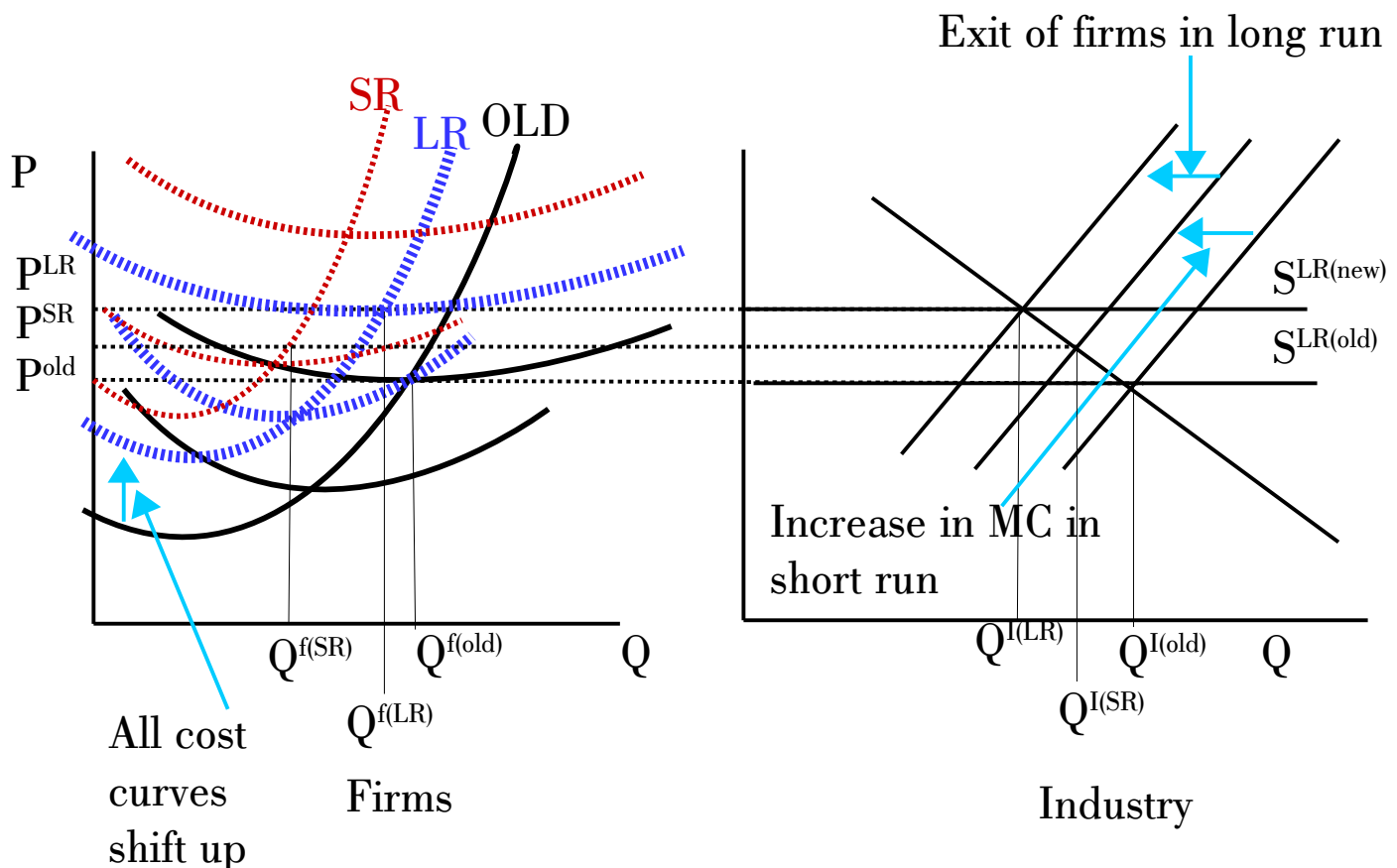


Figure 60: LRCC industry going to a new equilibrium when VC changes

Short run

- Industry price goes up: $P^{old} \rightarrow P^{SR}$
- Industry quantity goes down: $Q^{I(old)} \rightarrow Q^{I(SR)}$
- Firm quantity goes down: $Q^{f(old)} \rightarrow Q^{f(SR)}$

The short run price is below the firm's (**thinner red dashed**) AC. Profits are therefore negative in the short run. The long run effects are the same as they were when fixed costs increased. For completeness, we repeat the argument:

Negative profits generate exit of firms in the long run. As firms exit, fewer MC curves are added to the short run supply curve. This causes it to start shifting to the left. As a result, the intersection of the short run supply and demand shifts up and left toward lower quantities and higher prices. Exit stops when the price finally reaches the new zero profit price. *Note that since factor prices increased exogenously, this is a new, higher, long run industry supply curve.* This supply curve is still flat, but reflects the new, higher, zero profit price.

Again, without knowing the exact production function of the firms, we cannot tell exactly how the MC, AC, and AVC, move, but the minimum of the AC curve will be higher than it was before the factor price increase and lower than it was in the short run before the firms could adjust to optimal usage of the fixed factors. (Note that the long run adjustment will involve shifting into, rather than out of, fixed factors since a higher variable factor price implies lower relative fixed factor prices.) The **thick dashed blue** cost curves show one possible example of what might happen that is consistent with this. Thus:

Long run

- Industry price goes up: $P^{SR} \rightarrow P^{LR}$
- Industry quantity goes down: $Q^{I(SR)} \rightarrow Q^{I(LR)}$
- Firm quantity is uncertain: $Q^{f(SR)} \rightarrow Q^{f(LR)}$

The analysis of what happens for LRDC and LRIC cost industries is largely the same. The complication is that in addition to having the cost curves of firms move in the long run due to exogenous changes in relative factor prices, we also have to take into account that factor prices change endogenously due to positive or negative factor price effects.

Glossary

Accounting Profit: The difference between revenues and the direct out-of-pocket costs of production: $\text{Total Revenue} - \text{Direct Costs}$.

Aggregate or Market Demand: The total demand for a given good summed up over all agents in a specific market. This is equal to the horizontal sum of the demand curves of the relevant set of consumers.

Aggregate or Market Supply: The total supply for a given good summed up over all firms in a specific market. If the number of firms is fixed (as it is for example in the short run) and the market is competitive, then this equals the horizontal sum the part of each firm's marginal cost curve that is above the average variable cost. If entry or exit of firms is possible, and the market is competitive, the long run supply curve is the locus of zero profit price/quantity combinations in the market. If the market is not competitive, there is no such thing as a market supply curve.

Competitive Market: A market is competitive if all firms and all consumers are price takers.

Economic Profits Any accounting profit in excess of what is required to pay indirect costs and ordinary market rates of compensation other factors used in production that are not paid for out-of-pocket: $\text{Total Revenue} - \text{Opportunity Costs}$.

Equimarginal Principle: If a thing is worth doing at all, it is worth doing until the marginal benefit equals the marginal cost.

Factor Price Effect: A FPE is seen whenever industry-wide changes in output affect the price of any of its inputs. FPEs are generally positive (that is, increases in industry output drive up the price inputs) but can be negative in unusual cases.

Long Run Constant Cost Industry: Industries with a flat long run supply curve. For the zero-profit price to be the same for all levels of output, firms must have access to identical technology and there can be no factor price effects.

Long Run Decreasing Cost Industry: Industries with a downward sloping run supply curve. For the zero-profit price to be decreasing as output goes up, the industry must experience a negative factor price effect.

Long Run Increasing Cost Industry: Industries with an upward sloping run supply curve. For the zero-profit price to be increasing as output goes up, the industry must experience a positive factor price effect or firms must have access of different technologies (or both).

Long Run Shutdown Price: The minimum of the average cost of a firm.

Market Power: The ability of an individual firm or consumer to affect prices in the markets in which they participate.

Monopoly: A market with only one supplier.

Monopsony: A market with only one consumer.

Partial Equilibrium Analysis: Economic analysis which considers each market in isolation and ignores the effects that other markets might have.

Price Makers: Agents with market power.

Price Takers: Agents without market power.

Short Run Shutdown Price: The minimum of the average variable cost of a firm.

Sunk Costs:, Cost which can never be avoided. That is, a resource that has been committed and which can no longer be uncommitted or recovered even by going out of business.

Problems

1. What is a competitive firm? What condition must be satisfied for a firm in a competitive industry to choose to keep producing (that is, not to exit)? Assuming this condition is satisfied, prove that for competitive firms, supply curve and marginal cost curve are the same.
2. What is the difference between fixed and sunk costs? Can you think of real world examples of fixed costs which are not sunk and sunk costs which are not fixed?
3. Consider the hot dog industry. The technology is freely available to anyone and anyone can enter with this identical technology. Also, there are no factor price effects in the industry. Suppose that the price of hot dog carts goes up, and that carts are a fixed factor. What is the effect in the short run and long run on prices and quantities for the firm and industry? Use pictures to justify your answer.
4. Suppose the market price for Jeff Bezos Bobble-Heads (JFBH) is \$50. ACME Bobble-Head Consortium has a total cost function at fixed factor prices given by: $TC(Q) = 625 + Q^2$.
 - a. What is the marginal cost of producing JFBHs for this company? What is the optimal number of JFBHs for ACME to produce?
 - b. At this equilibrium quantity, what is the profit the ACME makes? What is the average cost of production?
5. Answer the following questions regarding a competitive firm.
 - a. Why is it that the long run profit for a competitive firm is zero?
 - b. Why would any entrepreneur start a business in an industry when he knows he would get zero profit as a result?
 - c. It is always optimal for a competitive firm to operate where price equal marginal cost?
 - d. What is the relationship between the long run marginal cost function and the long run supply curve for a price taking firm?
 - e. When does a competitive firm shut down in the long run? When does it shut down in the short run?
6. Suppose the price of butter goes up.
 - a. What happens to the demand for butter, and why?
 - b. Suppose butter is a constant cost industry. Using your conclusion about the change in demand from (a), what happens to the price and quantity of bread in the short run? What happens in the long run?
 - c. What would happen in the short run in the demand for butter goes up?
7. The major input to college teaching is labor. At any given time, there are only a certain number of qualified teachers available. Getting more teachers requires raising wages to the point that people are willing to switch careers into teaching and more students choose to enter teaching

programs. For the next decade, the demographic trends show that the college age population will increase each year. Assume that college education is provided in a strictly competitive market (I know this is not true, just assume it anyway). Given this information, what do you expect to happen in the long run and short run to the price and quantity of college degrees?

Chapter 6. Equilibrium

Section 6.1. Partial Equilibrium

We are now ready to put this together and see what a partial equilibrium competitive outcome looks like. Informally, a market is in equilibrium if, at the market price, the aggregate quantity supplied by firms equals the aggregate quantity demanded by consumers. Mathematically, this means we must find a P^{FM} that satisfies this equation:

$$S(P^{FM}) = D(P^{FM}).$$

Such a price is said to be **market clearing** because it leaves neither *excess demand* nor *excess supply*. Graphically, the quantity supplied and demanded both equal Q^{FM} at price P^{FM} . These free market prices are said to “clear the market”. On the other hand, if the price is somehow set at $P^{TooHigh}$, then the quantity supplied exceeds the quantity demanded and there is excess supply, or equivalently, a **surplus**. Conversely, if the price is somehow set at P^{TooLow} , then the quantity demanded exceeds the quantity supplied and there is excess demand, or equivalently, a **shortage**.

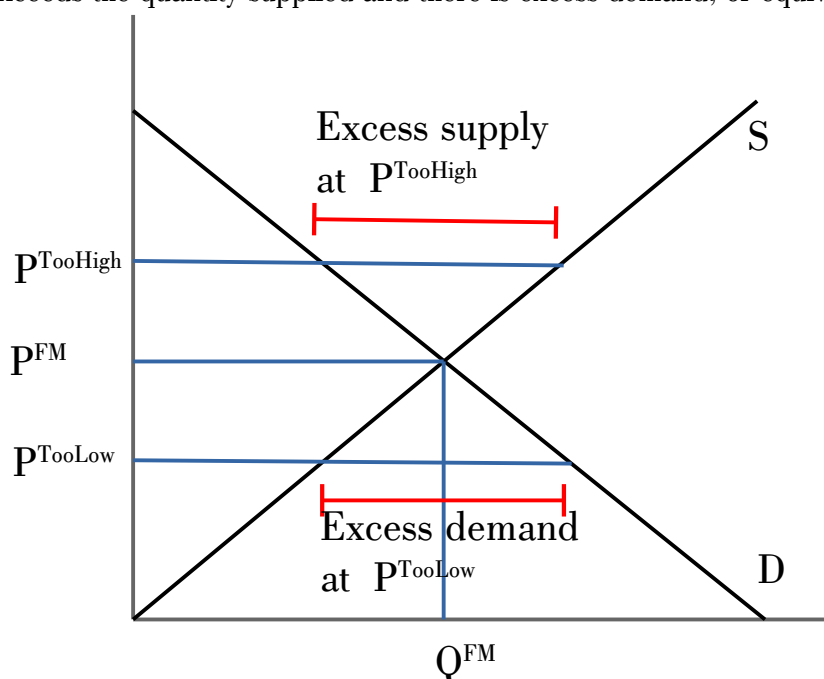


Figure 61: A market adjusting prices to get to equilibrium

To take this one step further, suppose that the government imposed a tax of Tx dollars per unit on consumers in this market. This is called a **sales tax**. If the tax were imposed on producers, it would be called an **excise tax**. Such taxes can be per unit, percentage, non-linear, or even more complicated. A universal effect of taxes is to drive a wedge between what the consumer pays for a good and what the producer receives.

In the case of a sales tax, the consumer pays P to the firm for a unit of the good, and then also has to pay a tax of Tx to the government. This means that the **consumer price** (the total cost to the consumer) is higher than the **producer price** (the revenue the producer gets to keep) by exactly Tx .

In the case of an excise tax, the consumer pays P to the firm for a unit of the good, but the firm only gets to keep $P - T$ after the firm pays the tax of Tx to the government. Again, the consumer price is higher than the producer price by exactly Tx .

Note that the amount paid by the consumer directly to firm (called the **nominal price** since this is shelf price, or what is written on the price tag) equals the producer price in the case of the sales tax, and the consumer price in the case of an excise tax.

To summarize:

Sales Tax: A tax that a consumer is legally obliged to pay to government when he purchases a good. We denote a tax as Tx regardless of whether it is a sales or excise tax.

Excise Tax: A tax that a producer is legally obliged to pay to government when he sells a good.

Consumer Price: The net money price paid by a consumer to all agents (including the government and the producer) for a unit of a good. This includes all taxes and subsidies and will be denoted P^{CON} .

Producer Price: The net money price received by a producer from all agents (including the government and the consumer) for a unit of a good. This includes all taxes and subsidies and will be denoted P^{PRO} .

Nominal Price: The amount of money that the consumer gives to the producer in exchange for a unit of a good (or, symmetrically, the amount of money a producer receives from the consumer for a unit of a good). This will be denoted P^{NOM} .

To illustrate, suppose the industry supply and demand curves are the following:

$$D(P) = 100 - 2P \quad \text{and} \quad S(P) = 3P.$$

We can easily solve for the equilibrium price and quantity:

$$100 - 2P = 3P \Leftrightarrow 5P = 100 \Leftrightarrow P = 20 \quad \text{and} \quad Q = 60.$$

Now suppose that the government imposes a sales tax of $T = \$10$. Then the producer price is the same as nominal price (since the producer keeps everything paid to him by the consumer),

while the consumer price equals the nominal price plus the tax: $P^{CON} = P^{PRO} + Tx = P^{NOM} + Tx$ (since the consumer has to pay to nominal price to the producer and also the tax to the government). We can solve for the after tax equilibrium as follows:

$$D(P^{NOM}) = 100 - 2(P^{NOM} + 10) = 3P^{NOM} = S(P^{NOM}) \Leftrightarrow 5P^{NOM} = 80 \Leftrightarrow P^{NOM} = 16 \text{ and } Q = 48.$$

By imposing a tax, the government raises the cost of the good to consumers while lowering the revenue to producers. We know that the gap is equal to Tx , but what part of this is absorbed by consumers in the form of higher prices compared to by firms through lower revenue? In our example, consumer price went from 20 to 26, while producer price went from 20 to 16.

Thus, consumers bear 60% of the tax burden while producers only bear 40%. Note that the slope of the demand curve is -2 , while the slope of the supply curve is 3 . This means that at any given price/quantity pair, the own price elasticity of demand is lower than the own price elasticity of supply. The general rule is this: *the relatively less elastic side of the market pays a greater fraction of any tax.*

Intuitively, a consumer is less elastic if it is difficult to find substitutes for given good. For example, a diabetic needs a certain amount of insulin. Dropping the price of insulin would not induce a diabetic to demand more (why would he need it?) and increasing the price would not induce him to consume less (he would get sick or die without his required dose).

Thus, he is inflexible and insensitive to price changes. He bears more of the tax burden as a result. If there are many substitutes for a given good, a consumer tends to be more elastic. For example, if the price of nectarines moves above the price of peaches, the consumer may choose to eat only peaches and no nectarines. Such a consumer is very price sensitive and flexible and so reduces the quantity demanded of a good by great deal when price goes up. He bears less of the tax burden as a result.

A price elastic supply curve is flatter and is usually seen when the costs of production do not vary much as the quantity changes (consider a long run constant cost industry, for example). A price inelastic supply curve is steeper and is usually seen when it is difficult to change the quantity produced regardless of the price. (Consider tickets for a football game. The number of available seats can neither be increased nor decreased without very costly renovations to the stadium.)

Graphically, the effect of a sales tax is to lower the demand curve (as a function of the nominal price) by exactly Tx since paying $P - Tx$ to the firm under the sales tax system is exactly like paying P to the firm before the tax was imposed. You can see that the producer and nominal prices are the equal since the shelf price (nominal price) is paid directly by the consumer to the producer. Since the producer has no tax liability, he gets to keep it all.

On the other hand, the consumer price equals what he has to pay the producer (the nominal price) plus the sales tax. Thus, the consumer price is above the nominal price by exactly Tx . Over

all, the effect is to drive a wedge between the producer and consumer price equal to T_x , and also to lower both the quality supplied and demanded in the market.

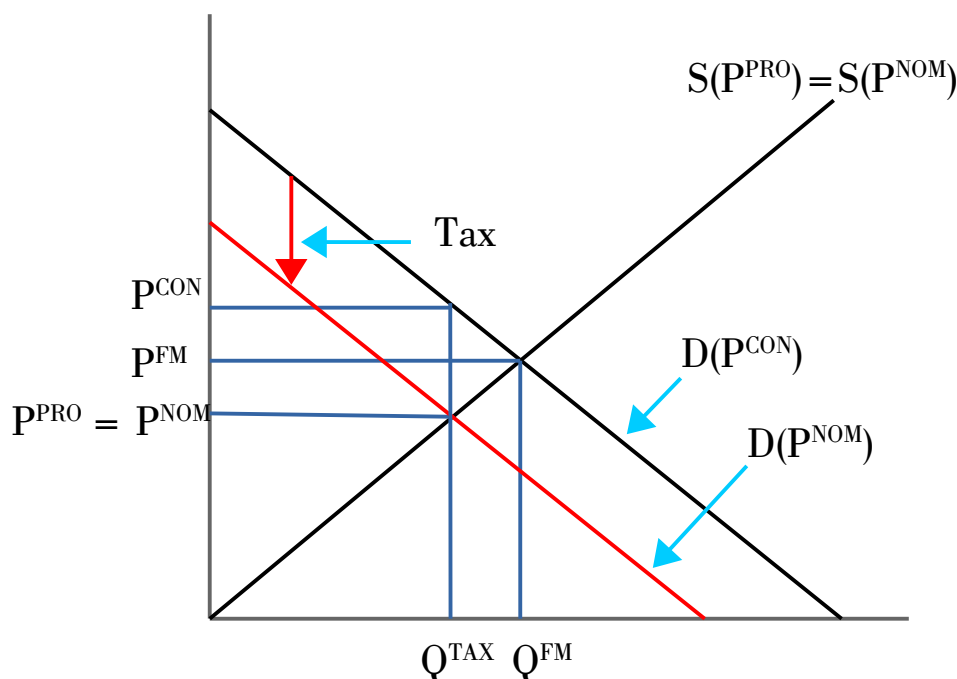


Figure 62: The effects of a sales tax

The effect of an excise tax is to raise the supply curve (as a function of the nominal price) by exactly T_x since receiving $P + T_x$ from the consumer under the excise tax system is exactly like receiving P from the consumer before the tax was imposed. You can see that the consumer and nominal price are equal since the shelf price (nominal price) is paid directly by the consumer to the producer. Since the consumer now has no tax liability, this is all that he has to pay to purchase the good.

On the other hand, the producer price equals what he receives from the consumer (the nominal price) minus the excise tax. Putting this together, an excise tax in the amount of T_x affects the both consumer and producer prices in exactly the same way as a sales tax of T_x . It also lowers both the quality supplied and demanded in the market by exactly the same amount. The only difference is that the nominal price becomes the equal to the consumer price instead of the producer price. The graph below illustrates this.

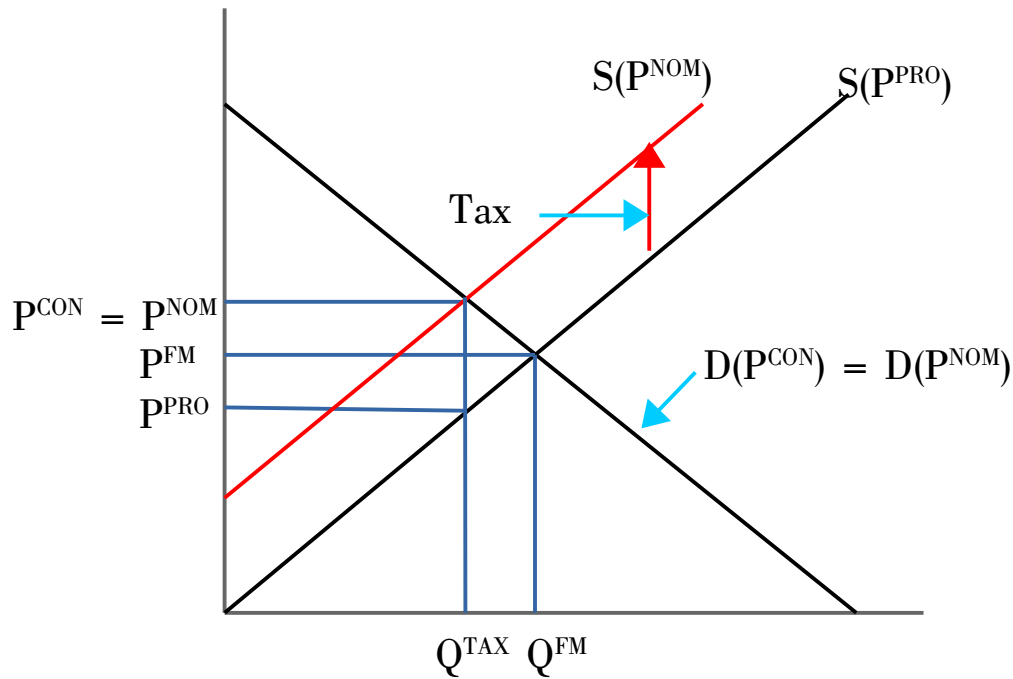
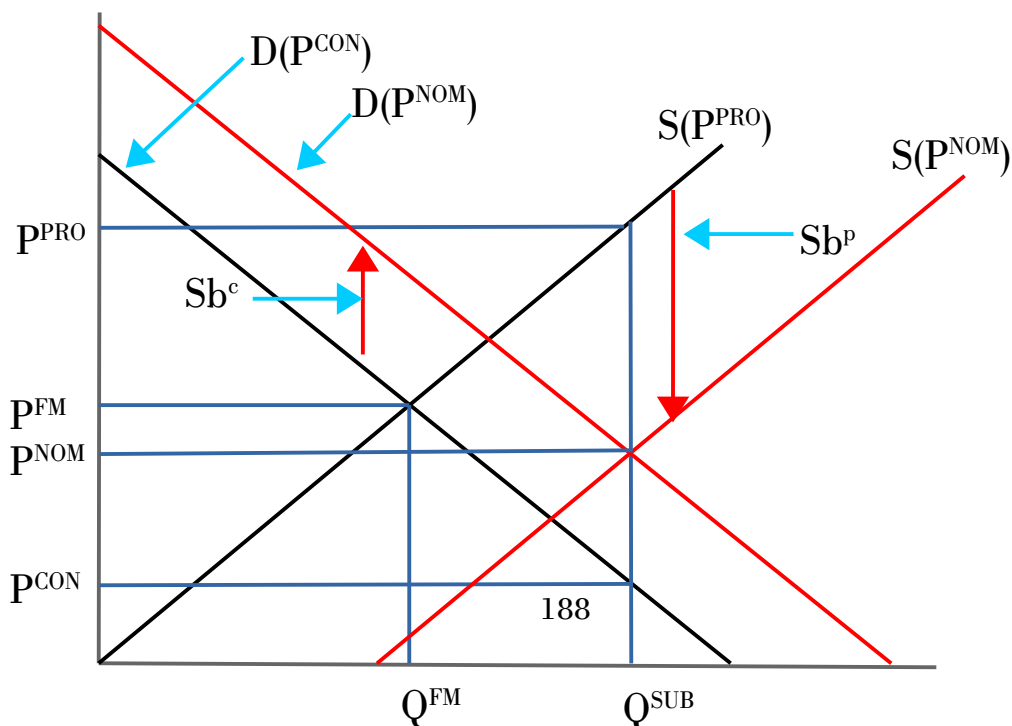


Figure 63: The effects of an excise tax

An important implication of this argument is that from an economic standpoint, it really does not matter who has the legal responsibility of paying the tax. The impact on the real prices paid by agents on each side of the market is the same, as is the overall quantity. Only the amount of the tax makes a difference. We summarize this in a key economic principle:

**THE ECONOMIC INCIDENCE OF A TAX IS INDEPENDENT OF ITS
LEGAL INCIDENCE.**



As you can see, the consumer price goes down, the producer price goes up, the nominal price does not equal to either one, and the quantity goes up. To understand this, consider the following argument:

- The nominal price is found at the intersection of the distorted supply and demand curves (which are functions P^{NOM}). Note that this is also true in the two tax examples above. The difference is that only one curve was distorted instead of two.
- If the consumer pays P^{NOM} to the producer and then receives a subsidy of Sb^c for each unit he buys, the net price he pays is $P^{CON} = P^{NOM} - Sb^c$, which means that quantity demanded is Q^{SUB} .
- If the producer receives P^{NOM} from the consumer and also receives a subsidy of Sb^p for each unit he sells, the net price he receives is $P^{PRO} = P^{NOM} + Sb^p$, which means that quantity supplied is Q^{SUB} .
- We conclude that under this policy of subsidizing both sides of the market, a nominal price of P^{NOM} causes supply to equal demand.

Other cases to consider are markets where both sides are taxed (social security taxes are partly paid by workers, and partly paid by employers, for example) and where one side is taxed, and the other is subsidized (tobacco farmers are subsidized, but cigarettes are taxed). We leave these as exercises for the reader.

Section 6.2. Decentralization

Markets depend on choices made by individual economic actors without any centralized direction. An alternative approach is to centralize control and run an economy by command and control. To see how this might work, consider a problem that might have been faced a central planner in the now defunct USSR. Suppose that the planner wanted to produce 1,000,000 telephones for the people using the three telephone factories that Comrade Stalin ordered to be constructed. How would he do this? He might propose a production plan given by the red dashed lines and italic numbers (in thousands) in the figure below:

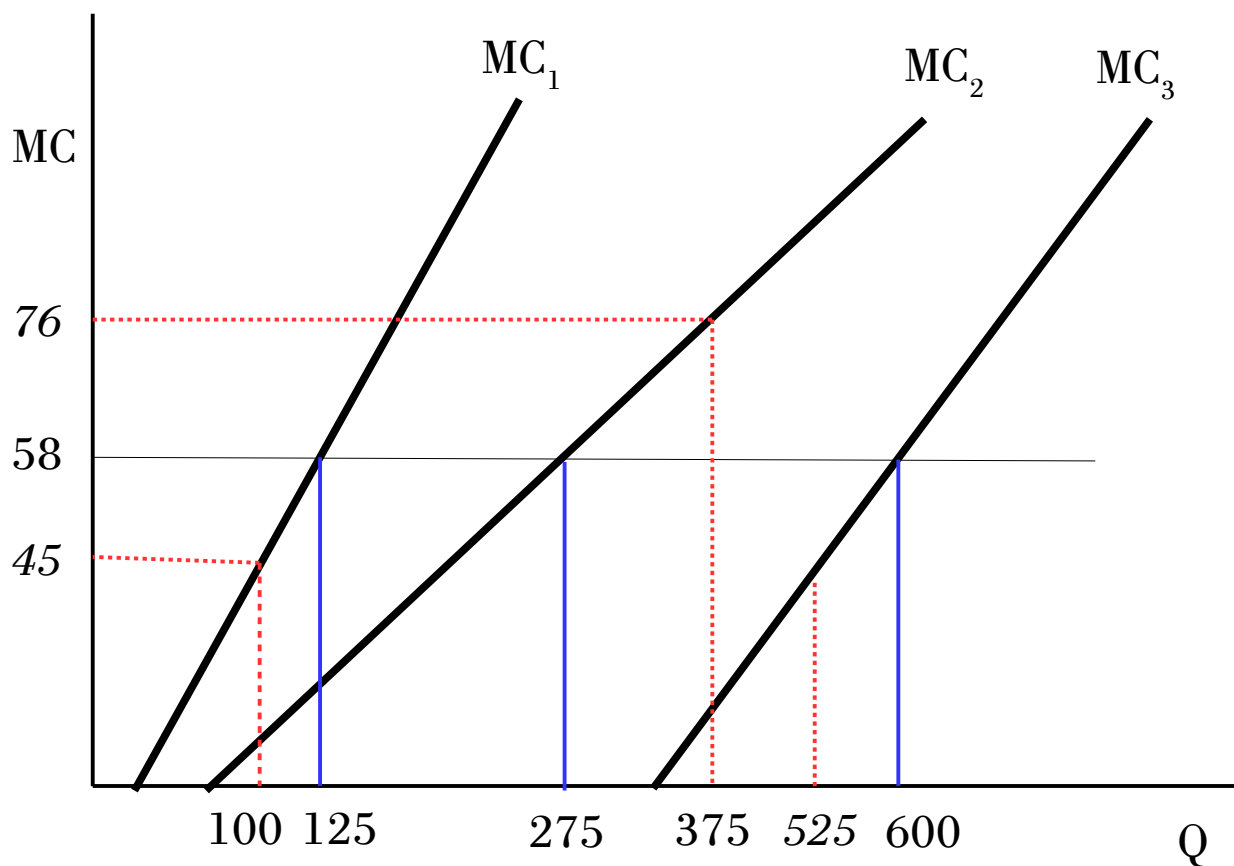


Figure 65: Optimally choice of production levels over factories

Now suppose the planner somehow knew the MC curves of each factory. What if he transferred the production of one telephone from factory 2 to factory 1? This would save 76 rubles because of the reduced output in factory 2, but increase spending by 45 rubles because of the increased production in factory 1. Thus, the USSR would save 31 rubles in total by moving production from one plant to the other. You can immediately see that as long as marginal costs are different between factories, it is possible to produce the same number of phones at lower costs by transferring produc-

tion from high MC factories to low MC factories. When MC is equal at all factories, then the planner is producing his target number of phones at the least possible cost. This optimal production plan is shown as the blue solid lines in the figure above.

Although the idea of equating marginal costs over factories is rather straightforward, it is quite difficult for a central planner to do so. First, he has to know the exact shape of each factory's MC curve. Even if the factory manager wanted to report this, he probably would not know what it looked like. A factory manager might know his current costs, but asking him to speculate on what costs would be if he doubled or halved production is not likely to get very high quality answers.

Second, even if the managers knew their costs, there is no incentive to report them correctly. Remember that managers of factories in a centrally planned economy are just government employees. They do not get to keep any profit the factory might make. Thus, they have little incentive to minimize costs. Instead, they are likely to **feather-bed**, meaning that they will claim that production takes more resources than it really does and then use the surplus inputs in other ways. This might be by selling the surplus on the black market, allowing workers to quit early or arrive drunk, or just not taking care to manage the production process very well and simply wasting the excess inputs.

If the planner could overcome these problems somehow, he would still need to figure out how many phones to produce. How important are phones compared to radios or cars? The planner has to make such choices since resources are limited. Once the planner decides on a production target, he has to decide how to allocate the phones. Does he give one to the prime minister, his wife, his mistress? Does each police station, fire station, and hospital get one? What about bars, barber-shops, and bicycle factories? If he asks people to tell him how much they need a phone, they are likely to lie since the more they say they need one, the more likely it is that they will get one.

The upshot is that it is an all but impossible task to get the information needed to make an optimal social decision. The information is often unavailable, too costly to collect, and not likely to be honestly revealed to the planner. How can the planner figure out how to balance production over different goods, where they should be produced, and who should be allowed to consume them in such an environment? As the Russians say: Забудьте об этом ! Давайте выпьем водки !³

What about planners in competitive economies? But wait, there are none! Then how are these decisions made? The answer is that the decisions are made by thousands of individuals and coordinated through the price system. This process is called *decentralization*.

Decentralization means that agents in the economy are not coordinated by any central authority but instead make autonomous individually optimal decisions based on equilibrium prices.

To see how this works, suppose the market price of a telephone is 58 rubles. Observe the following:

- Each of the factories in a decentralized economy takes this price as given and maximizes profits. This means that each factory produces telephones exactly to the point where $MC =$

3 Forget about it! Let's drink vodka!

58. Thus, the price system coordinates phone production over factories to get any given output level at the least possible cost. No information needs to be collected by any central authority.

- Firms have an incentive to minimize costs since the managers or stockholders get to keep the profits.
- On the demand side, if the price is 58, agents buy a phone if and only if the marginal benefit they get from owning one is at least 58 rubles. Low value users have no incentive to lie about how much they need a phone if they have to pay for it.
- The quantity demanded also shows how many people value phones this much. If, for some reason, more people wish to buy phones at this price than are produced by firms, it is a signal that too few phones are being produced compared to cars or radios. Fortunately, the planner does not have to figure out how to move resources around in response. If demand exceeds supply, the price goes up. Firms increase production (and firms in other industries decrease production as the rising price of inputs signals that those resources have more valuable uses elsewhere). Consumers decrease consumption as well since the higher price signals users on the margin that they are tying up resources that are worth more in their alternative use than the marginal benefit they would get from phone ownership.

Here is the amazing thing: all of this complicated coordination requires no central planner and no collection of information from agents. The market can “figure out” how much of each good to produce, where to produce it, and who should consume it, and also coordinate all producers and consumer to implement this “plan” using only one simple number: a price. This process of agents serving the collective good in apparently coordinated way enough though they are only doing what is best for themselves in reaction to market prices is sometimes referred to as the “invisible hand” As the Americans say: Awesome! Let’s get a beer!⁴

4 Потрясающе ! Давайте получить пиво!

Section 6.3. Competitive Markets and General Equilibrium

So far, we have considered single markets in isolation. We ignored any cross-linkages between markets. In this section, we recognize that markets are in fact linked. Changes in prices, costs, demands, and so on, in one market will have effects in other markets. For example, if gas prices go up, we would expect that the demand for cars would go down, and so the price and quantity of cars would decrease, which in turn would affect the market for steel, coal, rubber, etc.. We refer to these as “general equilibrium effects”.

In our partial equilibrium analysis, we said that the market was competitive if all agents were price takers. This is really a bit too reductionist. In the general equilibrium context, many real world conditions can lead to a noncompetitive outcome and market failure. Below we list the major requirements for markets to work along with the most common things that cause these requirements to be untrue in the real world.

1. Small agents (no market power)
 - a. Monopoly or oligopoly
 - b. Increasing returns to scale
 - c. Monopsony
 - d. Thin markets
2. No transactions cost
 - a. Frictions in markets
 - b. Fees (for example permit fees, stock and real estate brokerage fees)
 - c. Search costs
 - d. Cost of accessing a market (e.g. a farmer may have to walk miles to get to market)
3. Complete information
 - a. Sellers may have private information about the quality of what he is selling
 - b. Experience goods
 - c. Fraud
 - d. Not knowing where buyers, sellers, products, or markets are located
4. Complete markets
 - a. High transactions costs or incomplete information may cause markets never to open

- b. Government may legislate that certain markets be closed (illegal drugs, human organs)
- c. Incomplete property rights may shut markets because of nonexcludability (fishing rights, the Russian mafia taking profits from potential business owners)
- d. Public goods and externalities often cause market failures because, in effect, they imply that certain markets are effectively incomplete

Under these four conditions, competitive markets work well in general. To see how, we consider a simple model of a general equilibrium economy and assume that all the conditions above are satisfied in the next section.

Section 6.4. Pure Exchange Economies and the Edgeworth Box

Most of the basic ideas can be shown graphically in what is called a pure exchange economy with two agents and two goods. In such an economy, there is no production, so agents start with endowments and simply exchange goods with one another in a marketplace.

Cartesian Product

N-fold Cartesian Product: The set of ordered N-tuples where the n^{th} component is an element of the n^{th} set in a concatenated list:

$$(x_1, \dots, x_I) \equiv x \in X \equiv X_1 \times \dots \times X_I \equiv \prod_{i \in \mathcal{I}} X_i.$$

As an example the consumption plan is an allocation of each good for each agent and is denoted x . Since the consumption vector for each agent must be in that agent's consumption set, the plan as a whole must be in Cartesian product of these consumption sets over the index set for agents: $X_1 \times \dots \times X_I \equiv X$. while the combined consumption vector over all agents and goods is written:

$$(x_{1,1}, \dots, x_{1,n}, \dots, x_{I,1}, \dots, x_{I,n}) \equiv (x_1, \dots, x_I) \equiv x \in X_1 \times \dots \times X_I \equiv \prod_{i \in \mathcal{I}} X_i \equiv X \in \mathbb{R}^{I \times N}.$$

Note that we are making use of another piece of notation: $\prod$. This is very much like the summation operator, $\sum$, you are already familiar with, but instead of summing over an index set, $\prod$ is the product operator and so tells you to multiply the elements in the index set together. Thus, we indicate a production plan for the economy as:

$$(y_1, \dots, y_F) \equiv y \in Y \equiv Y_1 \times \dots \times Y_F \equiv \prod_{i \in \mathcal{F}} Y_f.$$

Subsection 6.4.1. The Model

Formally, the economy is defined as follows:

- Two agents A and B.
- Two goods 1 and 2.
- Agent A is endowed with $\omega_A = (\omega_{A,1}, \omega_{A,2})$ and similarly for agent B.
- We will denote the total endowment of each good as $\omega_1 \equiv \omega_{A,1} + \omega_{B,1}$
- We will denote the consumption of agents as $(x_{A,1}, x_{A,2})$ and $(x_{B,1}, x_{B,2})$
- We will denote the price of the goods as p_1 and p_2

We begin by drawing what is called an **Edgeworth box**. The height of the Edgeworth box is set equal to the total endowment of good 2 while the width is set equal to total endowment of good 1.

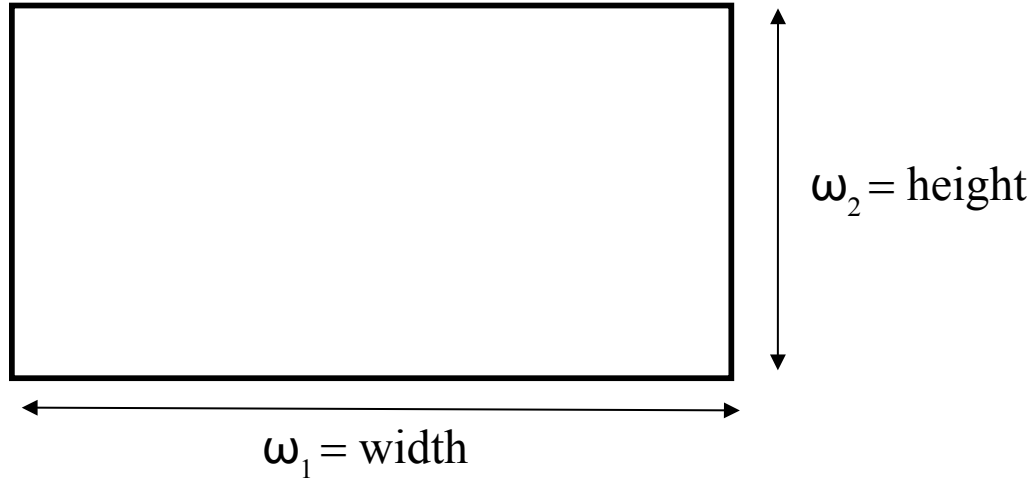


Figure 66: The dimensions of an Edgeworth box

We represent the endowment as a point in the Edgeworth box. The origin for agent A is taken as the lower left-hand corner, while the origin for agent B is taken as the upper right-hand corner. We measure the amount of good 1 in agent A's endowment as the horizontal distance from agent A's origin to the endowment point.

Similarly, we measure the amount of good 1 in agent B's endowment as the horizontal distance from agent B's origin to the endowment point. Notice that by construction, $\omega_1 = \omega_{A,1} + \omega_{B,1}$. The same is true for good 2, but we measure vertically.

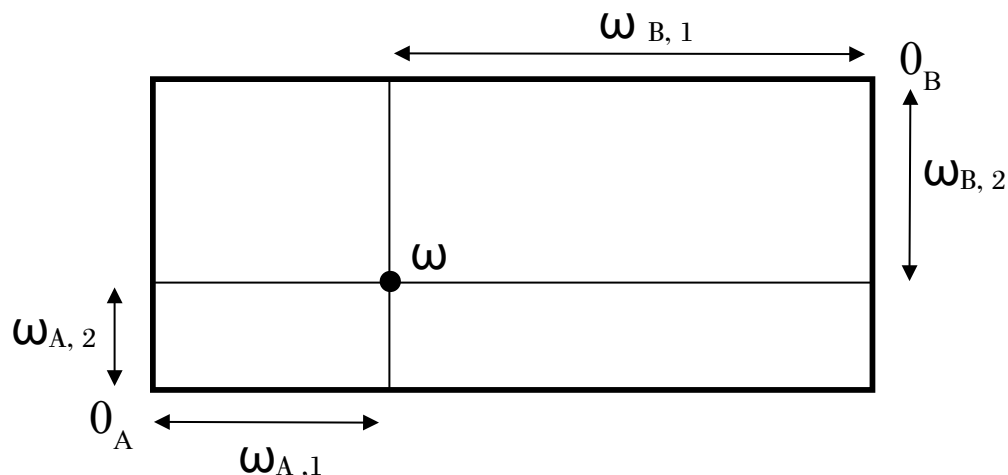


Figure 67: The endowment point in an Edgeworth box

The key thing about the Edgeworth box is that every point in the box represents a **feasible allocation** of the goods. For example, at point β it must hold by construction that $\omega_1 = \beta_{A,1} + \beta_{B,1}$ and $\omega_2 = \beta_{A,2} + \beta_{B,2}$

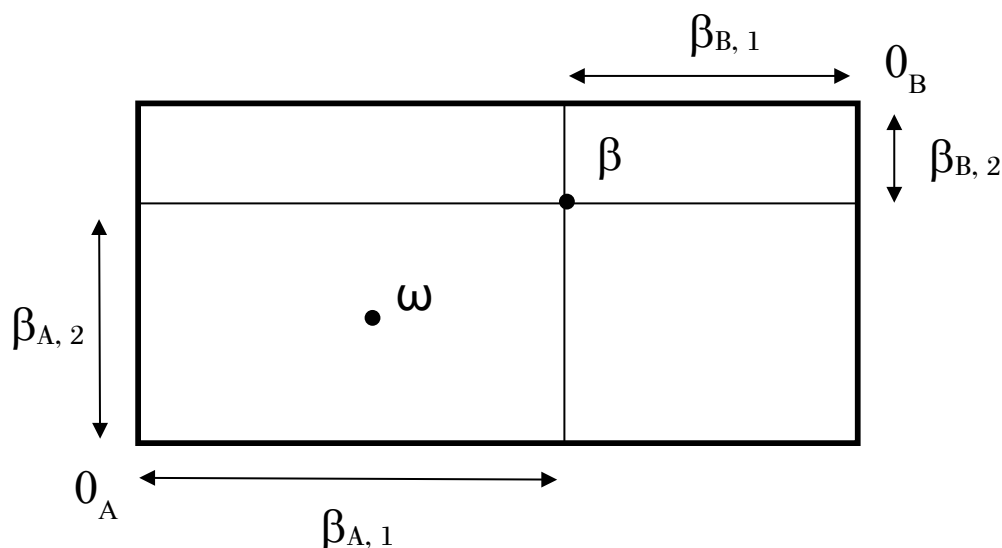


Figure 68: A feasible allocation over agents in an Edgeworth box

Subsection 6.4.2. Prices and Competitive Equilibrium

Prices are represented by a budget line. In the picture below, we show a budget line with prices $p_1 = 20$ and $p_2 = 10$. $p_1 = 20$ and $p_2 = 10$. Agent A trades from the endowment point ω to point β by buying two units good 2 for \$10 dollars each, and selling one unit of good 1 for \$20 dollars to pay for it. Thus, he moves up (buys good 2) and then left (sells good 1). Agent B trades from the endowment point to point β by selling two units of good 2 for \$10 dollars each and then buying

one unit of good 1 for \$20 dollars with the proceeds. He moves up (sells good 2) and then left (buys good 1). Notice that prices for each agent are represented by a common line. Their movements along this line are complementary. If they jointly move from the endowment point to any other common point on a given budget line, the sales of one agent exactly offset the purchases of the other. As a result, the market clears!

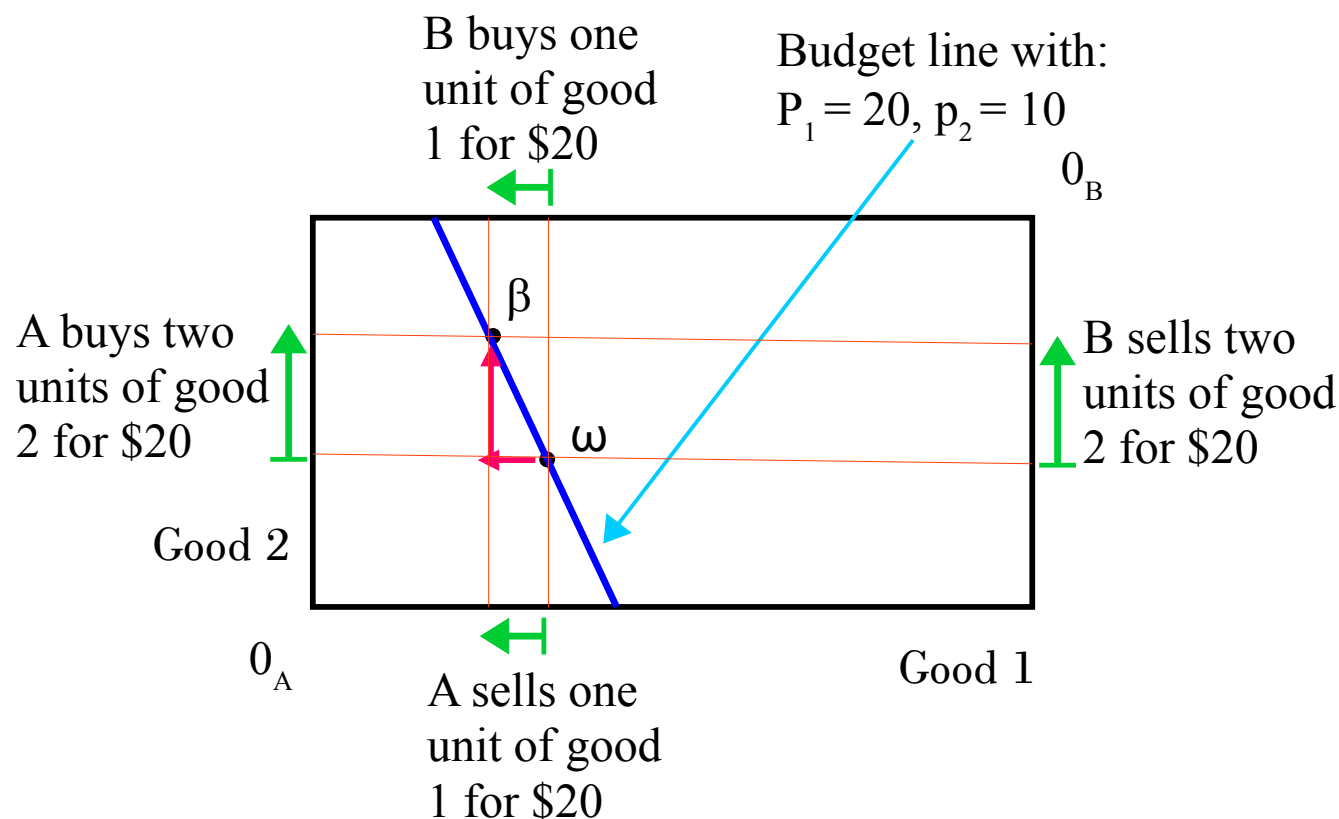


Figure 69: A budget line in an Edgeworth box

Next we need to introduce preferences into the box. For agent A, we just use the normal representation of preferences in a goods space using the lower left-hand corner of the box as the origin. For agent B, take the normal representation, but then flip it over by rotating it 180 degrees and using this to complete the box.

Putting this altogether, agent A has the red dashed indifference curves and most prefers agent B's origin. Agent B has the solid blue indifference curves and most prefers the agent A's origin since this gives him all of both goods.

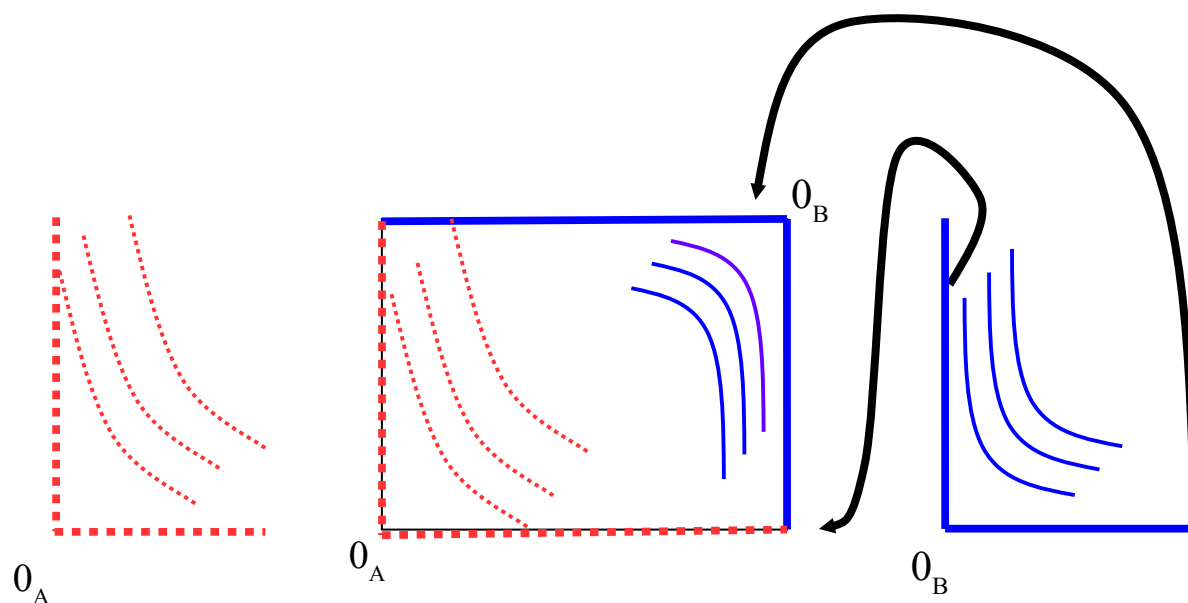


Figure 70: Constructing an Edgeworth box

Now consider the endowment point. All consumption points above agent A's indifference curve through the endowment (labeled IC_A in the picture) are preferred to the endowment point by agent A. Similarly, all consumption points above agent B's indifference curve through the endowment (labeled IC_B in the picture) are preferred to the endowment by agent B. Note that "above" in the case of agent B means in the direction of agent A's origin and away from agent B's origin.

This means that the red-hatched football-shaped area directly above the endowment point between the two labeled indifference curves contains feasible allocations that are preferred by both agents. This is called the **region of mutual improvement**. By moving anywhere into this region, both agents reach a higher indifference curve and so are better off.

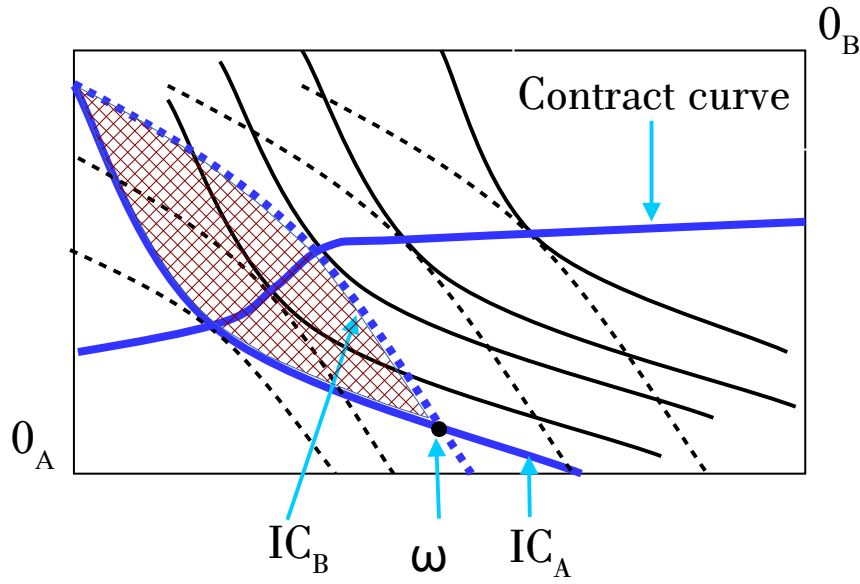


Figure 71: Pareto optimality and the contract curve

More generally, choose any allocation in the Edgeworth box and suppose that the indifference curves of the two agents through any specific allocation point cross each other. Then this point must have a region of mutual improvement associated with it. Suppose instead that indifference curves through a specific allocation point were tangent to each other at this point. Then the football of mutual improvement disappears! Collectively we call these points of tangency the **contract curve**. We say that an allocation is **Pareto optimal** if no other feasible allocation makes both agents better off. The contract curve is therefore the collection of all of these Pareto optimal allocations. More formally:

Let x and $\bar{x}$ be two feasible allocations in the two-person exchange economy. They stand in one of four possible Pareto relationships to one another:

Weak Pareto improvement: $\bar{x}$ is a weak Pareto improvement over x if $\forall i \in \{A, B\}, \bar{x}_i \geq x_i$ and $\exists j \in \{A, B\}$, such that $\bar{x}_j > x_j$.

In words, all agents i are at least as well off at $\bar{x}_i$ as they are at x_i , and at least one agent j is strictly better off at $\bar{x}_j$ as compared to x_j .

Strong Pareto improvement: $\bar{x}$ is a strong Pareto improvement over x if $\forall i \in \{A, B\}, \bar{x}_i > x_i$.

In words, all agents i are strictly better off at $\bar{x}_i$ as compared to $\hat{x}_i$.

Pareto equivalent: $\bar{x}$ is a Pareto equivalent to x if $\forall i \in \{A, B\}, \bar{x}_i \sim x_i$.

In words, all agents are exactly as well off at $\bar{x}_i$ as compared to $\hat{x}_i$.

Pareto unranked: $\bar{x}$ is a Pareto unranked with respect to x if $\forall i \in \{A, B\}$, and $\exists i, j \in \{A, B\}$, such that $\bar{x}_j > x_j, x_i > \bar{x}_i$, and $i \neq j$.

In words, at least one agent is strictly better off at $\bar{x}$ while at least one other agent is strictly better off at $\hat{x}$. Note that Pareto unranked is not the same as Pareto equivalent.

To give some additional vocabulary: If $\bar{x}$ is a **weak (strong) Pareto improvement** over x , we will also say $\bar{x}$ **weakly (strongly) Pareto dominates** x , and $\bar{x}$ is **weakly (strongly) Pareto superior** to x . We will also use **Strict** and **Strong** interchangeably. Pareto **equivalent** and **Pareto Indifferent** are sometimes used interchangeably as well.

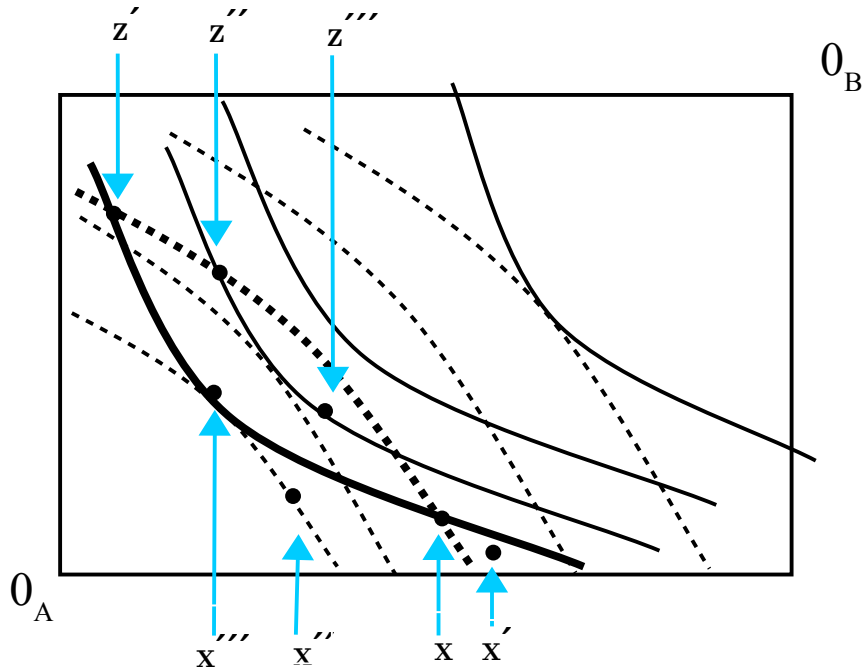


Figure 72: Strong and weak Pareto improvements

For example, consider allocation x . Relative to this:

- x' is strictly worse for both agent A and B. Thus, x' is strongly Pareto inferior to x since x strongly Pareto dominates x' .
- x'' is better for agent B, but worse for agent A. Thus, x'' is Pareto non-comparable to x . Note that this is an example of an incomplete binary relation since not all allocations can be ranked by the Pareto criterion.
- x''' is just as good for agent A, and strictly better for agent B. Thus, x''' weakly Pareto dominates x .
- z' is just as good for both agent A and B. Thus, it is Pareto indifferent to x .
- z'' is strictly better for agent A, and just as good for agent B. Thus, z'' weakly Pareto dominates x .
- z''' is strictly better for both agent A and B. Thus, z''' strongly Pareto dominates x .

Given this:

Weak Pareto optimality (WPO): A feasible allocation x is **weakly Pareto optimal** if there does not exist another feasible allocation $\bar{x}$ such that $\bar{x}$ is a strong Pareto improvement over x .

Strong Pareto optimality (SPO): A feasible allocation x is **strongly Pareto optimal** if there does not exist another feasible allocation $\bar{x}$ such that $\bar{x}$ is a weak Pareto improvement over x .

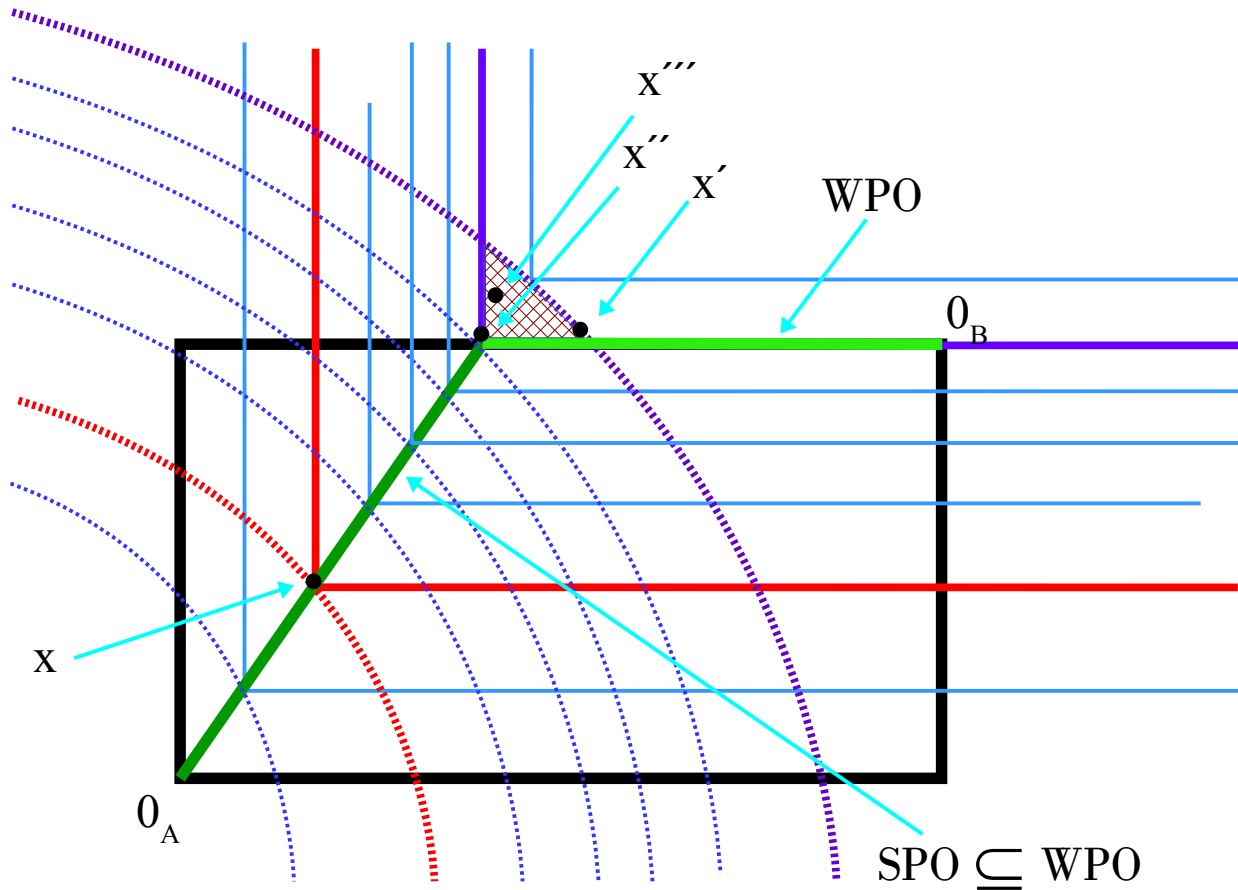


Figure 73: Strong and weak Pareto optimality

In the figure above, the SPO is the dark green line, while the WPO is the union of the dark green line and the light green line.

To see this note that we have drawn the indifference curves of both agents extend beyond the borders of the Edgeworth box. Although it is not feasible to give allocations outside the box (since this would require a larger quantity of some goods than exist in the entire social endowment), agents still have preference over such bundles. For example, you might prefer to have flying over heat-vision as your superpower. Sadly, you are not from the planet Krypton, and as a mere Earthling, you will never have either one. This does not change the fact that it would be really cool to fly!

Now consider the allocation labeled x . All the allocations that are at least as good as x for agent A are on or above the heavy, solid red indifference curve through x . All the allocations that are at

least as good as x for agent B are on or above (that is, in the southwest direction of) the heavy, dashed, red indifference curve through x .

Note that there does not exist any allocation that is strictly above one of these indifference curves while being on or above the other. Thus, there does not exist a weak (or strong) Pareto improvement over x . It follows that x is SPO and is therefore also WPO. This argument holds for all the other allocations on the dark green line through x .

Finally, consider the allocation labeled x' . All the allocations that are at least as good as x' for agent A are on or above the heavy, solid, purple indifference curve through x' . All the allocations that are at least as good as x for agent B are on or above (that is, in the southwest direction of) the heavy, dashed, purple indifference curve through x' . All allocations such as x''' which in the interior of the red hatched triangular region are therefore strictly better for both agents.

Unfortunately, all of these allocations are also outside the Edgeworth box and so are not feasible. Allocations like x''' are therefore *not* Pareto improvements over x . On the other hand, you can check that all allocations on the lower and left-hand boundaries of the triangle, excluding the corners, are better for agent B while being just as good as x' for agent A. The left-hand boundary is above the Edgeworth box and is therefore infeasible. However, the lower boundary is at the very top of the Edgeworth box and is therefore feasible. It follows that all the allocations between x' and x'' are weak Pareto improvements over x' .

Thus, point x' is WPO (since there are no feasible strong Pareto improvements), but not SPO (since there are feasible weak Pareto Improvements). The same argument can be made for any point to the right of allocation x'' on the upper boundary of the Edgeworth box. The allocations on the light green line are therefore all WPO.

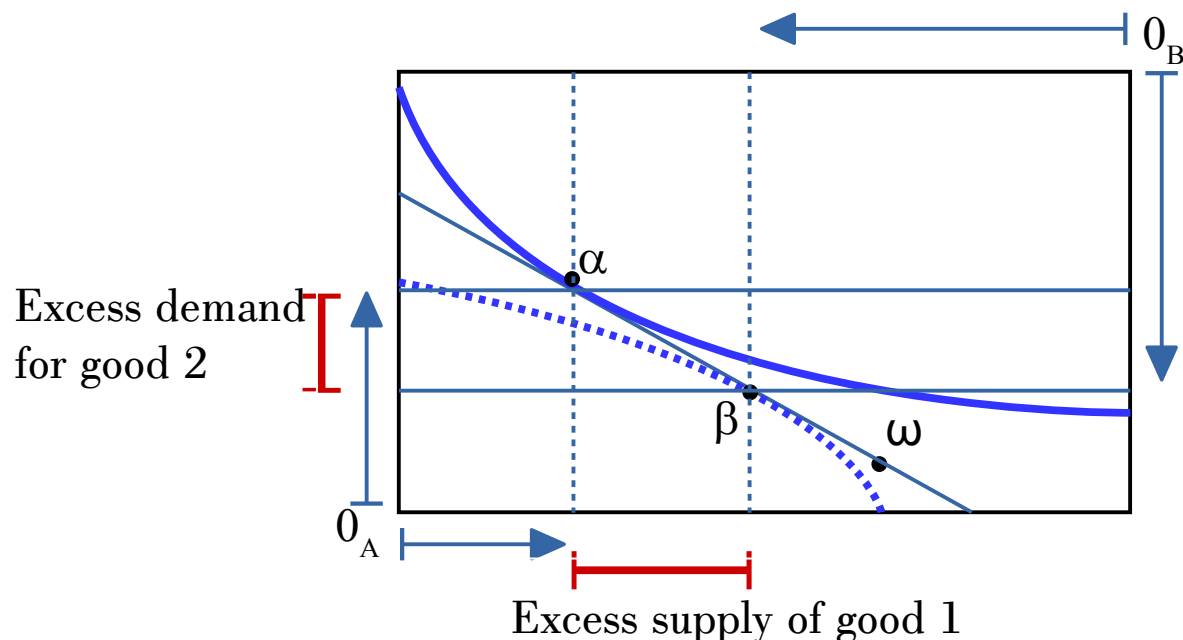


Figure 74: An example of nonequilibrium prices and markets that do not clear

Next, we add markets. When faced with the prices in the figure, agent A's most preferred point is α while agent B's most preferred point is β . As a result, when the agents each chose the most preferred points they can afford under these prices, there is an **excess supply** of good 1 and an **excess demand** for good 2. This implies that the prices shown in the figure below are not equilibrium prices since the market does not clear. We seem to have failed!

What we need are prices such that when the agents maximize utility, they end up asking for a *common point on the budget line*. This would mean that there is no excess supply or demand and markets clear. Under such prices we end up at what is called a **competitive equilibrium allocation** (CE).

In the figure below, it looks like there is only one CE starting from the endowment given in the figure. This could be the case, but it is also easy to show that there can be several competitive equilibria starting from a single endowment if preferences have the right shape.

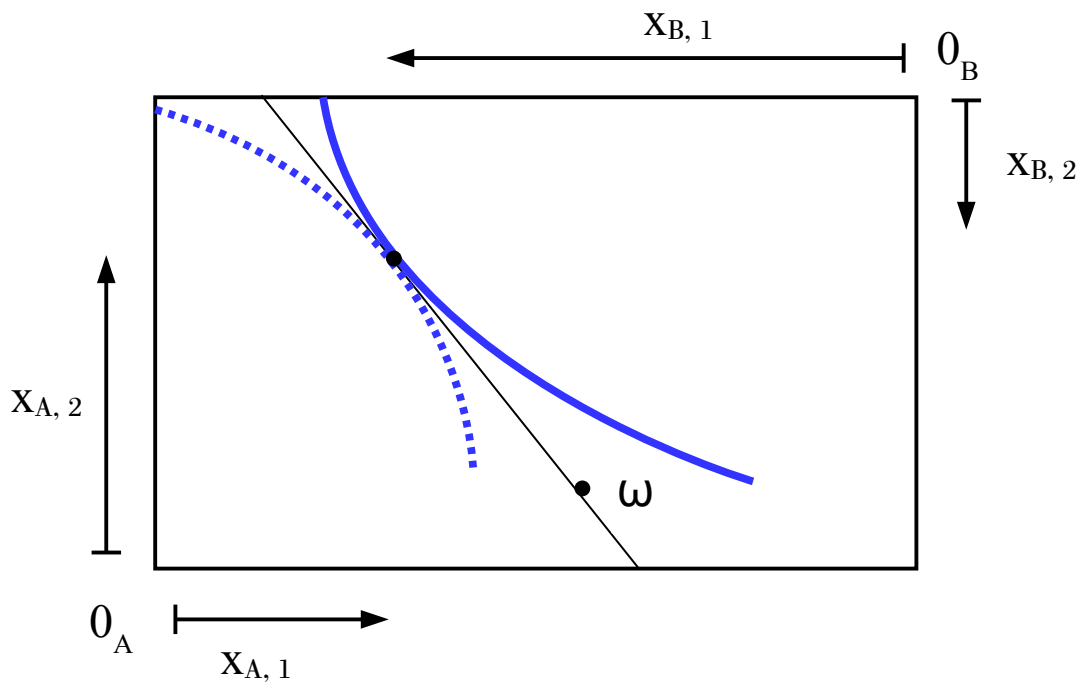


Figure 75: An example of a competitive equilibrium allocation and price vector

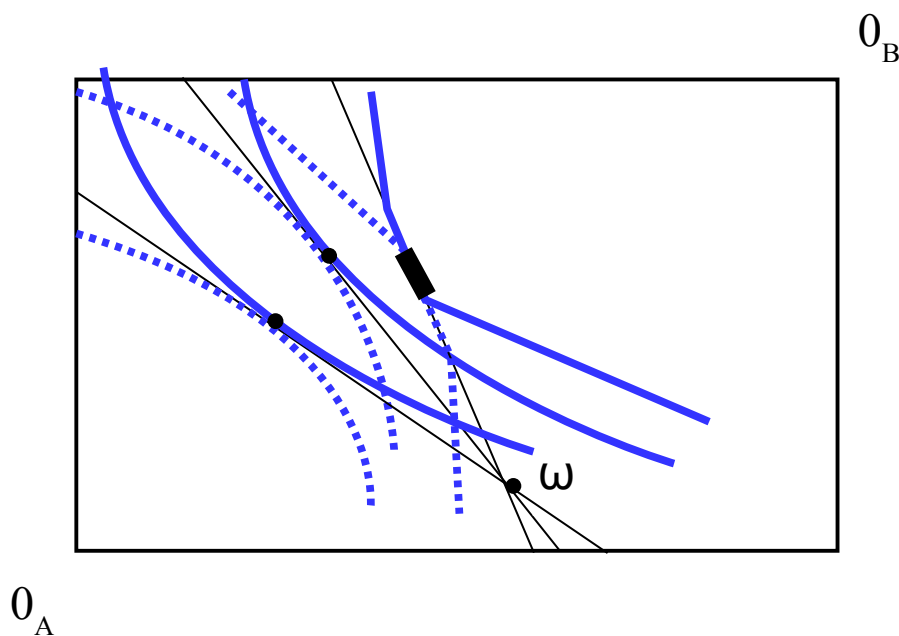


Figure 76: An example of a multiple CE from the same endowment point

For example, the Edgeworth box below shows three budget lines through the endowment point. The two shallowest cases show how we can have distinct CE at different prices. In the last case, we see an example with indifference curves that have flat portions and so a whole region of the budget line contains competitive equilibria (In fact, there are an infinity of CE here). The indifference curves are convex, monotonic, transitive, and in every other way admissible. It turns out that this example generalizes. Only under very restrictive conditions can it be shown that an economy has only one CE.

Subsection 6.4.3. Welfare Theorems and the Core

At last, we are able to state and show the two fundamental theorems about general equilibrium economies.

First Welfare Theorem: All competitive equilibria are Pareto optimal.

It is easy to see this. If an allocation is a CE then it must be that the indifference curves through the allocation are tangent to a common budget line at this allocation. It is immediate that the indifference curves must be tangent to each other at this point as well. Thus, the allocation must be PO.

(An aside: under relatively weak assumptions, ALL CE are WPO. Under somewhat stronger assumptions, all CE are SPO. The details are beyond the scope of this course, and so in such cases, we will simply use the term Pareto optimality (PO) and not further specify exactly which type we have in mind.)

The importance of this is that it implies that the free market does not waste resources. It is never possible to improve the welfare of all agents starting from a CE. If one agent gets a better allocation, it must come at the expense of another. This is a strong endorsement of the market mechanism. Finding a PO allocation in the abstract is quite difficult. Notice that the contract curve is a line. Since a line has no dimension, zero percent of the Edgeworth box is taken up by the contract curve. Thus, if you were to choose a feasible allocation at random, the odds that it would be PO are zero!

The fact that the CE is always PO is good, but one can still object. Notice as you move leftward along the contract curve, you move from allocations that are more preferred by agent A to those that are more preferred by agent B. In other words, it is nice that the CE is PO, but you may prefer a different PO allocation. What can we do about this?

Second Welfare Theorem: Any Pareto optimal allocation can be supported as a competitive equilibrium allocation for some reallocation of endowments.

To see this, choose any point you happen to like on the contract curve (which is to say, any PO allocation). Now, draw the tangent line to the two indifference curves through this point and take

this as the new budget constraint. Finally, choose any point on this budget constraint as the new endowment point. You can see that the point on the contract curve you chose is now a CE from this new endowment point.

In the figure below, ω is the initial endowment and CE is the initial competitive equilibrium. If you happen to like a point on the contract curve such as CE', all you have to do is move the endowment to a point like ω' or any other point on the budget line tangent to both indifference curves at CE'. Moving the endowment to a point like ω'' , on the other hand, supports the Pareto efficient point, CE'', as a competitive equilibrium which favors agent A rather than agent B .

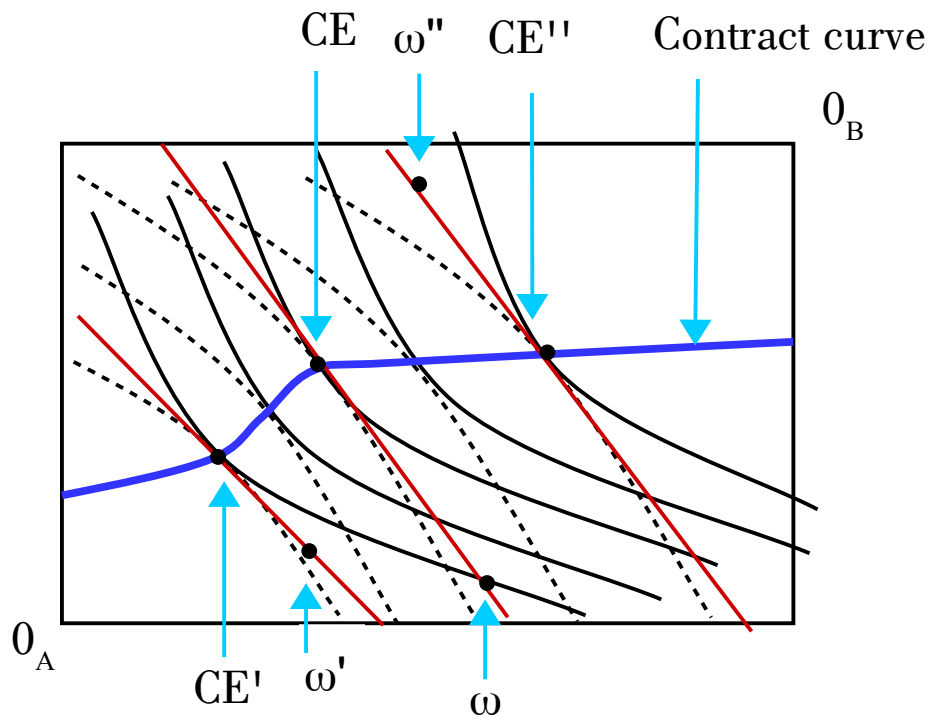


Figure 77: A graphical demonstration of the Second Welfare Theorem

The importance of this is the following: if you do not think the market allocation is fair or equitable, what you are really saying is that the initial allocations are unfair, not the market mechanism itself. The market is simply a neutral mechanism that selects a PO allocation. No matter what your social or ethical preferences happen to be, it should not be controversial that if it were possible to make *all* agents better off, then this would be a good thing. Thus, the market is a mechanism that not only takes society to a PO allocation, but in fact, can take society to *any* PO allocation you like, provided that you can adjust the initial allocations appropriately.

There is one important caveat to this. The second welfare theorem requires that it be possible to reallocate endowments between agents. If these are physical goods, there is no problem. However, most of the wealth you see around you is really based on human capital. Most rich people are rich because they are selling their skills, not because they inherited or own a great deal of physical property. To the extent that this is true, it is not possible to reallocate wealth! Smart, talented, and educated people will always have a better endowment and do better in the market.

The Pareto set is large, and the second welfare theorem says we can get any PO allocation as a CE if we reallocate wealth. Given this, what allocations are we likely to see in real world? To narrow this down, consider an equilibrium concept called the “core”. Roughly speaking, we say an allocation over agents is a **core allocation** if it is stable against all possible coalitional deviations.

A coalition would choose to deviate from a feasible allocation $\bar{x}$ if there exists an allocation that is feasible for the deviating coalition that all of its members prefer to $\bar{x}$. If such an allocation exists, it is said to **block** $\bar{x}$.

A feasible allocation $\bar{x}$ is a **core allocation** if it cannot be blocked (we give a formal and more general definition of the core latter in the text).

In the two-person exchange economy there are only three coalitions possible: $\{A\}$, $\{B\}$ and $\{A, B\}$.

If an allocation is not PO, then clearly it can be blocked by the grand coalition $\{A, B\}$. The two, one agent collations can block any allocation that puts them on a lower indifference curve than the one through their endowment. We say that an allocation is **individually rational** (IR) if all agents get at least as much utility as they could if they were to go off by themselves and form a one-person coalition. This means that in a two-person exchange economy, a feasible allocation is in the core if and only if it is both Pareto optimal and individually rational. It must therefore be an allocation that is both on the contract curve, and above the indifference curves through the endowment from the perspective of each agent. The figure below illustrates the core of an exchange economy.

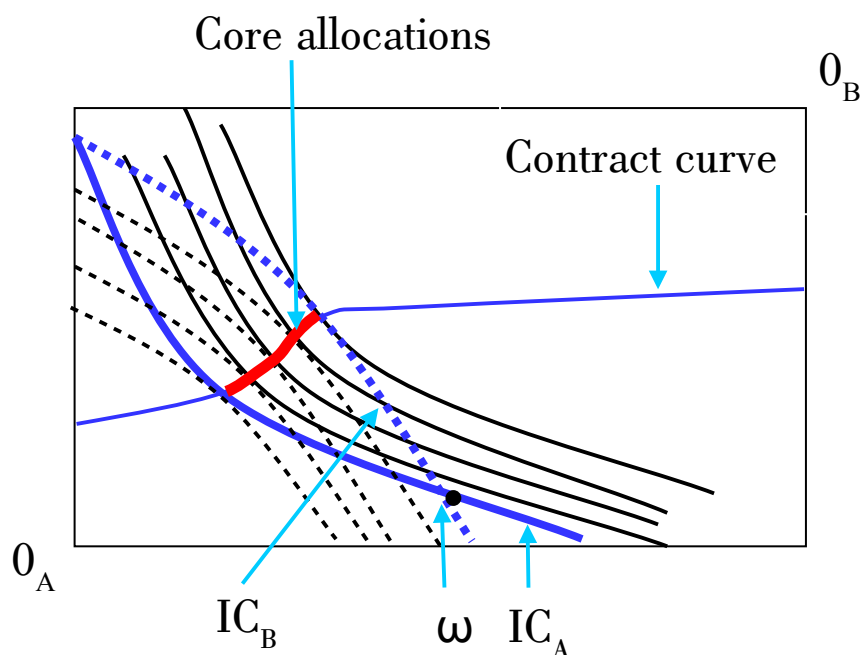


Figure 78: The core allocations of an exchange economy

In economies with more than two agents, we would also have to check that coalitions with more than one agent, but fewer than all the agents, could not block. In any event, this is important because of the following two theorems:

Theorem: All CE allocations are elements of the core.

Theorem: As an economy gets large in the sense of adding more firms and agents, the set of core allocations shrinks until it contains only the CE allocations. (The core converges to the CE.)

In other words, for large economies, only the CE allocation is stable against coalitions seeking to improve the well-being of their members. Unless we believe that agents or coalitions of agents will choose to forego these opportunities, we have reason to believe that regardless of how a society chooses to allocate resources (whether it uses the market mechanism or not), if society tries to impose something other than a CE allocation, it risks social unrest and instability.

Section 6.5. A Simple Economy with Production

We are about at the limit of what we can do with two-dimensional graphs. We have only one last thing to show before we resort to a general algebraic treatment of an economy.

Suppose Robinson Crusoe is trapped on a desert island by himself. He has twelve hours of daylight each day that he divides between catching fish and collecting coconuts.

If Crusoe spent all his time fishing, he might catch 100 fish, but collect no coconuts. If he decides to switch one hour over to collecting coconuts (a) he chooses to give up the hour when it is hardest to catch fish, and (b) he spends the one-hour collecting the lowest hanging coconuts that are closest to his camp. Thus, he might catch 97 fish and collect 15 coconuts. If he decides to devote a second hour to coconut collection, it must be one that is better suited to fishing, and he also must collect coconuts that are higher in the trees and further from camp. Thus, he might catch 92 fish and collect 25 coconuts in total.

From the discussion above, we draw two conclusions: First, $(100, 0)$, $(97, 15)$, and $(92, 25)$, are all feasible production choices for Crusoe. If we graph all of these feasible choices, we get something called the **production set**. This is just another example of a feasible set like the budget set discussion the previous sections. The boundary of this production set is called the **production possibility frontier (PPF)**. Second, there are diminishing marginal returns to labor in both fishing and collecting. As a result, the production set is convex. Another way of saying this is that the slope of the PPF, which we call the **marginal rate of transformation (MRT)**, is diminishing. The figure below illustrates this.

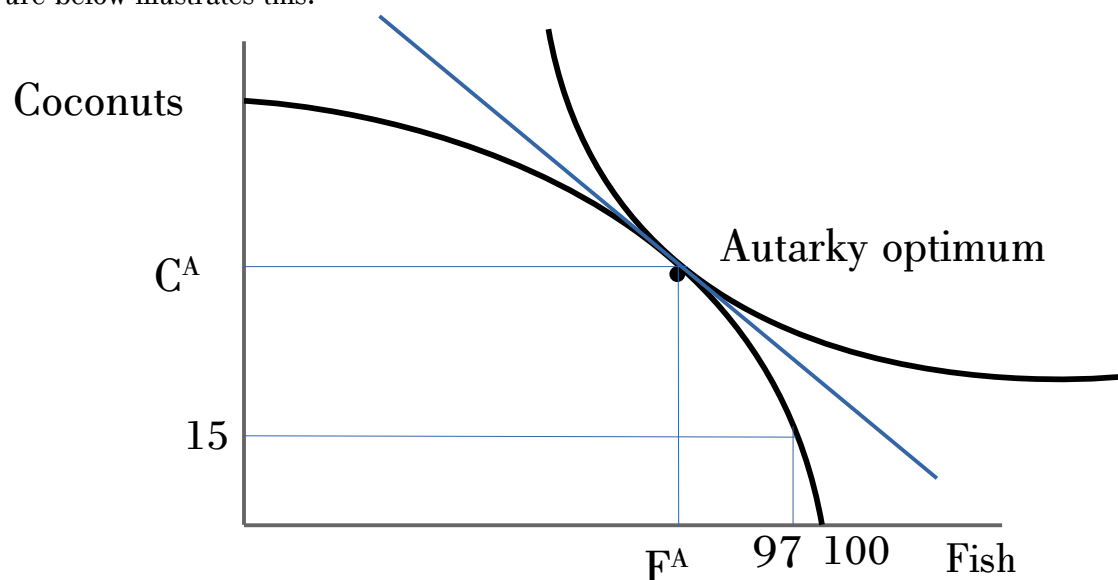


Figure 70: A PPF and optimal choice under autarky

If Crusoe is truly alone, he simply maximizes his own preferences taking the production set as his feasible set. The solution to this maximization is called the **autarky** optimum. Autarky in this context means “the government of one’s self”. Of course, the slope of the indifference curve and the PPF will be the same at the optimum ($MRS = MRT$).

Now suppose that Crusoe finds a raft and paddles to another island. There he finds a market at which he can trade fish and coconuts at market prices, and that these prices happen to be different from his own MRS. What should he do? Have a look the figure below and note the following:

- Suppose he chooses some point A on the PPF as his production plan. If he takes this output and rows over to the market, his trading opportunities can be represented by a budget line through point A having a slope equal to the relative prices in this market.
- His choice set has therefore expanded to include not only the allocations on the PPF, but all the allocations on this new budget line as well. His next step, therefore, is to choose the most preferred point over this entire set and then trade from point A to this new consumption bundle.
- Since his choice set depends on his production point, Crusoe’s best strategy is to pick a production plan that maximizes the expands his choice set as much as possible. In other words, he should produce point B since he can sell this for the largest amount of money at the market, and thereby have the highest budget possible to choose from. Of course, $MRT =$ market price ratio, at this optimum. Therefore, when Crusoe can trade he has two completely independent maximizations to perform. First, he chooses the production plan in his production set that will bring the most money in the marketplace. It does not matter at all what his preferences are. Second, he takes the budget line implied by the market prices and the income from selling his output, and chooses a utility maximizing consumption bundle. The figure below illustrates this:

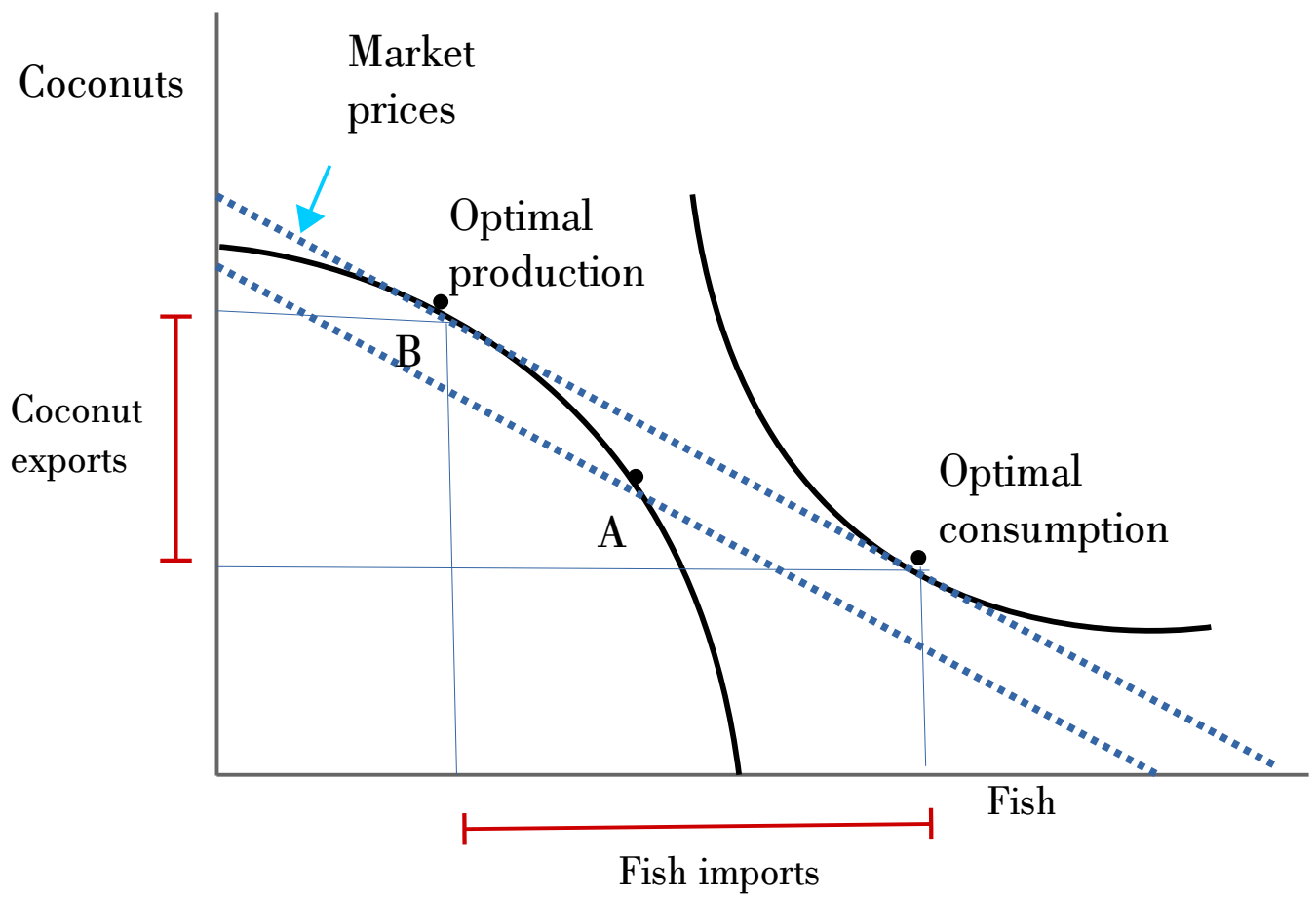


Figure 79: Optimal production and consumption choice with trade

Section 6.6. General Equilibrium, More Generally

The Edgeworth box is good for developing intuition. Fortunately, these results hold in a much more general context. In this section, we give a general mathematical statement of a competitive economy. Formally, an **Arrow-Debreu-McKenzie general equilibrium economy with production** consists of:

$n \in \mathcal{N}$	goods
$i \in \mathcal{I}$	agents
$\omega_i \in \mathbb{R}^N$	an endowment for each agent
$X_i \in \mathbb{R}_+^N$	a consumption set for each agent
$u_i: X_i \Rightarrow \mathbb{R}$	a utility function for each agent
$f \in \mathcal{F}$	firms
$Y_f \subset \mathbb{R}^N$	a production set for each firm

You may wish to have look at this section of the appendix to remember notation: [A.2.4.2: General Equilibrium Economies](#). Note that we could also have specified that each agent had a preference relation over the consumption set, instead of a utility function.

We denote a **production plan** for firm f as: $y_f \in Y_f$. By convention, if a component of the production plan is *negative*, $y_{f,n} < 0$, we interpret this as meaning that good n is an *input* to the firm since the net output is less than zero. If a component is *positive*, we interpret this meaning that the good is an *output* for the firm.

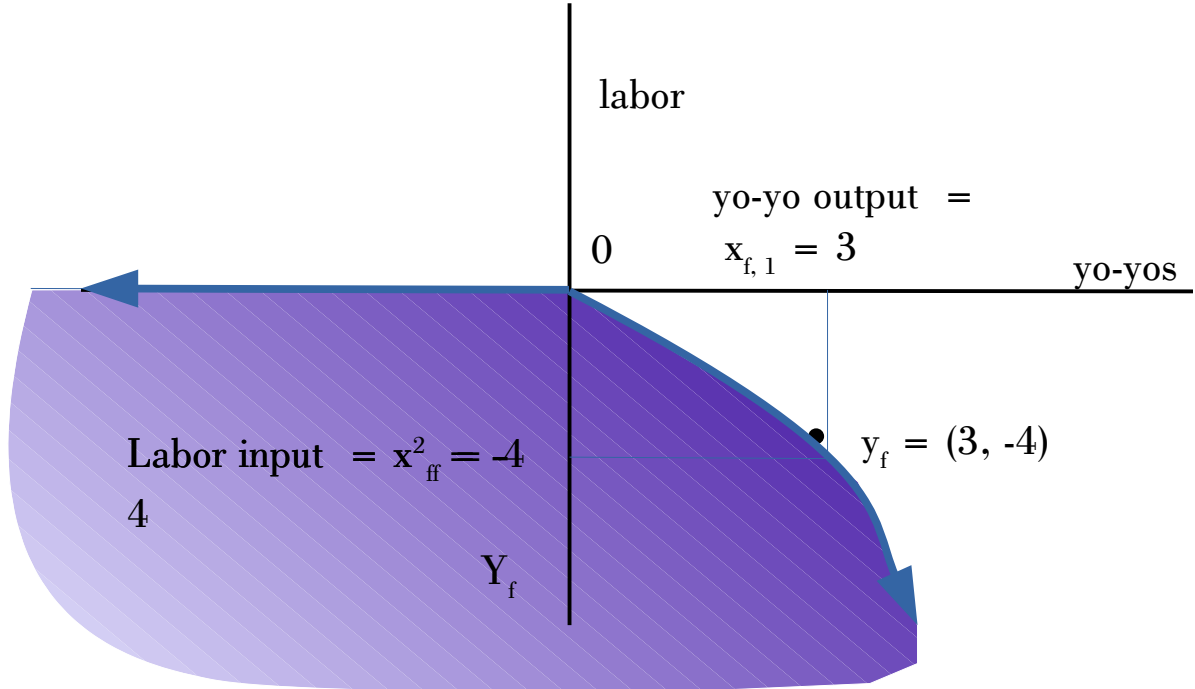


Figure 80: A production set for a firm

Firms are owned by agents. To keep track of **ownership shares** we use the following notation: θ_{fi} is the fraction of firm f owned by agent i . Thus, for any particular firm, the sum of ownership shares over agents must equal one and no share can be negative. More formally, the whole vector of ownership shares is denoted $\theta \in \Theta$ where:

$$\Theta \equiv \{(\theta_{1,1}, \dots, \theta_{1,I}, \dots, \theta_{f,i}, \dots, \theta_{F,1}, \dots, \theta_{F,I}) \in \mathbb{R}_+^{F,I} \mid \forall f \in \mathcal{F}, \sum_{i \in \mathcal{I}} \theta_{f,i} = 1\}.$$

For example, if $\theta_{fi} = 1$ and zero for all other agents, then firm f is owned entirely by agent i . If $\theta_{fi} = \theta_{fj} = .5$ and zero for all others, then firm f is owned by a partnership of agents i and j . If $\theta_{fi} = .00000000893$ and similar small numbers for a large group of agents in $\mathcal{I}$ and zero for the rest, the firm is owned by a widespread set of stockholders.

We know our previous discussion that all that matters are the relative prices of goods, not the absolute prices. This means we have a kind of trivial indeterminacy in equilibrium prices since if p is an equilibrium price vector, then so are $2p$, $10p$, and so on, because multiplying by any positive number keeps relative prices the same.

To pin this down, we often **normalize** prices. The easiest way to do this is to simply set the absolute price of the first good equal to one and then allow all other prices to be determined relative to this. When the **normalization** $p_1 = 1$ is imposed, we call good 1 the **numéraire good**. This normalization is generally used when preferences are quasilinear in the first good: $u_i(x_i) = x_{i,1} + v_i(x_{i,2}, \dots, x_{i,N})$. In this case, good 1 is often called a **transferable good**. One problem with this normalization is that the set of all admissible prices is unbounded. That is, we

might need to allow the price of some goods to get arbitrarily large to capture all possible relative prices.

An alternative normalization is to require prices to be an element of the **(N-1)-dimensional unit simplex** defined as:

$$\Delta^{N-1} \equiv \{ p \in \mathbb{R}_+^N \mid \sum_{n \in \mathcal{N}} p_n = 1 \}.$$

We will call this the “(N-1)-simplex” for brevity. If there are two goods, then the 1-simplex is simply the line segment between $(0,1)$ and $(1,0)$. Note that every element of this line segment has the property that the components are nonnegative and sum up to one. The reason this is called the 1-simplex is that a line segment is a one dimensional object (technically, a manifold) in a two-dimensional space. If there are three goods, the 2-simplex is a kind of two-dimensional triangle connecting $(0,0,1)$, $(0,1,0)$, and $(1,0,0)$. In general, the advantage of this normalization over using a numéraire good is that although it still allows all relative prices, these prices are now taken from a compact set which is sometimes useful in proving the existence of equilibrium.

Whether we choose to normalize prices or not, it is easy to find the profits that firms make using our notation. Consider any production plan for firm f : $y_f \in Y_f$. Recall that our convention is that this is net-output vector and so inputs used up by the firm in production are negative, while outputs produced by the firm are positive. This immediately means that we can calculate **profits** as follows:

$$\pi_f = p y_f$$

since any positive net output, when multiplied by its price, is a positive number and so becomes revenue, while any negative net output (an input), when multiplied by its price is a negative number, and so becomes a cost that is deducted from revenue. What is left after costs are subtracted from revenues is, of course, profit.

Given all this, we can at last define a competitive equilibrium.

Competitive Equilibrium:: A **competitive equilibrium** for endowments $(\omega_1, \dots, \omega_I) \equiv \omega$ and ownership shares $\theta \in \Theta$ consists of:

- (a) an allocation $(x_1, \dots, x_I) \equiv x \in X$
- (b) a production plan $(y_1, \dots, y_F) \equiv y \in Y$
- (c) a price vector $(p_1, \dots, p_N) \equiv p \in \mathbb{R}_+^N$

such that:

- (1) $\forall i \in \mathcal{I}, \quad px_i \leq p\omega_i + \sum_{f \in \mathcal{F}} \theta_{f,i} py_f$
- (2) $\forall \hat{x}_i \in X_i$ such that $p\hat{x}_i \leq p\omega_i + \sum_{f \in \mathcal{F}} \theta_{f,i} py_f$, it holds that $u_i(x_i) \geq u_i(\hat{x}_i)$.
- (3) $\forall f \in \mathcal{F}$ and $\forall \hat{y}_f \in Y_f$ it holds that $py_f \geq p\hat{y}_f$
- (4) $\sum_{i \in \mathcal{I}} x_i = \sum_{i \in \mathcal{I}} \omega_i + \sum_{f \in \mathcal{F}} y_f$

Condition (1) says that agents can afford their equilibrium allocation given their total income from endowments and profit shares.

Condition (2) says that no affordable alternative goods bundles is strictly preferred to an agent's equilibrium allocation.

Condition (3) says that firms maximize profits.

Condition (4) says that total consumption by agents equals net production by firms plus the total endowment. In other words, demand equals supply. This is sometimes called the material *balance condition*.

Glossary

Arrow-Debreu-McKenzie General Equilibrium Economy with Production: An abstract description of an economy with a fixed (but arbitrarily large) number of consumers, firms, and goods. Each agent is described by his endowment of goods, his ownership shares of the firms, and his preference relation, while the firms are described by their production sets.

Autarchy: Autarchy is “the government of one’s self”. In economic contexts, it often refers to a situation when an agent or a country is unable to trade, and so must consume only the things he produces himself (or that the country produces domestically).

Blocking Allocation: An $\bar{x}$ allocation which is feasible for some coalition $S \subseteq \mathcal{I}$ blocks an allocation x , which is feasible for the grand coalition if all agents $i \in S$ are better off consuming $\hat{x}_i$ than x_i .

Competitive Equilibrium: A triple consisting of a set of prices for each good, an allocation of each good for each agent, and a net production plan of each good for each firm collectively satisfying equilibrium conditions given a set of initial endowments of goods and ownership shares of firms over agents. First, that all agents find that their allocation in the competitive equilibrium is utility maximizing given their budget sets. Second, that all firms find that their production plans at the competitive equilibrium are profit maximizing. Third, that the sum of endowments over agents minus the sum of the consumption levels allocated to agents at the competitive equilibrium equals the sum of net production plans summed over firms. More briefly, agents maximize utility, firms maximize profit, and when they do, the resulting allocation is feasible.

Consumer Price: The net money price paid by a consumer to all agents (including the government and the producer) for a unit of a good. This includes all taxes and subsidies.

Consumption Plan: An allocation of all the goods over all the agents in an economy $(x_1, \dots, x_I) \equiv x \in X \subset \mathbb{R}^{N \times I}$.

Contract Curve: The set of weakly Pareto optimal allocations in pure exchange economy. This is called the contract curve since under equilibrium prices, agents will agree to contracts that take them from their endowments to some place on the contract curve.

Core: A feasible allocation $\bar{x}$ is a **core allocation** if it cannot be blocked.

Core Convergence Theorem: As an economy gets large in the sense of adding more firms and agents, the set of core allocations shrinks until it contains only the CE allocations.

Decentralization: The amazing fact that the activity of all agents in a large price taking economy can be coordinated and will lead to Pareto optimal outcomes with only a single price for each good. No central command and control is required, and choices are entirely decentralized.

Edgeworth Box: A graphical representation of a two person, two good pure exchange economy without production.

Excise Tax: A tax that a producer is legally obliged to pay to government when he sells a good.

Experience Good: A good whose consumption value cannot be determined until it is experienced or consumed. For example, one cannot know how much one will enjoy a meal at a new restaurant until one tries it, and one cannot look at speaker and know how good it will sound. It is costly to obtain full information about the desirability, and thus, one willingness to pay for such goods, and is detracts from market efficiency.

Featherbedding: The practice on the part of managers or supervisors of firms of claiming that production takes more resources than it really does and then use the surplus inputs in other ways. They might share these with their employees, not make them work as hard, sell the surplus, or just manage production poorly. Unless a monitoring mechanism in place that makes this practice difficult to get away with, feather-bedding will arise and detract for full cost minimization.

First Welfare Theorem: All competitive equilibria are Pareto optimal.

Frictions: Market frictions impede and slow transactions. More generally, anything that drives a wedge between the net value a buyer gives, and a seller receives, from a transaction. This includes taxes, brokerage and transaction fees, non-monetary costs involving the time and effort to search, gathering information about the goods to be bought and sold, the cost of delays, direct costs involved in conveying ownership and so on. All transactions costs are frictions. Incomplete and asymmetric information may generate frictions, but are not frictions in themselves.

General Equilibrium Analysis: General equilibrium approaches economic problems by modeling a large number of consumers and producer interacting in a large number of markets in a completely self-interested and non-strategic way. Prices are taken as given by all agents and are the only things that coordinates their behavior.

Incomplete Information: When agents do not know facts about the nature of the markets, information is said to be incomplete. For example, this includes not knowing about the quality, type, location, prices, or number of goods being offered, the type of buyers or sellers they are likely to encounter, and so on. This differs from imperfect information which is related to imperfect knowledge of the actual actions of other agents in the economy.

Incomplete Markets: Even though there may be N goods in an economy, it might be that only $M < N$ of them are bought and sold. This could be because the government forbids certain goods from being transacted (drugs, human organs), that property rights are incomplete and so title to a good cannot be given fully to another person, or that various kinds of frictions, transac-

tions costs, and imperfect information make it unprofitable for market to open. The result is that there are foregone opportunities to gain from trade and the first welfare theory fails.

Incomplete Property Rights: A property right gives one agent full and exclusive ownership and control over something of value. To be complete, the property right must be fully transferable to other agents. For example, property rights are incomplete when many agents have the right to use something (a fishery, a common pasteurizer) but none of them can exclude the others. In some countries, people build houses on land to which they do not have clear title in hopes that the owners and the government will not evict them. There are many other examples, but the effect is market failure as agents over-utilize the commonly owned resources, or fail to improve or protect the partial owned land.

Individual Rationality: We say that an allocation is **individually rational** (IR) if all agents get at least as much utility as they could if they were to go off by themselves and form a one-person coalition.

Marginal Rate of Transformation (MRT), The slope of the production possibility frontier at any given point. That is, the rate at which the technology permits production of one good to be transferred to another good.

Market Clearing Prices: Mathematically, this means we must find a P^{FM} that satisfies this equation: $S(P^{FM}) = D(P^{FM})$. Such a price is said to be **market clearing** because it leaves neither excess demand nor excess supply.

Nominal Price: The total amount of money that the consumer gives to the producer in exchange for a unit of a good. The consumer may have to pay taxes or receive subsidies in addition to this price. Symmetrically, the total amount of money a producer receives from the consumer for a unit of a good. The producer may have to pay taxes out of what he receives or may be eligible to get subsidies in addition to the nominal price.

Numéraire Good: When the **normalization** $p_1=1$ is imposed we call good 1 the **numéraire good**. This normalization is generally used when preferences are quasilinear in the first good: $u_i(x_i) = x_{i,1} + v_i(x_{i,2}, \dots, x_{i,N})$. In this case, good 1 is often called a **transferable good**.

Ownership Shares: Firms are owned by agents. We use the following notation to express this: θ_{fi} is the fraction of firm f owned by agent i . Thus, for any particular firm, the sum of ownership shares over agents must equal one and no share can be negative. More formally, the whole vector of ownership shares is denoted $\theta \in \Theta$ where:

$$\Theta \equiv \{(\theta_{1,1}, \dots, \theta_{1,I}, \dots, \theta_{f,i}, \dots, \theta_{F,1}, \dots, \theta_{F,I}) \in \mathbb{R}_+^{F \cdot I} \mid \forall f \in \mathcal{F}, \sum_{i \in \mathcal{I}} \theta_{f,i} = 1\}.$$

Private Information: The things one agent knows, and others do not are called private information. Note that when the information available to all agents is not the same, we say that information is asymmetric. When not all information is known to all agents, we say information is in-

complete. Private information implies that information is incomplete, and asymmetric. but incomplete or asymmetric information does not imply that information is private.

Producer Price: The net money price received by the producer from all agents (including the government and the consumer) for a unit of a good. This includes all taxes and subsidies.

Production Plan: A vector that specifies the net output of each good for each firm in the economy.

Production Possibility Frontier (PPF): The upper boundary of a production set is called the production possibility frontier.

Production Set: The set all net output vectors that are feasible for a firm given its technology is a firm's production set. The set of all feasible production levels for an economy once the social endowment of goods is added to the sum or the production sets of the firms in the economy is the economy's aggregate production set.

Pure Exchange Economy: An economy with no firms and no production. Agents simply trade their initial endowments with one another. The Edgeworth box is a graphical representation of the two person, two good case.

Region Of Mutual Improvement: The football-shaped area between the indifference curves through each agent's endowment point in an Edgeworth box. This region includes all the strong and weak Pareto improvements over the individual endowments.

Sales Tax: A tax that a consumer is legally obliged to pay to government when he purchases a good.

Search Costs: The opportunity cost in time, effort, money and other resources spent finding information about commodities to be bought or sold. This includes such things as price, location, and quality information about the commodities, and also information about the buyers and sellers to the extent that it affects the value that agents get from the transaction.

Second Welfare Theorem: Any Pareto optimal allocation can be supported as a competitive equilibrium allocation for some reallocation of endowments.

Simplex: The (N-1)-dimensional unit simplex is the set of N-dimensional vectors such that all the components are non-negative and sum up to one:

$$\Delta^{N-1} \equiv \{p \in \mathbb{R}_+^N \mid \sum_{n \in \mathcal{N}} p_n = 1\}.$$

The "(N-1)-simplex" is a manifold of one-dimension less than the space in which it is embedded.

Shortage: If the market price of a good is such that the quantity demanded exceeds the quantity supplied, then there is excess demand, or equivalently, a shortage.

Strong Pareto Improvement: A movement from an initial allocation that leaves all agents strictly better off.

Strong Pareto Optimality (SPO): A feasible allocation x is strongly Pareto optimal if there does not exist another feasible allocation $\bar{x}$ such that $\bar{x}$ is a weak Pareto improvement over x .

Surplus: If the market price of a good is such that the quantity supplied exceeds the quantity demanded then there is excess supply, or equivalently, a surplus.

The Core: A feasible allocation $\bar{x}$ is a **core allocation** if it cannot be blocked.

Thin Market: A market with a small numbers of potential demanders and potential suppliers. That is, both sides the market have a small number of agents. For example, only a few people might want to buy a certain object on eBay, or a certain type of house in a given city, while only a few sellers may offer the object for sale on eBay, or try to sell this type of house in the city.

Transactions Cost: The wedge between the net value a buyer gives, and a seller receives from a transaction. See frictions.

Weak Pareto Improvement: A movement from an initial allocation that harms no agent, but benefits at least one agent.

Weak Pareto Optimality (WPO): A feasible allocation x is weakly Pareto optimal if there does not exist another feasible allocation $\bar{x}$ such that $\bar{x}$ is a strong Pareto improvement over x .

Problems

1. Explain how the following things would change the equilibrium price and quantity of electric vehicles, and why.
 - a. The price of crude oil goes down.
 - b. It is expected that new battery technology will be invented within a couple of years that will make EVs better and cheaper.
 - c. The current consumer subsidy of \$7500 per car is cut in half.
 - d. Cities invest heavily in public transportation systems.
2. Suppose that there is a \$1 per pack sales tax on cigarettes paid by consumers. At the same time, there is a subsidy on the production of tobacco that amounts to \$1 per pack. In a diagram, show the following things: the price and quantity of cigarettes consumed *before* either the tax or subsidy is applied to the market, and the consumer's price, producer's price, nominal price, and quantity of cigarettes *after* both the tax and subsidy are placed on the market.
3. Let the supply curve for solar panels be given by the following equation: $Q = 3p$, and the demand curve be given by: $Q_d = 1000 - 2p$.
 - a. What is the equilibrium price and quantity?
 - b. Suppose the government decided to subsidize producers at the rate of \$50 per unit. What is the consumer price, producer price and quantity after the subsidy.
 - c. What is the elasticity of demand for solar panels in both the free market and subsidy equilibrium? (Your answers should be numbers.)
4. There a dangerous neurological condition caused by paying too much attention to the Dow-Jones Industrial Average loose in the world called Micro-Fever (MF). It causes its unfortunate victims to enroll in intermediate microeconomics course and in more serious cases, become economics majors. This is very worrying, because it is found that drinking beer tends to make people even more likely to develop this condition and decide they like economics. Fortunately, a treatment called the "green pill" has just been developed. The demand for this pill is:

$$D_g(p_g, p_b, w) = 4w - 100p_b - p_g,$$

while the supply is given by:

$$S(p_g) = 4p_g$$

where p_g and p_b are the price of the pill and beer, respectively, and w is the average income. Suppose average income is \$250 and the price of beer is \$1.

- a. What is the free market quantity and price of the MF blue pill?
- b. The government is understandably concerned with the effect that economics might have on America's youth (think how it might affect their voting in the future!) It decides to subsidize

the manufacture of the MF blue pill at \$10 each. What is the new consumer price, producer price, and quantity?

- c. What is the income elasticity of demand at the subsidized equilibrium quantity? What is the own price supply elasticity? (Your answers should be numbers).

5. Ken McSubstitute and Ron O'Complement were flying to a fast food festival in Fiji when an unexpected storm forced their plane to ditch in the middle of the Pacific. Miraculously, they are washed up on a desert island. Ken finds that he has only 5 slightly wet hamburgers and 15 orders of fries in his pockets. Ron discovers he has 15 hamburgers and 5 orders of fries. Ken only cares about how much he gets to eat. His utility function is:

$$u_K(x_h, x_f) = x_h + x_f.$$

Ron, on the other hand, believes that it is uncivilized to eat hamburgers without French fries or French fries without hamburgers. His utility function is:

$$u_R(x_h, x_f) = \min \{x_h, x_f\}.$$

In an Edgeworth box, show the endowment point, the Pareto optimal allocation(s), and the competitive equilibrium. Is the competitive equilibrium Pareto optimal?

6. Suppose that an economy is at a competitive equilibrium.
- By the First Welfare Theorem, this allocation is Pareto optimal. Suppose we took a poll and asked each citizen if he or she is happy with the current allocation, or would prefer a different one. What percentage would want to stay and what percentage would want to move to a new allocation?
 - Suppose that technological advances are made that make it possible produce enough to make all agents better off. The market finds a new equilibrium. Is it possible that any agent would prefer the old equilibrium to the new one?
 - Suppose there were market failures that prevented the market from reaching equilibrium. A social planner imposes an allocation on the economy without using the market. Fortunately, the social welfare, and even the consumer surplus, is larger at this allocation than it was previously when markets failed. Is it possible that any agent would prefer the market failure allocation to the social planner's allocation?
7. Wayne Newton is the only survivor of a tragic cruise ship accident in the South Pacific. He is washed up on the shore of an isolated island where the only foods available are coconuts and fish.
- Wayne can allocate his time to fishing or coconut gathering. Putting coconuts on the X-axis, draw a convex production possibility frontier between fish and coconuts, and show the optimal consumption and production choices that Wayne Newton makes.
 - Suppose that a series of tragic cruise ship accidents leaves the islands surrounding Wayne's island full of stranded lounge singers and that they set up an "interisland market" to trade fish and coconuts with one another. Suppose also that in general, fish are more abundant on the other islands, and so the relative price of fish at the interisland market is lower than the

domestic price of fish on Wayne's island before trade was opened. What happens the production and consumption of fish on Wayne's island compared to part (a) if fish are a normal good?

Chapter 7. Welfare Economics and Social Optimum

Section 7.1. What is Welfare Economics?

So far, we have explored how self-interested agents behave as individuals or in groups. We have not asked whether these the outcomes are in any sense desirable or socially optimal. We have kept to what is called the *positive* and away from the *normative* side of economics. Recall:

Positive Economics: Statements about what is. No opinions are offered, just facts and analysis that follow from assumptions.

Normative Economics: Statements about what should be. This involves value judgments based on political, religious, philosophical and ethical beliefs.

As a behavioral science, economics is focused on positive analysis. Economics tells us absolutely nothing about how to make normative judgments. Nevertheless, economists are often called upon to suggest a “good policy” how to “improve” an economic outcome. Although economists are no more qualified than any other citizen to say what is good or desirable for a society, we do have tools available that can be used to make these choices clearer. This section develops some of these tools.

Section 7.2. Consumer and Producer Surplus

Suppose that my individual demand curve for oranges is given by the following equation:

$$D(P) = 10 - P.$$

This tells us, for example, that if the price of oranges is \$6, I will choose to consume 4 oranges. More precisely, it tells us that if I have already purchased and consumed 3 oranges at \$6, I would be willing to buy exactly one more at this price, but not two more. Why is this? Recall, that the demand curve shows the optimal consumption level of a good after I consider all alternative uses of my income under the current prices. In other words, I am willing to pay \$6 for my fourth orange *exactly because* the incremental benefit I get from eating that fourth orange is worth precisely the benefit I would get from spending \$6, on the next most beneficial use (that is, the opportunity cost of \$6).

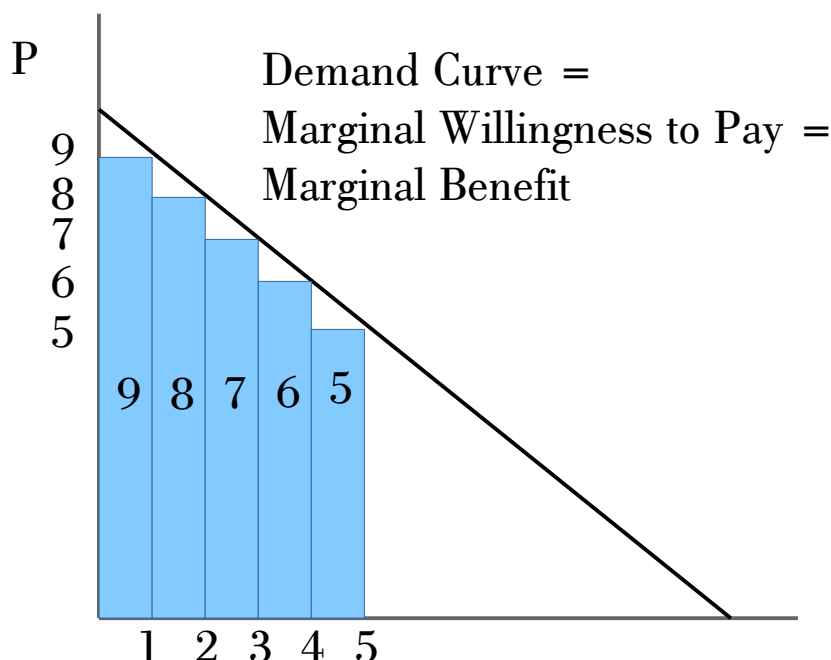


Figure 81: Approximating marginal and total benefit of consumption

To be more concise: My **reservation price** for the fourth orange is \$6, because this is my marginal benefit of the fourth orange. *Thus, my demand curve is really the same as my **marginal benefit** curve.* I can therefore use this curve to find out how much benefit agents get from their consumption choices.

As you can see in the figure, the MB to me of the first orange is \$9. This means that I am indifferent between one orange and \$9. They are precisely as good as one another to me. The MB of the second orange (given that I have already eaten the first) is \$8. Thus, the total value of a basket of 2 oranges is \$17. In a similar way, a basket of 4 oranges has a total value of \$30.

Of course, this is really just an approximation. The real value is the area under the demand curve. This is because, just as the derivative of total benefit with respect to quantity is the marginal benefit, the integral (geometrically, the area beneath the curve) of the marginal benefit curve with respect to quantity is the total benefit.

$$\int_Q MB = TB, \text{ and } \frac{\partial TB}{\partial Q} = MB .$$

Thus, the total benefit equals the triangle plus the rectangle, $\frac{1}{2}(4 \times 4) + 6 \times 4 = 8 + 24 = 32$,

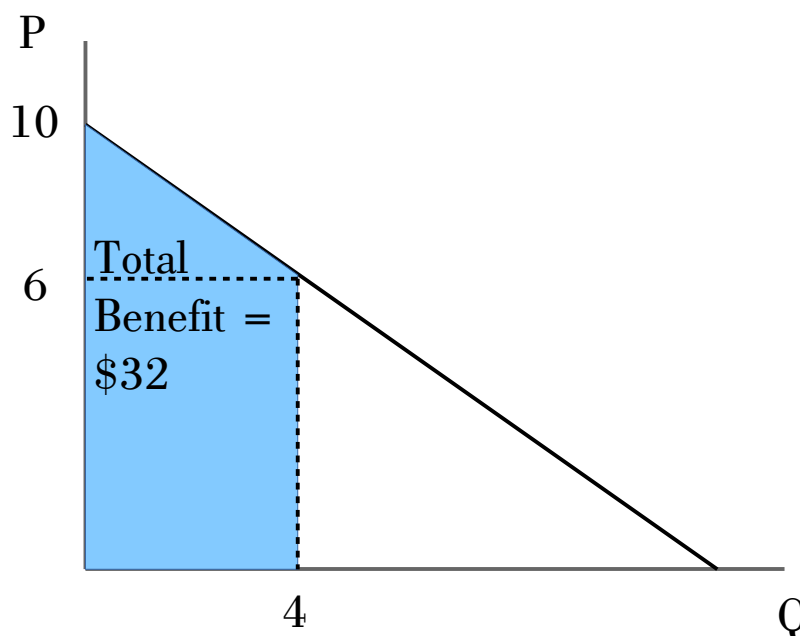


Figure 82: Total benefit of consumption

Unfortunately, oranges are not free. Suppose the market price of oranges is \$6, then I have to pay a total of \$24, to purchase 4 of them. This gives me a net benefit of $\$32 - 24 = 8$. We call this net benefit an agent's **consumer surplus (CS)**. One way to think about this is that it is worth \$8, to me to be given access to a market in which I am allowed to buy as many oranges as I like at \$6, each. In other words, I would pay as much as \$8 for a map to such a market.

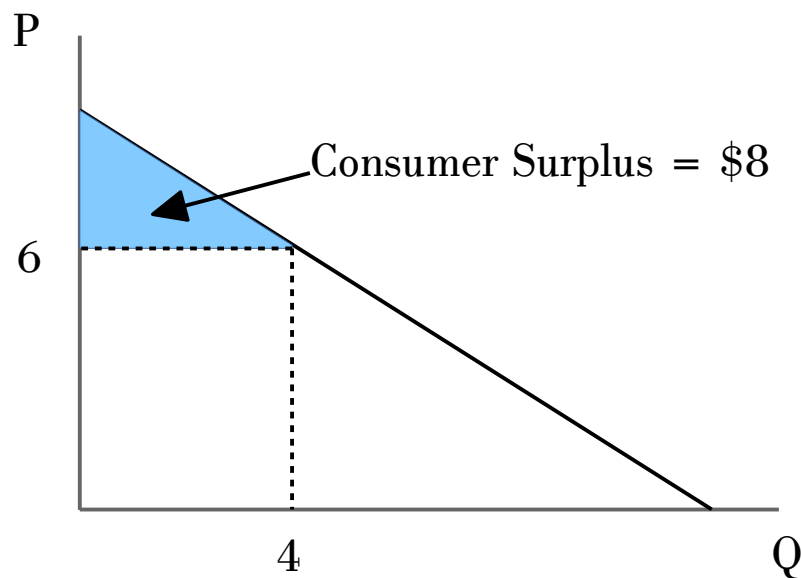


Figure 83: Consumer surplus for an agent

Now, consider a market demand curve. Recall that this is the horizontal sum of individual demand curves. This means that when looking at the market demand curve, each column of surplus really has an agent attached to it. Thus, Sue is the first one to be willing to buy an orange as the price finally drops to, \$9, She has a higher marginal benefit of orange consumption than anyone else. When the price drops to \$8, Bob is willing to buy an orange, and when it drops to \$7, Sue finds she is willing to buy a second orange. Finally, as the price drops to \$6, Joe jumps in and finds that oranges are finally at or below his reservation price (MB) for the first time.

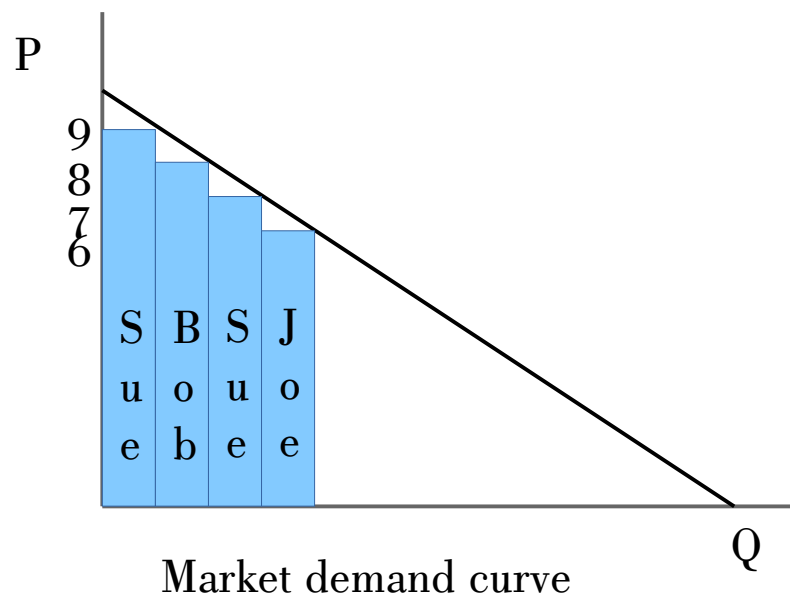


Figure 84: Consumer plus in that market

This is even easier from the producer side. The total benefit producers get is the total revenue, $(P^{FM} \times Q^{FM})$. The total cost is the integral of the marginal cost, that is, the area beneath the marginal cost curve. The difference between TR and TC is the **producer surplus (PS)**. In the example, $(P^{FM} \times Q^{FM}) - TC = PS$. Note that similar to consumers,

$$\int_0^Q MC = TC, \text{ and } \frac{\partial TC}{\partial Q} = MC.$$

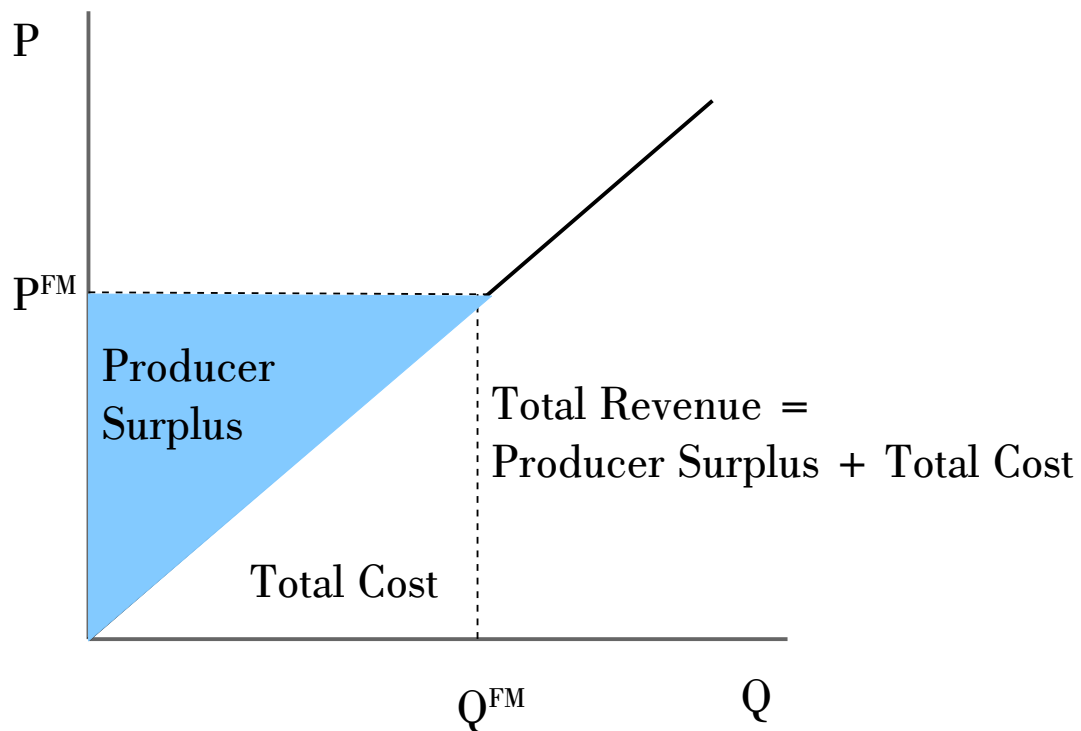


Figure 85: Producer surplus and total cost in the market

As an aside, the producer surplus need not be positive as we show in the figure. Recall that the MC is “U-shaped” in general. Thus, in the special case of a competitive market with free entry, first part of the MC will be above the free market price and so area between the MC and price are therefore losses. The second part of the MC Will be below the price and are therefore gains. In equilibrium, these will be exactly off-setting and so $PS = 0$.

In any event, if a market is competitive with both buyers and seller being price takers, we get the following:

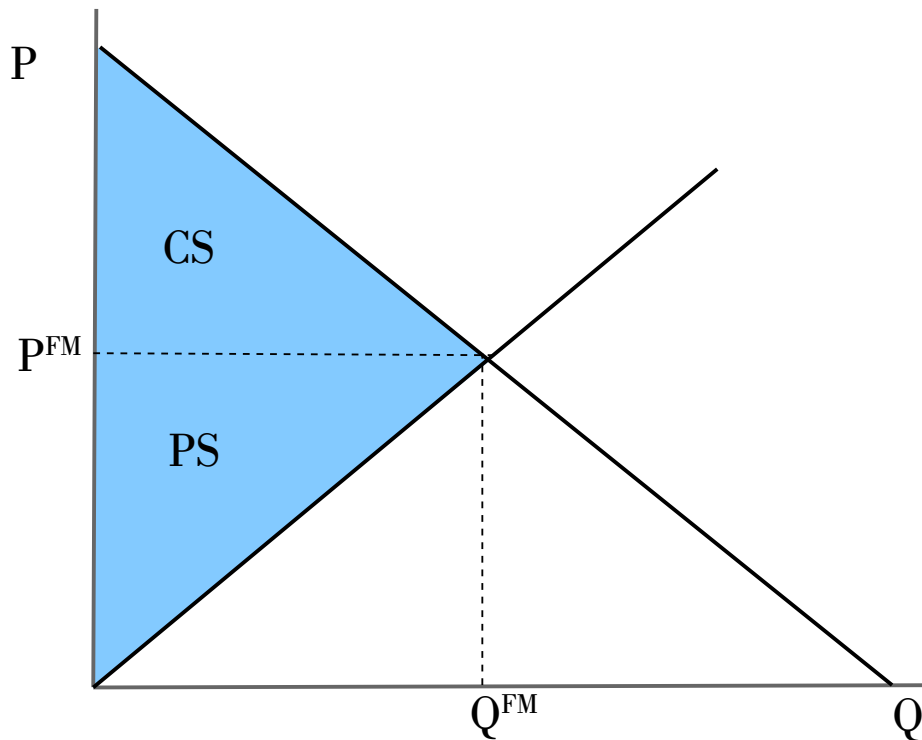


Figure 86: Consumer and producer surplus in the market

The sum of the net surplus to all agents in the economy is called the **Social Surplus (SS)**. (We will use *social gain* and *social welfare* interchangeably with social surplus.) Thus:

$$CS + PS = SS$$

In general, the rule is:

Consumer Surplus: The area beneath the demand curve, out to the quantity, and above the price.

Producer Surplus: The area beneath the price, out to the quantity, and above the supply curve.

Section 7.3. Social Optimum and Positive Economics

It is worth discussing the concept of social surplus for minute. We defined this as is the net benefit measured in dollars that all agents (producers, consumers, the government, and everyone else) realize in a given market. One might think that the objective of government should be to maximize social gain. This may be true in many cases, but some caution is needed.

The notion of consumer surplus, in particular, incorporates the willingness of agents to pay for goods (that is, the reservation price). Thus, Jeff Bezos may be willing to pay \$100,000,000 for the *Venus de Milo* with the idea that he will gild it and make it into a fountain to go in his master bathroom. It might be that the richest museum (acting to express the preferences of its patrons) would only be willing to pay \$20,000,000 to display it. Is it a social good that Bezos gets this piece of our culture's heritage? Thousands of people much poorer than Bezos would have their lives enriched if they could see the *Venus de Milo* in the museum. If we had majority voting, the *Venus de Milo* would certainly be in a museum, for example. A justification for thinking about social surplus as a relevant measure of the value of social policy is the so called:

Kaldor-Hicks Criterion: This says that a policy or allocation is socially beneficial if there is a potential for the winners to fully compensate the losers and yet still be better off themselves. In other words, a policy is socially good if it allows for a **potential Pareto improvement** (more on Pareto improvement below.)

In our hypothetical case, Bezos could take the *Venus de Milo* while compensating the museums \$21M (which leaves museums and their patrons better off than if they had gotten the statue), and still be better off himself. If net social gain is positive, then the *potential* exists to reallocate the surplus in such a way that all agents benefit. If such compensation were actually paid, then there would be unanimous agreement that the policy is good, and we would be able to endorse it from a strictly positive standpoint.

Of course, just because all agents *could* gain if compensation were paid after the market concludes does not mean that agents actually *will* be compensated in real life. A market or governmental policy will usually leave winners and losers. In this case, positive economics provides no criterion to evaluate the desirability of the policy. Some people like it, and others do not. Positive economics provides no guide for choosing who should win and who should lose, or weighing such losses and gains against one another.

Despite this, one possible motivation for maximizing social surplus is that the gains and losses to various policies might be randomly distributed over members of the society. In this case, all agents will come out ahead on the average over time, even though any particular policy might be to their detriment. On the other hand, if the system is rigged such that new policies systematically benefit some favored group at the expense of the rest of society, then maximizing social surplus is equivalent to making non-Pareto improving transfers, and so cannot be justified on a positive basis.

Section 7.4. Policy Analysis

Subsection 7.4.1. Sales Tax

All those caveats and disclaimers having been said, we can use the tools outlined above to analyze the impact of various governmental and other policies. The figure below shows the simple case of a sales tax. We see that consumers and producers are made worse off, but that the government is better off since it gets the tax revenue. However, the losses to the losers are larger than the gains to the gainers, so total social welfare goes down.

If one policy (a sales tax, for example) has a smaller social gain than another (the free market for example), the difference between the social two gains is called the **Dead Weight Loss (DWL)**. This means that the policy that maximizes social gain has zero DWL by definition and so is best by the Kaldor-Hicks Criterion.

In the next figure, you can see that the sales tax of T causes a DWL of equal to the area labeled HI relative to free market. It also reduces the quantity supplied and demanded from the free market level, raises the price consumers pay, lowers the price producers receive, reduces both consumer and producer surplus, but gives the government **Tax Revenue (TR)**, which is part of the social surplus. You can easily verify that an excise tax of T would cause exactly the same DWL and distribution of CS and PS and TR.

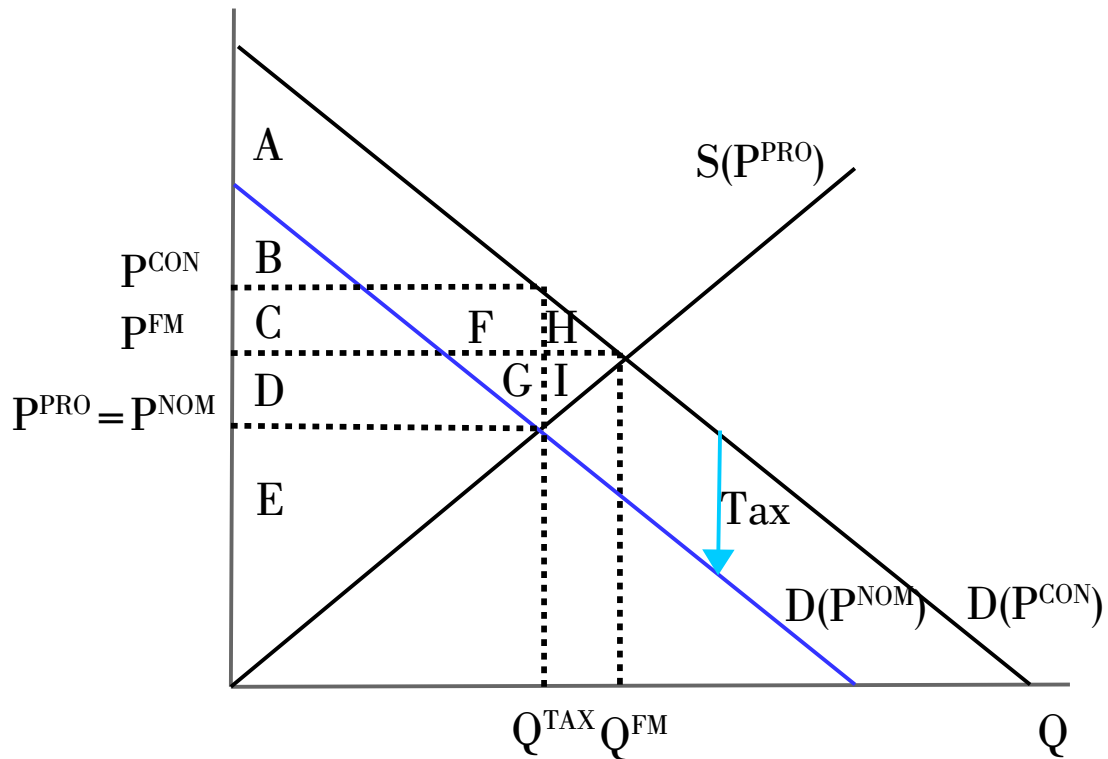


Figure 87: Graphical analysis of a sales tax

	FM	Sales Tax
CS	ABCFH	AB (or BCD)
PS	DEGI	E
TR	-	CDFG (or AFG)
SS	ABCDEFGHI	ABCDEFG
DWL	-	HI

Table 2: Welfare analysis of a sales tax

Note the following:

- The areas are labeled with letters to help us do the accounting. The letters really stand for dollar amounts of surplus that are divided over agents. For example, areas A plus B form a triangle. If we knew the base and the height, we could calculate its numerical value. If $A + B = \$13,490$, for example, whichever set of agents count A and B as part of their surplus are getting \$13,490 of value.
- There are two demand curves shown. The $D(P^{\text{CON}})$, shows the quantity demanded as a function of the **consumer price**, while $D(P^{\text{NOM}})$, shows the quantity demanded as a function of the **nominal price**. These are separated by a vertical distance exactly equal to the per unit sales tax. Suppose, for example, that consumer wished to buy 100 units when the price is \$40 and there is no tax. If the government imposes a sales tax of \$10 per unit, then consumers will wish to buy exactly 100 units at a nominal price of \$30 since in this case, the consumer price again equals \$40 which is the nominal price plus the sales tax.
- We can calculate consumer surplus as the area beneath $D(P^{\text{CON}})$, out to the Q^{TAX} , and above $D(P^{\text{CON}})$, (that is, AB), or the area beneath $D(P^{\text{NOM}})$, out to the Q^{TAX} , and above P^{NOM} (that is ACD). You can easily verify that these two triangles have identical area. The first triangle is just T units above the second.

Subsection 7.4.2. Subsidy

Subsidies are equivalent to negative taxes. When an agent buys or sells a good, the government gives the agent an amount of money equal to the subsidy instead of demanding that the agent pay an amount of money equal to the tax. Real world examples include, solar tax credits, and educational scholarships, crop payments to farmers, and research and exploration credits to companies. The figure below shows an example of a consumer subsidy.

You can see that consumers and producers are better off, but that the government is worse off since it pays the **subsidy cost (SC)**, which must be subtracted from social surplus. However, the losses to the losers are larger than the gains to gainers, so total social welfare goes down and there is a DWL of HI. The subsidy also increases the quantity supplied and demanded from the free market level, lowers the price consumers pay, and raises the price producers receive. Again, you can easily verify that a producer subsidy of S would cause exactly the same DWL and distribution of CS, PS, and SC. Also note that there are two equivalent ways of expressing the producer surplus: $DEL = BCDFG$ and the subsidy cost: $BCFGHIJK = EHIJKL$. That is, the dollar amounts these two pairs of geometric areas represent are the same.

The subsidy shown in the figure below is a given to producers rather than consumers. As a result there are two supply curves shown. The $S(P^{PRO})$ shows the quantity supplied as a function of the **producer price**, while $S(P^{NOM})$ shows the quantity supplied as a function of the **nominal price**. These are separated by a vertical distance exactly equal to the per-unit producer subsidy.

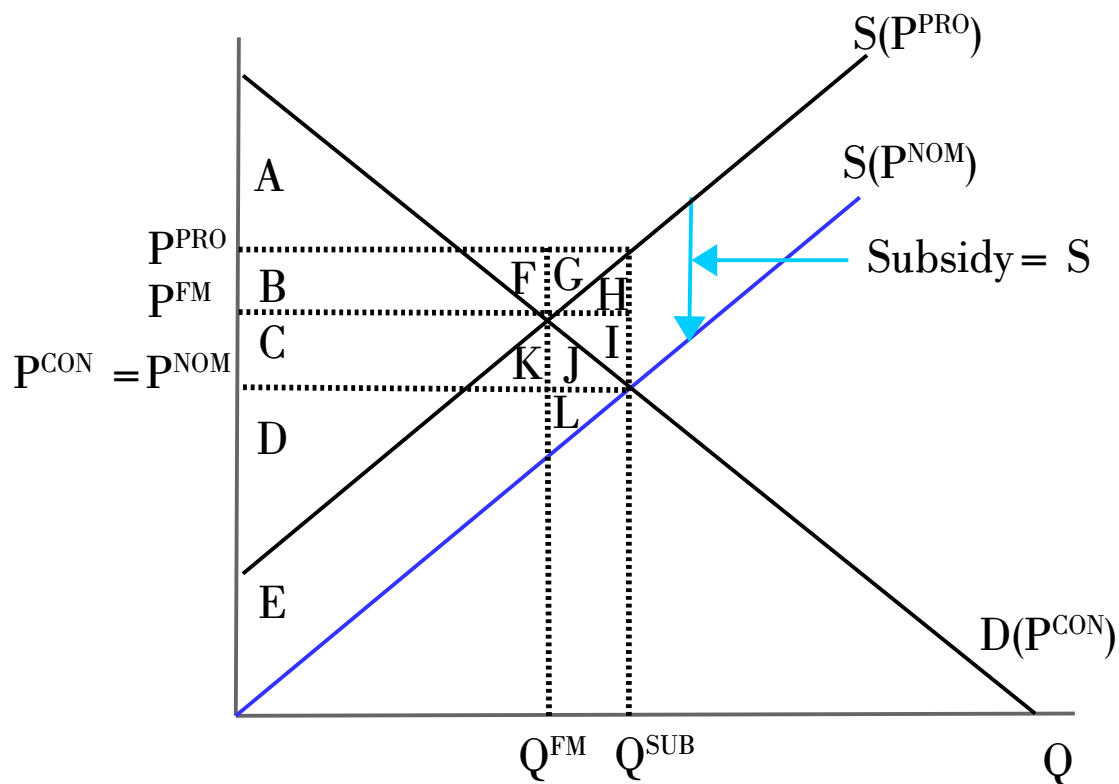


Figure 88: Graphical analysis of a producer subsidy

	Free Market	Subsidy
CS	AB	ABCKJ
PS	CD	DEL (OR BCDFG)
SC (subtracted from SS)	-	BCFGHIJK (or EHIJKL)
SS	ABCD	ABCD – HI
DWL	-	HI

Table 3: Welfare analysis of a producer subsidy

Subsection 7.4.3. Price Ceiling with Shortage

A **price ceiling** is a maximum price that firms are allowed to charge by government mandate. If price ceiling is set below the free market price (as it generally is), then consumers wish to buy more goods at the artificially low price than firms are willing to sell. This means that price ceilings generate excess demand, also called **shortages**. If you see people standing in line to buy a good, you know it must be priced below the free market level.

Examples of price ceilings include rent control and legally mandated maximum prices for basic goods, especially in socialist countries and during wartime. Tickets to popular sporting events (the Super Bowl), and shows (Miley Cyrus concerts) are also sometimes priced too low, although not by government mandate.

It matters how the shortage that results is resolved. Here are four major possibilities:

Very Smart Government (VSG): The government allows only the agents who are on the highest part of the demand curve (and so have the highest MB of consumption) to buy the goods. This means the government has to be magically smart enough to know everyone's reservation price.

Standing In Line (SIL): A policy thought experiment where shortages of goods are resolved by selling them on a first come, first served, basis. This induces rent seeking behaviors like search, bribery, or literal standing in line in order to be allowed to buy the good that is in short supply.

Lottery with Resale (L with RS): A policy thought experiment where a lottery is held to determine who is allowed to buy a good in short supply at a below market price. Allowing resale means that the winners of the lottery can either sell the winning ticket to other agents who value the good more than they do (and therefore the right to buy at a below market price), or sell the good directly at whatever price the market will bear. Note that a lottery can also be used to result a surplus by only allowing the winners to sell goods at the artificially enforced above market price.

Lottery Without Resale (L without RS): A policy thought experiment where a lottery is held to determine who is allowed to buy a good in short supply at a below market price. Disallowing resale means that the winners of the lottery must consume any goods that they choose to buy and cannot sell them to agents who value the good more than they do.

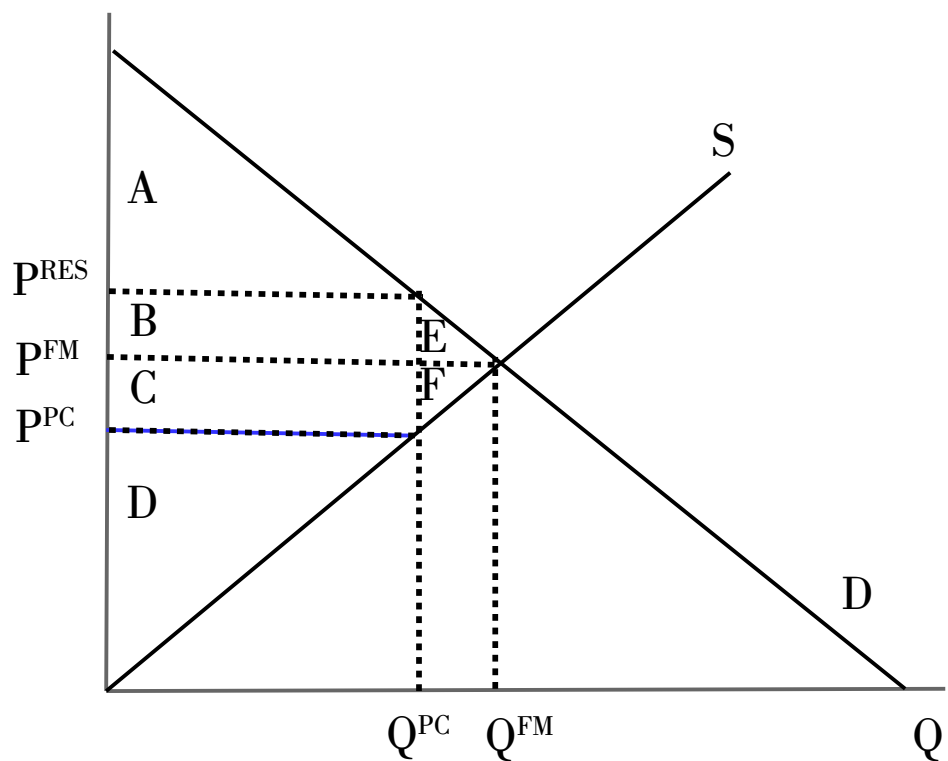


Figure 89: Graphical analysis of a price ceiling with shortages

	FM	VSG	SIN	L with RS	L without RS
CS	ABE	ABC	A	A + BC	<u>CS</u> < ABC
PS	CDF	D	D	D	D
SS	ABCDEF	ABCD	AD	ABCD	<u>SS</u> < A BCD
DWL	-	EF	EF + BC	EF	<u>DWL</u> > EF

Table 4: Welfare analysis of a price ceiling with shortages

Explanation

VSG: This is a completely unrealistic base case. Obviously, there is no such thing as a very smart government. However, if the government could somehow figure out who had reservation prices above P^{RES} and allow only those agents to buy the good at price P^{PC} then the CS would be ABC while, the PS would be D Note that this results in a DWL of EF This is because the price ceiling causes firms to stop producing while MB is still greater than MC. Thus, price ceilings cause society to forego the opportunity to produce goods that would have generated more benefits than costs.

SIL: In addition to literally standing in line, SIL includes many other kinds of costly effort aimed being allowed to buy the good below such as spending time and resources searching for the good and engaging in costly behavior that convinces the government or owner of the good to sell to you instead of someone else.

We call this kind of activity by agents **rent seeking** in economics. If there is a profit of some kind to be made, then agents do whatever they need to get it for themselves. The “rents” here might be surplus obtained from getting plum contracts, jobs paying above market wages the government awards, or the right to buy goods for less than they are worth, as in the case of a price ceiling. The problem is that agents are willing to put as much effort into getting this rent for themselves as the value of the rent. They pay bribes, lobby, fight, search, stand in line, or whatever else is required. If it is possible to get the rent for less effort than this, then some other agent will put in the little extra effort needed to win. In the end, the agent who wins walks away with no net rent at all since it is all dissipated by his rent seeking efforts.

To understand this more clearly, consider a simple example: Suppose everyone in the world makes \$20 per hour. Miley Cyrus books a hall with 1000 seats and puts tickets on sale for $P^{\text{PC}} = \$60$ However, $P^{\text{RES}} = \$100$ which means that she could have priced them at \$100 and still sold out. *I claim that the line for tickets will have to be exactly two hours long.*

The logic is this: If the line is two hours long then anyone who wants a ticket has to give up two hours of time as well as \$60. But if the wage rate is \$20 per hour, then the value of the time spent in line is \$40, making the opportunity cost of a ticket \$100 in total. At this net cost, exactly $Q^{\text{PC}} = 1000$ people find that it is worth their while to stand in line for tickets since they get a MB of \$100 or more from seeing Miley.

If the line were longer than this, fewer than 1000 would show up and tickets could be purchased at face value since some would be unsold. More likely, people would show up later at the ticket office instead of getting there early to be sure they were one of the first 1000. If the line were shorter than this, more than 1000 people would show up or some people would show up earlier and spend more time waiting.

The point is that two hours of wasted time is the exact amount necessary to make the market clear. The length of the line will automatically adjust to this through the decentralized, rational, private actions of utility maximizing Miley Cyrus fans.

From this we see that even though the money cost of tickets to consumers is P^{PC} , the opportunity cost is P^{RES} . This means that those who stand in line and go to the concert only get a CS of A. However, since Miley does not get any direct benefit from her fans standing in line, this effort is simply wasted and becomes part of DWL. The supplier still only gets a PS of D when goods sell at P^{PC} . Thus, we now see DWL coming from two sources. First, EF is lost due to under-production of tickets. Second, AD is wasted in allocating tickets only to the most avid fans.

L with RS: To understand this, let's consider an example similar to the one above. Suppose again that at price P^{PC} , there are $Q^{PC} = 1000$ units of good supplied, and so 1000 tickets given away by lottery. Note that there are exactly 1000 people on the demand curve with reservations prices of P^{RES} or above (by construction). Call these agents *high demanders* and the rest of the agents *low demanders*. We have no idea how the lottery turns out, but some high demanders will win a ticket. To be concrete, suppose there happen to be 350 high demand winners and so 650 high demand losers. This implies that the remainder of the lottery tickets had to have been won by low demanders. Thus, there must be 650 low demand winners, one for each high demand loser.

From here, the logic is similar to the story above. *I claim that the market clearing resale price of a lottery tickets is $P^{RES} - P^{PC}$.* The proof is easy:

1. Consider the high demand losers. They all have a MB of P^{RES} or above. If the price of a lottery ticket is what I claimed it was, then the opportunity cost of a unit of the good equals the cost of the good at the price ceiling *plus* the cost of the lottery ticket for a total of:

$$P^{PC} + (P^{RES} - P^{PC}) = P^{RES}$$

At this opportunity cost, all 650 high demand losers are willing to buy a ticket since their MB is larger than this opportunity cost.

2. By a similar argument, each of the low demand winners have a MB at or below P^{RES} , and so if they had to pay P^{PC} for a unit of the good, their personal consumer surplus would be less than $P^{RES} - P^{PC}$. Clearly then, all 650 low demand winners are better off selling their lottery tickets and pocketing $P^{RES} - P^{PC}$ than using the lottery ticket themselves.

3. We conclude that at a price of $P^{RES} - P^{PC}$ we find that 650 lottery tickets are both supplied and demanded and so this is a market clearing price for lottery ticket.

Putting this together, the agents who end up actually buying and consuming the good (all the high demanders) pay an opportunity cost of P^{RES} and so get a collective consumer surplus of A. It does not matter if a high demander happens to be a lottery winner or not. The opportunity cost is still P^{RES} since if a high demand winner buys a unit of the good, he must use up his lottery ticket and therefore *foregoes the opportunity to sell it for $P^{RES} - P^{PC}$.*

Lottery winners (both high and low demanders), on the other hand, get a collective CS of BC , which is $Q^{PC} \times (P^{RES} - P^{PC})$. Note that some of these lottery winners will also be high demanders, so they get a double helping of surplus, but the surplus is still coming for two different sources: consuming a good with a MB larger than its opportunity cost, and winning a valuable right to buy a good at a below market price in a lottery.

We conclude that a lottery with resale actually implements the VSG allocation in that only high demanders end up consuming the good. However, the CS is now spread out to low demanders as well if they happen to be lottery winners. We get rid of the DWL due to the allocative inefficiency of the SIN approach, but we still have the DWL from underproduction. This is a general result for lotteries with resale.

L without RS: In this case, lottery winners can buy the good for P^{PC} but cannot sell the right to do so. If by good fortune, only high demanders win the lottery, then the outcome is the same as the lottery with resale. However, to the degree the low demanders win, the CS and SS will be lower, and the DWL higher. Exactly how much deviation from the VSG levels we see depends on exactly how the lottery turns out.

Examples

Now let us consider some real world situations in light of this analysis:

Rent control: Rent control is put into place to make housing affordable in cities like San Francisco and New York. Given what we have learned so far, why would this be a popular idea? To get an apartment at the controlled price, people have to search, stand in line, pay bribes to landlords or employ rental agencies at high cost to “find” them place to live.

In all cases, the effective rent (in this case, “rent” is used in more conventional sense meaning the cost of leasing something rather than a surplus over which agents compete) is even higher than it would be in the free market $P^{RES} > P^{FM}$. Even worse, in the longer run, landlords, who get only $P^{RES} < P^{FM}$, will choose not to maintain or repair their properties, and investors will choose not to build new apartments to replace those that burn down or to match population growth. Landlords will also do their best to convert their rentals to condominiums. Thus, the number of rentals available actually falls, the quality declines, and renters do not see a lower opportunity cost of leasing apartments in the end.

This sounds like a disaster for all concerned. This neglects only one thing. The effects described occur in the long run. In the short run, however, people who are already in apartments get the benefit of a lower rental cost without the need to search! It is only new arrivals to the city and future renters who suffer. Of course, current renters are also the current voters. Future renters have no say in local policy. Thus, rent control can be seen as a law that benefits current renter/voters at the expense of future renters and which lowers the quality and quantity of a

city's rental housing stock in the longer run. There is no free lunch, but rent control allows current voters to eat the lunch of future voters.

Price controls for gasoline: Prices of gasoline go up periodically for a variety of reasons: wars in oil producing countries, oil embargoes, breakdowns in refineries or pipelines, hurricanes or other natural disasters that make its normal distribution possible, to name just a few. Price increases in such crises is often perceived as “gouging” and someone is sure to suggest price controls.

Unfortunately, this leads directly to the effects we outlined above. The gasoline that is available during the crisis under price controls is priced in such a way as to create a shortage. Long lines are the inevitable result. Consumers and producers are unambiguously hurt. In addition, consumers end up driving around searching for open gas stations, and when they find one, most people keep their cars running while they wait to fill-up. Thus, not only is time wasted in searching and waiting, but also gasoline, making the shortage even worse. As always, people with low value of time are more likely to wait in line, but this means people with high value of time do not get as much gas.

Thus, people with high economic value work from home, or do not work at all. This may not be the best way to allocate what gas is available. In addition, producers are likely to try to divert as much gas as they can to regions or countries where prices are not controlled or hold back supplies waiting for the controls to be lifted. This makes the shortage even worse.

These types of price controls make the black market very profitable. Supplies are likely to be diverted to private sales at high prices. Not only are prices higher for these black market purchases, but this further decreases the supply available at the controlled price, which makes the lines, and therefore the opportunity cost even higher. All told, it is difficult to make the case that this kind of policy response is helpful in such a crisis.

Price controls for basic goods: Goods such as bread, potatoes, toilet paper, soap, apartments, and gasoline are often provided by the government at low cost (sometimes even below production cost) in socialist countries like Cuba, Venezuela, and the former Soviet block, in developing countries like Egypt and Mexico, and in oil rich countries in the Middle East.

Sometimes the reason given is philosophical: it is a basic human right to have food, shelter and medical care. In other cases, it is an attempt to prevent unrest in the masses of poor people. It is generally the case that the government does not supply as much as people want at these low prices.

The quality of goods provided is also usually quite low. The poor and unemployed have low opportunity cost of standing in line, and so this can perhaps be seen as a way to make transfers to these groups. The rich can afford to buy better goods on the black market, and also would not find it worth their while to wait hours to buy ten pounds of sad, elderly potatoes for a few kopecks. It also gives the poor and unemployed something to do with their time instead of protesting their condition.

The analysis is a bit different from what we gave above since now the government is supplying goods by fiat (really, buying or producing them at market costs and then selling them off more cheaply) rather than depending on private suppliers. If the government happens to supply all the way out to the quantity that consumers demand, there is no shortage (see the analysis in the next section). If there is a shortage (as is typical) then the analysis we gave above holds. The opportunity cost can be even higher than it would be in the free market, less is supplied, and agents are clearly worse off.

The only modification here is what we mentioned above: the poor can stand in line more cheaply than the rich, and so the opportunity cost the poor might be lower than a higher money price in the free market. However, these are very costly programs and are difficult for government to sustain. In addition, the poor would be better off if the government simply transferred the money spent on these programs to them directly instead of through artificially cheap, low quality goods.

The bottom line is this: (a) we can imagine price controls like this that would benefit the poor in theory, but (b) in practice, most programs like this do not benefit the poor once the quality, quantity, and high opportunity cost of price controlled goods are taken into account, and (c) this aside, there are almost always alternatives that would cost the government the same amount and be more beneficial to the poor.

Admission to exclusive schools and universities: Top universities and exclusive preparatory schools have excess demand at the tuition price they charge. One might wonder why they do not raise prices to exploit this. There is no law that says that Princeton has to charge a measly \$60,000 a year in tuition and reject 95% of the applicants. Such schools seem to impose a price ceiling on themselves. This means that there are rents available to potential students since admission is clearly valued more than the cost of tuition.

We should therefore see rent seeking, and indeed we do. Applicants work hard in high school to get good grades, and to excel in sports, music, and public service. Some students are also naturally smart or otherwise talented. Others have rich or influential parents who can make donations or provide other benefits to the school. Princeton values these high quality students because if it did not have them, it would cease to be a high quality university. If Princeton just let tuition rise to the free market price, smart, interesting, talented, and influential students would be offered a lower price by competing universities. In short, there is excess demand for slots in top universities.

Markets clear through rent seeking on the part of students. Sometimes these are costly actions like studying or doing community service (at least past the point that a student would choose to do so if he could gain admission without such efforts). Sometimes these are in the form of bringing natural talents and intellectual gifts to the university. In this case, students are paying for their admission using their inborn endowments. Finally, parents may pay for admission by making gifts or using influence.

In the end, an equilibrium price is paid to Princeton for admission in some form or another. The best and the brightest, the hard-working and the influential, and the rich and the fortunate end up being concentrated at a few universities. Others who bring less to the table end up in less desirable universities. Is this a good thing or a bad thing? It is not for us to say. However, we can see the effects of this admission system using our analysis, and so can at least begin to engage the question.

You might ask yourself how giving applicants advantages for such things as diversity, equity, inclusion, legacy status, state or country of origin, and other such things, differ from basing admissions purely on grades, testing, activities, or personal accomplishments. All of these characteristics are endogenous or exogenous to various extents.

Getting good grades, and testing well require hard work (and to that extent are endogenous), but also intelligence, family or cultural support or encouragement, and opportunity (which are largely exogenous, and beyond the control of the applicant). On the other hand, race, gender, sexual orientation, religion, legacy status, and state or country of origin, tend to be more exogenous. While not completely out of an applicants control, they are much less mutable.

Thus, the question is, how does it matter if resources, status, or opportunities are given on the basis of endogenous, or exogenous and immutable characteristics. Does it matter how those characteristics relate to the item that is awarded. For example, how does giving land or money to someone of noble birth, compare to giving that same person an MD, and license to practice medicine because of this lineage? What about how endogenous characteristics that agents build relate to the item awarded? Should we give MDs to people who have lots of extracurricular activities, or were rich enough to have extra tutoring for classes or test preparation?

Concert tickets: We often see lines of various kinds for sporting and entertainment events. We discussed this above, but we did not address the question of why Miley Cyrus or the NFL would voluntarily charge something less than the maximum price they could get for tickets. Both fans and the entertainer/team are worse off because tickets are allocated by line standing rather than through higher prices. It is unlikely that this is simply a mistake. Popular bands can depend on selling out, and the Super Bowl does so every year.

What is the explanation? Notice that there is one thing a bit different about event tickets as compared to other goods: the venues have fixed capacities. Once a location for the event is chosen, there are a fixed number of tickets that can be sold. In other words, the supply curve becomes vertical instead of upward sloping. This does not materially change the argument of course since the entertainer/team would still be better off if the free market price were chosen. Fans would still be indifferent between higher ticket prices and standing in line since opportunity cost would be the same under each system. Why are tickets sold this way? Here are some possibilities:

- If fans have different opportunity costs of line standing, then those with lower wages would be willing to wait longer, and so would get all the tickets. This means the young (or unemployed)

will be the majority of the audience. Why is this good? It might be that Miley Cyrus does not want a middle-aged economics professors in the audience. It would be creepy. She would prefer teenagers. The great thing about the young fans is that they buy CDs and concert t-shirts, they go back each time Miley is in town, and they talk to their friends and build buzz and therefore bankable popularity for Miley. Economics professors might buy a CD, but they would not want to be seen wearing a “Miley Rules!” t-shirt. If they happened to have friends, they certainly would not want to admit that they went to see Miley Cyrus. In other words, the low-priced tickets are a kind of **loss-leader** that preferentially allocates tickets to the type of fan that generates other revenue streams for the performer.

- Super Bowl and playoff tickets cannot be purchased directly by standing in a line. Instead, season ticket holders enter a lottery and a small fraction of them get the opportunity to buy one of these tickets at the official price. Thus, when a fan buys a season ticket, he takes this potential benefit into account when evaluating how much he is willing to pay. In other words, this lottery strategy by the NFL raises the demand (willingness to pay) for season tickets. Why not just get all the extra revenue by increasing the cost of playoff tickets? It may be because fans tend to overestimate both the odds that their team will be in the playoffs, and their own chances of winning the playoff ticket lottery if this happens. Thus, they overestimate the expected benefit, and so overpay for season tickets, which ends up bringing in more revenue than simply selling the playoff tickets.
- It might be that the publicity from having sold out concerts or having it reported that scalpers are getting fantastic prices for tickets to the Super Bowl reinforces the public idea that the band or football is popular and desirable. Any reporting about fans desperate to get tickets in the press is free advertising. Entertainers/teams are willing to give up some of the potential ticket price in exchange for this benefit.

Subsection 7.4.4. Price Ceiling without Shortage

A draft is a special kind of price ceiling in which the government forces agents to supply production up to the quantity demanded. For example, during the Second World War, the government wanted more soldiers than were willing to volunteer at existing military wages. It instituted a draft using a lottery (but without resale). In the Civil War, on the other hand, there was a draft, but one could choose to send a substitute in one's place. Thus, there was a kind of resale after the lottery. Again, it matters precisely how the draft policy is implemented. To understand this, consider the following figure:

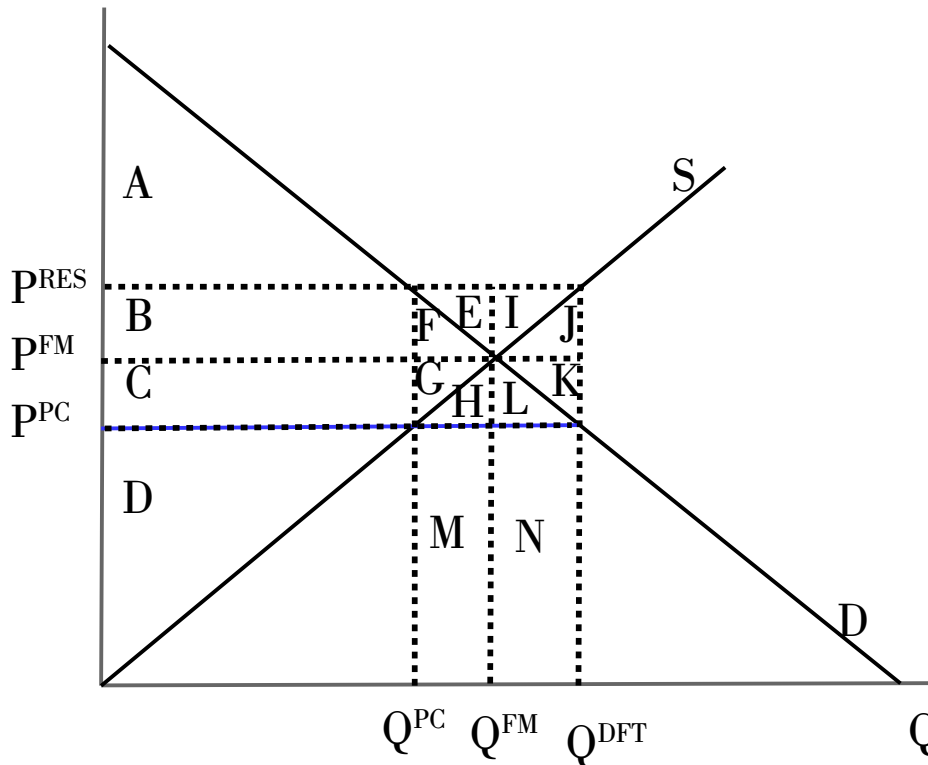


Figure 90: Graphical analysis of price ceiling without a shortage

	FM (Volunteer Army)	VSG	L with RS	L without RS
CS	ABF	ABCFGHL	ABCFGHL	ABCFGHL
PS	CDG	D – JKLH	D – JKLH	$\underline{PS} < D - JKLH$
SS	ABCDG	ABCDG – JK	ABCDG – JK	$\underline{SS} < ABCDG - JK$
DWL	-	JK	JK	$\underline{DWL} > JK$

Table 5: Welfare analysis of a price ceiling without a shortage

Volunteer Army: Here we pay recruits the equilibrium wage P^{FM} . This has the two effects. First, the army reduces its demand for soldiers to Q^{FM} which is where the marginal benefit equals marginal cost. Second, it causes only the people on the lowest part of the supply curve to choose to join the army. Thus, we get the right people in the right numbers and so the DWL is zero. Why would we use a draft rather than a volunteer army? Here are some possible explanations:

- It might be that we think that defending the country is an obligation that should be shared by all citizens. Both the rich and poor should serve and risk their lives if it is necessary. This is a normative statement of philosophical values. In a more formal sense, this is an assertion that the gains to the society from having all citizens understand that they have an equal stake in the country's defense, outweigh the losses from letting a grocery bagger stay at home while an accountant has to go off to war. Economics has nothing to say about whether this is true or false, it can only help assess the cost of this social policy.
- Raising the wage paid to soldiers to a high enough level to induce enough people to join the army could be very costly. Lots of new recruits are needed and there are other calls on government revenues in times of war, buying guns, tanks and bombs, for example. Increasing income, corporate profit, sales, and other taxes, to pay for all this might be very distortionary. That is, the high marginal tax rates needed to pay for the war would strongly discourage people from working and corporations from investing. In wartime, just the opposite is needed. A nation needs to mobilize all the resources it can. The draft is an alternative way of raising an army. Instead of taxing earned income, the government, in effect, imposes a lump-sum tax of four years of service for the act of being healthy, male, and between the ages of 18 and 25. While agents can avoid paying income taxes by not working as many hours, they cannot really avoid being young, male, and healthy. Thus, this lump-sum tax on labor (or more accurately, occupying a certain demographic and owning an endowment of labor) is not distortionary. Again, we might draft some of the wrong people, but the overall DWL to the economy might be less since other taxes would be lower under a draft than with a costly all volunteer army. This is a completely positive rational for a draft, and so can be evaluated with economic tools.

VSG: A very smart government tries to minimize the inefficiency of the draft. The way to do this is to draft only the people with the lowest opportunity cost, that is, with the lowest reservation price. This means agents with a reservation price of P^{RES} or below on the supply curve are drafted, while those with a higher opportunity cost are not. As a result, the DWL is JK. To see this, note that army chooses to draft Q^{DFT} people at wage P^{PC} . This means that MB to the army of having the agents between Q^{FM} and Q^{DFT} is below MC to these agents of serving. That is, this last group of recruits has high value (above P^{FM} in the civilian sector, but a lower value (below P^{FM}) to the army. As a result, we get the standard DWL see whenever any government policy (such as a subsidy) forces production and consumption past the point where MC equals MB .

L with R: All agents participate in the lottery. By the same logic we gave for price floor with a shortage, there will be exactly as many people not drafted with a reservation price of P^{RES} or below as there are people with a reservation price of P^{RES} or above who are drafted. In other

words, the government drafts a total of Q^{DFT} people. For every person on the supply curve below this quantity who “loses” the lottery and is not drafted, there must be exactly one person on the supply curve above this quantity who “wins” and is drafted. *I claim as before that $P^{\text{RES}} - P^{\text{PC}}$ is the market clearing price for lottery tickets.* In this case, however, this is a price paid by a “winner” to convince a “loser” to take the ticket off his hands, and with it, the obligation to serve in the army. Clearly, any agent with a reservation price of P^{RES} or below who is not drafted would take the offer since then his net wage would then be $P^{\text{PC}} + (P^{\text{RES}} - P^{\text{PC}}) = P^{\text{RES}}$, which is greater than his *MC* of joining the army given this next best alternative. On the other hand, any agent with a reservation price of P^{RES} or above has to pay $P^{\text{RES}} - P^{\text{PC}}$ to avoid the draft. But since his wage goes up by more than this if he works in the private sector rather than the army, he will happily pay it. Thus, the supply and demand of draft tickets are equal at this price. The effect of this is the same as we see in the VSG. The lowest opportunity cost people end up serving in the army, and DWL is JK since the army drafts too many people. The only difference is that some wealth is transferred from high value lottery winners to lower value lottery losers.

As we mentioned, this policy was actually used in the American Civil War. If you were drafted, you could send a substitute (usually paid) to serve in your place. Unfortunately, many of these substitutes would take their fee, enlist, desert, and then start the whole process over again with a new client. This is a good reminder that just because a policy improves welfare in theory does not mean that it will in practice if improperly implemented.

L without R: Now, if you are drafted, you serve. This means that some high value agents have to give up their jobs and join the army while some lower value agents are not drafted and continue work in the private sector. The net loss is the difference between their wages. For example, suppose that the army pays \$1000 per month and a plumber earning \$3000 per month is drafted while a telemarketer making \$1500 per month is not. If we had inverted this outcome, then society would gain \$3000 of services from the plumber, but only lose \$1500 from the telemarketer. Thus, society would be \$1500 better off while still having the same number of soldiers. Of course, it would be better not to draft the telemarketer either since the marginal benefit to society of the last soldier drafted is only \$1000. Thus, we still lose JK due to overuse of people by the army, but we also lose an additional amount because the wrong people are being drafted. How large this loss is depends on how the draft lottery actually turns out.

Subsection 7.4.5. Price Floor without Surplus

Price floors are a legal minimum price set by government mandate. The primary example is agricultural price supports of various kinds. The problem with requiring that the price of a good be above the free market equilibrium price is that producers want to produce more than consumers want to consume. This generates a surplus of goods which is generally resolved by the having the government buy up the excess agricultural production at the price floor in order to support the policy. The government is sometimes called the *buyer of last resort* since farmers are obliged to try to find a buyer in the regular market before turning to the government.

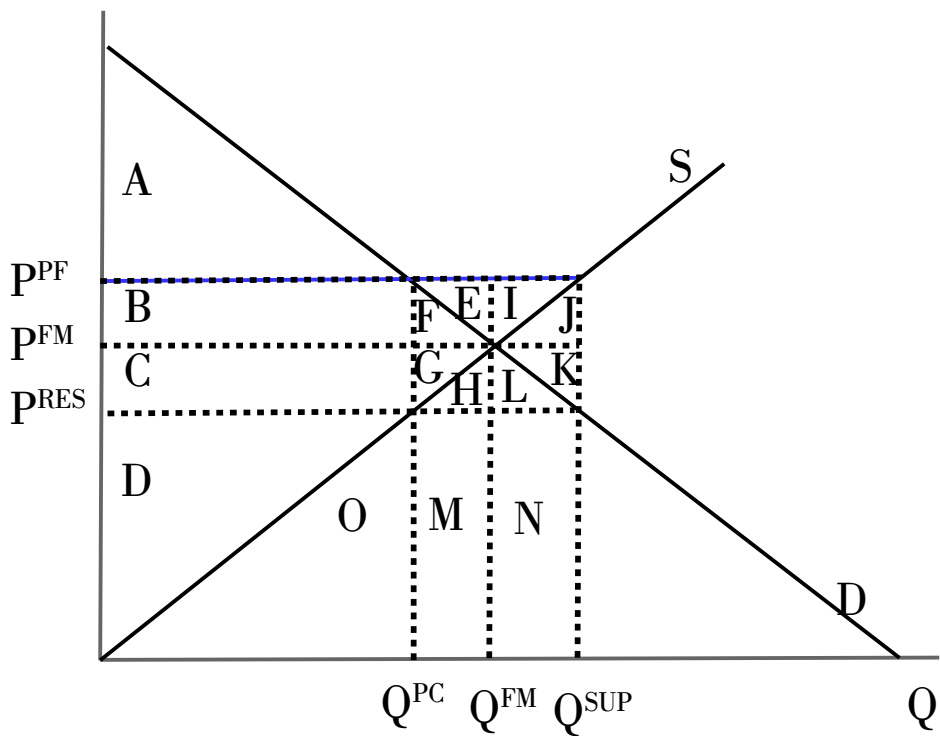


Figure 91: Graphical analysis of a price floor without a surplus

	FM	VSG/GDP	TO
CS	ABF	A + FGHLMN	A
PS	CDG	BCDEFGI	BCDEFGI
GC (Subtracted from SS)	-	EFGHIJKLMN	EFGHIJKLMN
SS	ABCDG	ABCDG - JK	ABCD - JK - HLMN
DWL	-	JK	JK + FG + HLMN

Table 6: Welfare analysis of a price floor without a surplus

How much DWL this causes depends on what the government does with the surplus. The two main alternatives are:

Give it to the Deserving Poor (GDP): A policy thought experiment where a surplus of goods are resolved by giving away for free to those agents who do not buy at the government imposed above free market price floor, but who are the next highest on the MB or demand curve. This is a form of a “very smart government” approach since it requires the government to somehow magically determine an agent’s reservation price.

Throw it in the Ocean (TO): A policy thought experiment where surplus of goods are resolved by are simply dumping or destroying any surplus that is not bought by the public at above equilibrium prices.

The results of these choices are the following:

GDP: Suppose the government can identify the agents with reservation prices between P^{PF} and P^{RES} gives them the good for free. In addition, the government can somehow prevent resale of these goods to agents who would choose to buy in the free markets (that is, with reservation prices above P^{PF}). Then agents who buy in the market pay a higher price for a smaller quantity of goods than before the price floor. Their surplus as a group drops from ABF to just A. However, the next group of agents on the demand curve get the good for free. The rule for calculating consumer surplus now becomes “below the demand curve, over (instead of 'out to') the quantity and above the price (zero in this case)”. All of these agents see an increase in their consumer surplus which now totals FGHLMN. The producers are clearly better off since they sell more at a higher price. The government, however has to buy up all the surplus (the quality between Q^{DEM} and Q^{PF} which costs $P^{PF} \times (Q^{PF} - Q^{DEM})$). Note that the cost here is a vertical rectangle rather than a horizontal one, which would have been the case with a subsidy. This policy ends up generating the exact same DWL that an equivalent subsidy would, JK, because of the overproduction that both induce.

TO: As above, people in the market buy Q^{PF} at the price floor and get a CS of A, while the government cost is EFGHIJKLMN, and must be subtracted from social welfare. Producer surplus goes up just as in the GTDP policy. However, rather than give this government purchased supply away and allowing the deserving poor to get a CS of FGHLMN, the surplus stock is destroyed. Thus, we cannot add this back into the social welfare, and DWL increases as a result. This is the worst policy possible from a social welfare point of view, Unfortunately, it is not at all uncommon. Surpluses of crops are often destroyed in a literal and deliberate way. Stocks may also be destroyed by storing them so long that they spoil, by converting them to different, lower value products (distilling surplus wine into alcohol fuel, for example) or by shipping them outside of the domestic market as foreign aid. The reason that surplus are destroyed may relate to the difficulty of identifying high marginal benefit consumers or of preventing resale.

We will leave the analysis of would happen if the surplus were given away using a lottery either with, or without, resale as an exercise for the reader.

Subsection 7.4.6. Price Floor with Surplus

If the government does not buy the surplus supply when it imposes a price floor, there is an unresolved surplus. Examples include the minimum wage and living wage laws.

To have any effect at all, the minimum wage or price floor must be above the free market wage. If this is the case, then more people want to work than employers wish to hire at the mandated wage. The resulting surplus of labor is what is called “unemployment”. How much DWL this causes depends on how it is determined which of the worker seeking these existing jobs end up finding then.

The three main ways that jobs end up being allocated are the following:

VSG: This unrealistic base case is pretty much as described above. The government somehow knows the reservation wage of each agent and allows only the agents with the lowest opportunity cost (a MC of P^{RES} or below) to take the minimum wage jobs.

SIL: Fixing the wage above the free market level creates an excess supply of labor. Agents engage in search, excess credentialing, line standing, and other types of rent seeking in order to obtain of the available jobs.

HBQW: Hire the Best Qualified Worker. People who apply for jobs differ in quality. Some are smarter, better looking, more charismatic, have good references, have no criminal record, have more or better education, and so on. More highly qualified workers also have better alternative opportunities and so tend to be on the higher part of the MC curve. To the extent that these characteristics are exogenous, other agents cannot acquire them. Getting a job is more a matter of having a lucky endowment than rent seeking in this case. Examples might be intelligence, looks, or charisma. On the other hand, if these desirable characteristics are endogenous and can be acquired by choice, they will become part of agents' rent seeking strategy. For example, agents will spend time and money getting more education, be more careful about getting in trouble with the law, be sure to please their previous bosses or supervisors, and so on. Whatever the case, employers may choose to hire the most highly qualified applicants from the pool of unemployed.

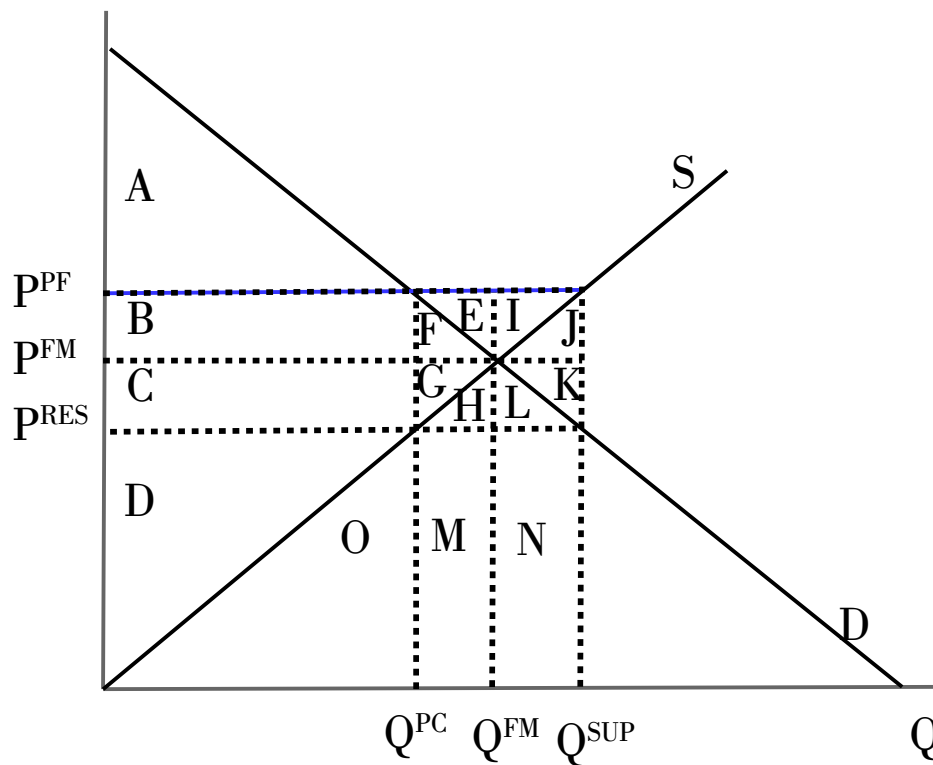


Figure 92: Graphic analysis of a price floor without a surplus

	FM	VSG	SIN	HBQW
CS	ABF	A	A	A
PS	CDG	BCD	D	D
TR	-	-	-	-
SS	ABCD	ABCD	AD	AD
DWL	-	FG	FG + BC	FG + BC

Table 7: Welfare analysis of a price floor without a surplus

VSG: If the government can match the lowest opportunity cost agents to minimum wage jobs, then the producer surplus (remember that workers are the producers of labor here) goes from the free market level of CDG to BCD. Thus, workers as a group gain D, but lose G. In the illustration, D is clearly bigger than G, so workers as a group gain. In general, this conclusion depends upon how high the minimum wage is set and how steep the supply and demand curves are. These gains and losses are not equally spread over workers. The lowest opportunity cost workers (with reservation wages below P^{RES}) benefit and see their wages go up from P^{FM} to P^{PF} . On the other hand, workers with reservation wages between P^{RES} and P^{FM} find themselves unemployed and suffer the loss of G in producer surplus as a group. Clearly, consumers (employers) are hurt and overall there is a DWL of EF. This is due to the minimum wage causing a reduction in labor consumption below the point where MB equals MC.

SIL: This is the most common way such surpluses are resolved. Specifically, workers undertake expensive rent seeking actions to get jobs such as spending extra time in school, getting a new suit or haircut, and doing research about the company before the interview. As before, the cost of these rent seeking actions must be exactly $P^{\text{PF}} - P^{\text{RES}}$ for the market to clear. This would imply a net wage of P^{RES} , which implies a labor supply of Q^{PF} , which is the same as labor demand at a minimum wage of P^{FM} . Unfortunately, this means that all workers are worse off than they would be in a free labor market. Some workers become unemployed, and those who keep their jobs end up with lower net wages after search and line standing costs are taken into account. Over all, the producer surplus drops to D and the DWL becomes $EF + AD$. The EF part is again due to the shrinkage of the labor market, but AD is due to the wasteful efforts that workers must undertake to get jobs when there is unemployment.

HBQW: SIL assumes that the worker can control the likelihood that they will find a job or be attractive enough to be hired by spending effort and money. Suppose instead that the things that make a worker attractive to employers are exogenous, and not under the control of workers. Employers still hire the workers they consider to be the best qualified, but now the choice falls on people who just happen to be smart, beautiful, or fortunate enough to be born with other attractive characteristics. Of course, these workers also tend to have the best outside opportunities and so have the highest reservation wage. Thus, employers hire a total of Q^{PF} workers but from the *highest part of the supply curve possible*. In other words, instead of starting from the worker at the lowest end of the supply curve and then hiring workers until he gets to Q^{PF} , the employer starts at the worker with a reservation wage of P^{FM} , and hires backwards from the higher to the lower part of the supply curve until a total of Q^{PF} workers are hired. What is the PS that the agents who are hired receive? It is the triangle under P^{PF} , down to the MC and running from quantity $Q^{\text{SUP}} - Q^{\text{PF}}$ to Q^{SUP} . Notice that this is exactly the same as the triangle under P^{RES} down to the MC and running from quantity 0 to Q^{PF} . Thus, the PS is the same regardless of whether getting hired in the face of surplus supply of labor depends on costly, endogenously chosen characteristics, or unalterable natural endowments. Again, no agent, worker, or employer is better off when the minimum wage is imposed. The difference is that if employers tend to value exogenous worker characteristics, those at the very bottom of the supply curve with the worst alternative opportunities end up being unemployed, while if employers value endogenous characteristics, these lowest MC workers crowd-out the agents at the top end of the MC curve who end up unemployed instead.

Discussion

Wishing to improve the well-being of the lowest paid workers may be a laudable goal (I think it is, but this is my personal opinion). What comes out of the analysis above is that it is very difficult to do so by attempting to manipulate markets. Such efforts generally create rents which are arbitrated away through actions chosen by self-interested agents (including the very agents the market manipulation is aimed at helping). Attempting to hold back this force is a little like trying to push back the tide. Even worse, the rent seeking opportunities these manipulations create often make things worse, not better for the very agents they were aimed at helping.

Milton Friedman said it this way: “One of the great mistakes is to judge policies and programs by their intentions rather than their results.”

This suggests what should be a general principle for all policymakers:

GOOD INTENTIONS ARE NOT ENOUGH. DO THE MATH.

In other words, make sure that the policy you advocate actually achieves your policy objectives, what ever they happen to be.

Let us return to the minimum wage/living wage debate. So far, the analysis suggests that minimum wage rules do not improve workers' welfare. Why do they have any political support at all given this? To solve this mystery we return to one of the first maxims of detective work: *Cui Bono?* This is Latin for “Who benefits?” (more accurately, “To whose benefit”).

- One group that benefits are less qualified workers who already have jobs and so will not have to stand in line or search after the minimum wage law is passed. This is a little like rent control. Current workers who keep their jobs benefit, but future workers compete these extra rents away and are actually worse off. Thus, we have a push for higher minimum wages from poorly paid workers and their advocates, but no countervailing push on behalf of future worker who do not yet exist. It is easier to worry about the immediate problems of flesh and blood people, than the future problems of young people not yet in the labor force.
- Unions and skilled workers also benefit since minimum wage laws make unskilled labor more expensive relative to skilled labor. Thus, employers switch out of unskilled labor in favor of skilled, and sometimes union, labor. This keeps the poorest of the poor and the least skilled workers from accessing the bottom rung of the employment ladder. In fact, it removes this rung altogether. Workers must start at the second or third rung, and only those with skills or higher education levels are able to make this leap.
- Minimum and living wage laws may even help reinforce racism and segregation. Suppose that one thing that employers value is the race of their employees. This might be because they feel more comfortable working with people who share their cultural backgrounds (we call this a form of **homophily** in economics), or they might think that their customers prefer to be

served by people of a certain race, or they may simply be racist. Race is one of those exogenous characteristics and so no rent seeking is possible. The disfavored race will end up having few outside opportunities, will therefore be on the lowest part of the supply curve, and thus, will end up unemployed. The favored race will be preferentially hired by employers at the higher minimum wage. The traditional economic story is that this outcome is difficult to sustain. It creates a pool of presumably equally qualified unemployed people of the disfavored race who can be hired for less, and thus, increase the profits of anyone who chooses to employ them. Living wage laws prohibit unemployed people of the disfavored race from undercutting the wage of someone of favored race, and so it short-circuits the market mechanism from correcting this inequality.

Suppose you wanted to help low-wage workers. Is there any set of circumstances where a living wage law might be effective?

First of all, it is hard to make the case that a single company or institution should adopt a living wage standard on its own. Suppose Vanderbilt University, for example, decided to pay all employees at least \$15 per hour with benefits. Then Vanderbilt would be the most attractive place in Nashville for low skilled workers to be employed. Every time Vanderbilt had an opening for a groundskeeper or janitor, it would receive hundreds of applications. The lucky applicant chosen for this job would see his salary almost double for doing the same work as before, plus he would get benefits.

How would, or should, Vanderbilt choose who to employ? The fairest thing might be to choose the person who is objectively the best qualified for the job. But then you would find Vanderbilt hiring unemployed people with degrees in botany as groundskeepers, experience in environmental engineering as janitors, and Ph.Ds in English as receptionists. This might even be good for Vanderbilt. Perhaps the higher quality, well paid and happy employees would contribute enough to make worthwhile to pay more than market wages. However, few of the people who currently fill these slots at Vanderbilt would pass this “best applicant” test.

In the short run, existing groundskeepers and janitors would benefit, but as they quit, retired, or were fired, they would be replaced by an entirely different group. Often (in my experience) such plum jobs go to the children, family, or friends of faculty members and other current employees. This makes it doubly difficult for someone at the beginning of his work life but with no connections to hope for one of these jobs.

What if instead we just had a lottery, perhaps even restricted to people with lower qualifications to make sure the living wage benefits the poor and not a bunch of professors' children? We would still have way too many applicants compared to jobs. The winners of the lottery come out ahead, of course, but most people would lose and get no benefits. The question then is this: Why would Vanderbilt want to look at an applicant pool of 500 deserving and sufficiently qualified people and choose, say 25, of them to give a bunch of money in the form of a well paying job? What makes these 25 people special, and why does it make sense to ignore the welfare of the other 475? If Vanderbilt wants to engage in charity, why not find the most deserving and instead of tying our charitable actions to either employment or random selection?

Second, suppose we decided to impose a living wage city-wide instead of only at a single workplace. Not only would we see the same kinds of problems outlined above (the lowest skilled workers left behind and unable to find employment, and rent seeking which dissipates the value of the higher wages), but we will also see low skilled jobs to leaving the city for other places. Why would you put a call center in a city where you had to pay \$15 per hour when you could put it somewhere else and pay only \$8 per hour? Thus, jobs that paid \$15 before would stay and workers would be no better off, but jobs that paid less than this would tend to be relocated to other cities and the workers that used to hold these jobs would end up unemployed.

Of course, this could not happen completely. Some low-skill jobs would stay. The city would still need restaurants, and restaurants need servers. Grass would still need to be cut, and groceries would still need to be bagged. There would be rent seeking and selection for these high pay/low-skill jobs, but they also would be minimized to the extent possible. For example, landscaping firms might invest in bigger mowers so they could do the same job with less labor, but more capital.

Finally, suppose we decided to impose a living wage at the national level. This probably has the best chance of doing some good. We still have rent seeking and selection problems, but it is much harder to move jobs completely out of a country than to the next town over. We would still see capital being substituted for labor and perhaps the wholesale exit of low-skill industries to other countries.

Given this, what is the best case we can make? One thing that will make a difference is how elastic the labor supply and demand curves are. In the figure below, both are quite inelastic near the equilibrium prices. Note that the minimum wage does not affect the quantity of labor consumed very much: Q^{PC} , Q^{FM} and Q^{SUP} , Are now almost right on top of one another. Thus, if:

- Employees are chosen on the basis of exogenous qualities instead of on endogenous ones (which would lead to rent seeking).
- Firms do not substitute out of low-skill labor when wages go up and so the demand curve for labor is relatively inelastic.
- This demand curve is also relatively stable over time (that is, not only is it steep today, perhaps because it is difficult to move a firm or change technology quickly, but also, the demand curve does not move in the long run when it is easier to change locations or substitute labor for capital).
- The supply curve for labor is relatively inelastic so that increases in the minimum wage does not encourage many new workers to try to enter the labor force and thereby increase unemployment.

Then: as before, the workers the top end of the supply curve are chosen over those at the bottom end, who become unemployed. In the figure, however, you can see that the gap between Q^{PC} and Q^{SUP} is very small. Thus, the number of unemployed agents are the very bottom of the supply curve is similarly small. The newly employed agents at the very top (between Q^{FM} and Q^{SUP}) are better off. The rest of the agents in the large middle of the supply curve were employed before and are still employed after the living wage is imposed. These agents used to make P^{FM} , but now they make P^{FM} ,

so they are all better off. This is only the case, however, if they do not dissipate their gains in rent seeking efforts. If they do, then the effective wage goes down to P^{RES} , and they are worse off.

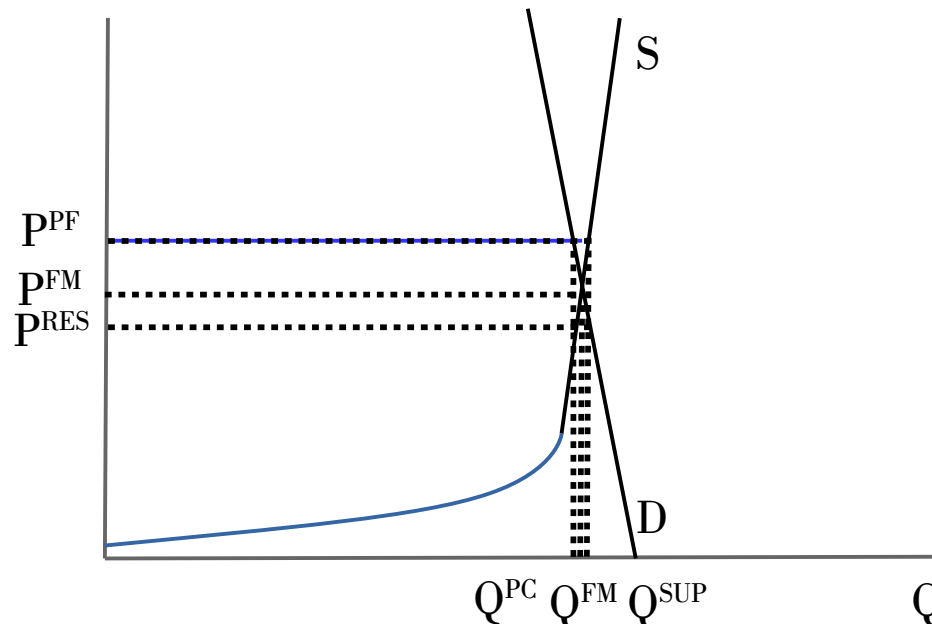


Figure 93: Graphical analysis of price floor with inelastic demand and supply

The basic difficulty here is that we are trying to produce a policy that allows workers get more than the marginal value of what they produce. Doing so without generating enough unemployment so that the marginal worker does in fact have this higher value, or causing competition for these over-compensated jobs which results in this rent being dissipated, is tricky. If it were possible, a more direct and stable approach is to make sure that workers have the skills and education so that their marginal value allows them to earn a living wage in the market. This would not require pushing against the tide.

Subsection 7.4.7. Tariff in a Small Country with no Domestic Production

A tariff is a tax imposed on the import of foreign goods. Similar domestically produced goods are not taxed at all. Tariffs are often set up with the objective of giving an advantage to domestic firms over foreign firms in an industry. Sometimes there is no domestic industry to protect, however. Let's begin by exploring the effects of a tariff in this most simple case.

We consider a small country with no domestic production. A country is “small” in this context if it is a price taker. That is, if it faces a world price for the imported good and its domestic consumption choices do not affect this price. This is the typical situation for most goods and countries. Spain does not affect the price of computers and Argentina does not affect the price of motorcycles. In addition, neither country makes significant numbers of either of these goods domestically. This means that the supply curve a small country faces is flat, that is, completely elastic. It is a price taker in the world market for the imported good.

It is conventional to show a tariff graphically as an excise tax paid by foreign firms to sell their goods in the domestic market. The effect of the tariff is therefore to raise the world supply curve as seen by domestic consumers by exactly the amount of the tariff. We can see that applying a tariff in this situation causes a loss of consumer surplus but a gain of tax revenue. For the country as a whole, the net loss is C. It is unambiguous that the country as a whole is worse off. This is because the country has no domestic industry to protect and has no market power that might allow it to lower the world price.

The only justification for a tariff in this case is as a revenue source. Gaining the tax revenue of B costs the country C. However, income, sales, or corporate profit taxes might have even less attractive ratios of revenue to DWL. Tariffs, in this case, are a just another fiscal instrument that should be used in balance with all the others.

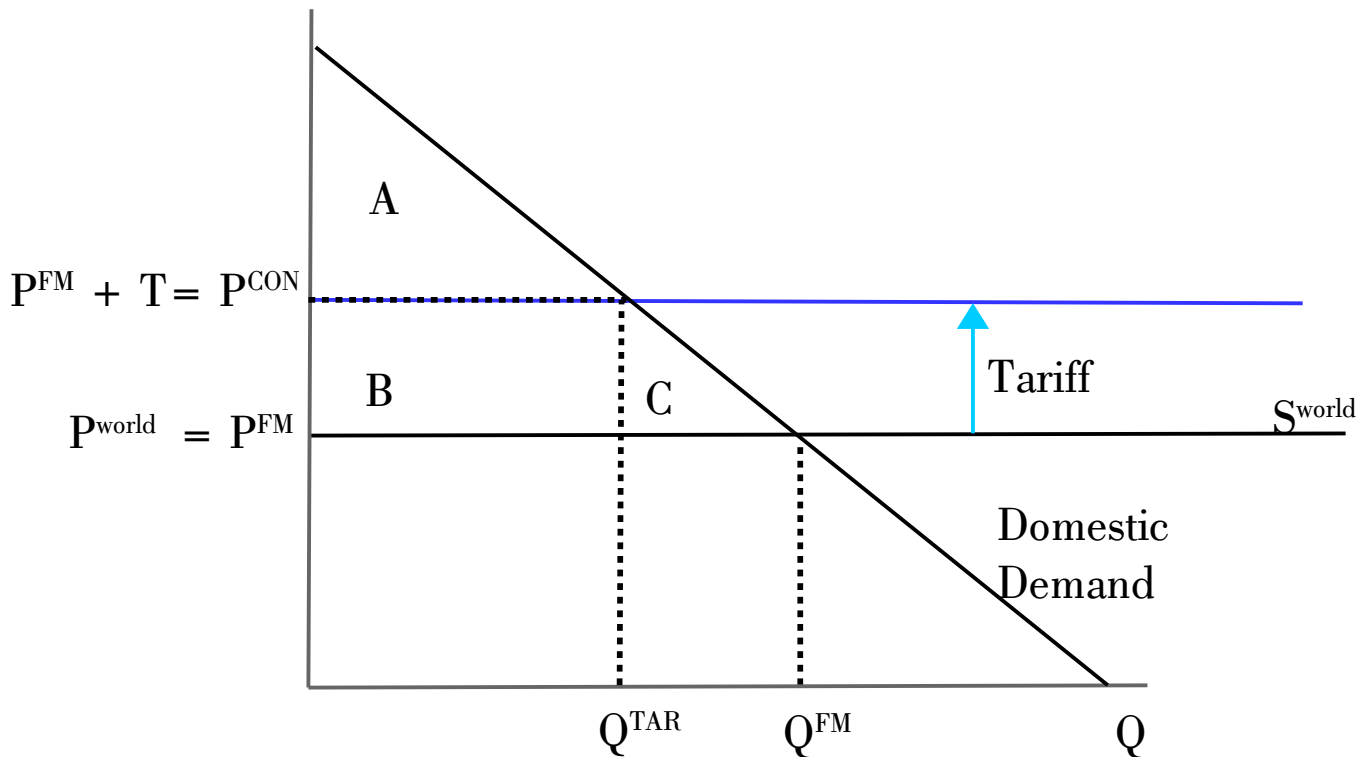


Figure 94: Graphical analysis of tariffs in a small country w/o domestic production

	FM	Tariff
CS	ABC	A
PS (no domestic producer!)	-	-
TR	-	B
SS	ABC	AB
DWL	-	C

Table 8: Welfare analysis of a tariff in a small country w/o domestic production

Subsection 7.4.8. Tariff in a Large Country with no Domestic Production

A country is “large” in this context if it can affect the world price of the imported good. That is, if it faces an upward sloping foreign supply curve. This is much less common case. First, it is rare that a country is big enough to affect world prices. This would require that a country consume a significant fraction of the world’s production of a given product. Even the US consumes only 20% of the world’s oil and less of most other commodities. Second, it is even rarer for such a country to produce none of a good that it imports in such high quantities. An example of where both conditions might be satisfied are Japanese imports of Atlantic Tuna.

We can see again that imposing a tariff in this situation causes a loss of consumer surplus, but a gain of tax revenue. For the country as a whole, there is a loss of H , *but there is also an offsetting gain of CG which comes at the expense of the foreign producer*. Whether this is a good or bad thing for the country as a whole depends on which of these areas is larger.

Intuitively, what is driving this result is that imposing a tariff artificially holds back domestic demand. In a sense, the tariff is a device to coordinate the actions of domestic consumers in the country’s collective interest. By holding back demand, the large country is able to move down the foreign (world) supply curve and get a lower net price for the good. The price actually paid to the foreign suppliers is $P^{PRO} < P^{FM}$. Such a tariff can be calibrated to be optimal in the sense that $CG - H$ is maximized.

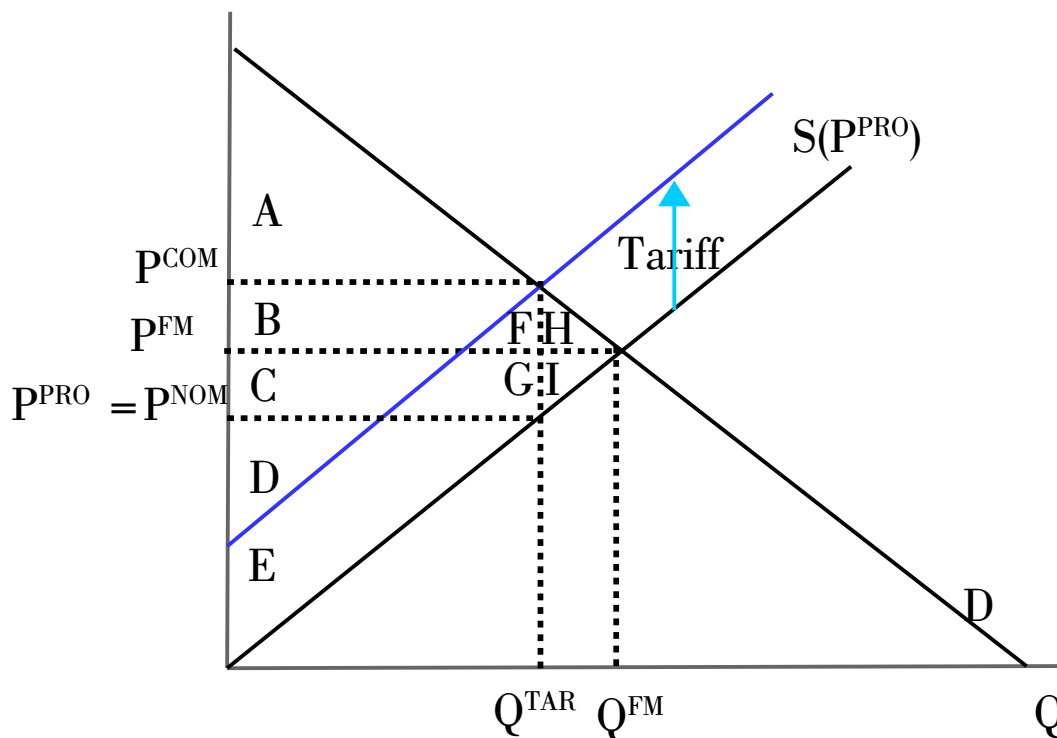


Figure 95: Analysis of tariffs in a large country w/o domestic production

	FM	Tariff
CS	ABCFH	A
PS (no domestic producer!)	-	-
TR	-	BCFG
SS	ABCFH	ABCFG
DWL (possible gain)	-	H – CG

Table 9: Welfare analysis of tariffs in a large country w/o domestic production

Subsection 7.4.9. Tariffs and Trade War Between Large Countries

So far, we see that small countries lose when they impose tariffs, but that large countries might gain. We assumed, however, that other countries do not respond to the imposition of tariffs by their trade partners. This is a rather doubtful assumption. Countries typically impose counter-tariffs and sometimes trade wars can result. Recent US administrations have negotiated several “free trade zone” treaties to head off trade wars and lower or remove tariffs.

To begin to understand the logic of this, suppose there are only two countries in the world who are engaged in trade. One exports good A and imports good B, while the other does the opposite. Neither country has a domestic industry that manufactures the good it imports, and neither country consumes the good it exports. Finally, assume the demand curves for good A and good B are identical, and the supply curves for good A and B are identical. Thus, we can use the same picture to show the effects in both markets, A and B, for each country. In one case, a given country is on the supply side and gets the producer surplus, and in the other case, that country is on the demand side and gets the consumer surplus and the tariff revenue.

In the free trade equilibrium, the domestic country gets more SS from the export market B but less from the import market A than it would under the trade war situation. Combining both markets, it gets the maximal sum of the producer and consumer surplus possible from a single market. In the trade war situation, however, the domestic country finds that the price of its export good is driven down by the foreign country’s tariff, but gains in the import market as compared to the free trade equilibrium because of its own tariff. Unfortunately, the sum of these gains and losses is a net loss of HI. Both countries would therefore be better off if they negotiated a free trade agreement, and ended the trade war.

How does this correspond with the real world? In my example, I assumed two completely symmetric markets. However, it may be that the US is more often an importer of goods than an exporter (consider our large trade deficits). It may also be that we have market power in more of our import markets than our trading partners have in our export markets since we are a relatively big country. To the extent that this is the case, we benefit more often (gains in our import markets) than we are harmed (losses in our export markets). We might be better off with a trade war than free trade as a result. This is an unlikely situation for most countries, however. It is hard to imagine that any, but the very largest countries could ever benefit from trade war. Most small and medium-sized countries would be better off with free trade pacts.

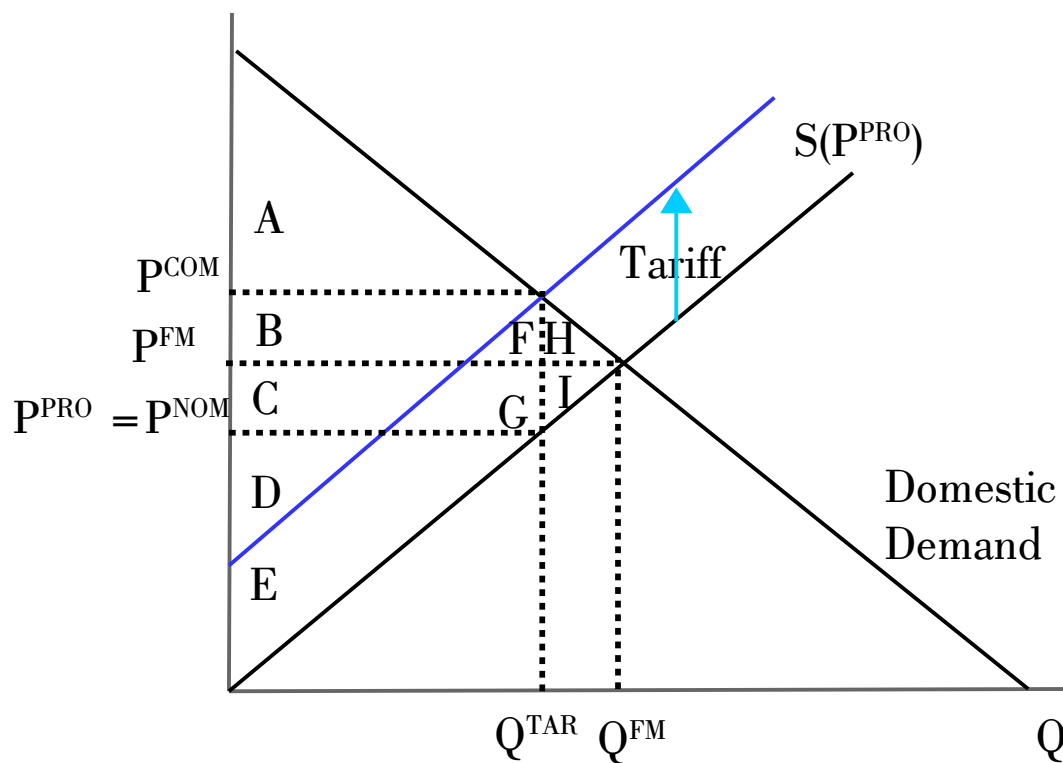


Figure 96: Graphical analysis of a trade war between large countries

	FM	Tariff and Trade War
CS (Market A)	ABCFH	A
PS (Market B)	CDEGI	DE
TR (Market A)	-	BCFG
SS	ABCDEFGHI	ABCEFG
DWL	-	HI

Table 10: Welfare analysis of a trade war between large countries

Subsection 7.4.10. Tariffs and Quotas in a Small Country with Domestic Production

The classic justification for tariffs is to protect domestic industry from foreign competition. As we mentioned above, most countries are “small” in the context of international trade in that they have little or no influence over world prices. However, many countries also produce things they import. Examples include wine, cheese, furniture, seafood, etc.

First, consider the free market. The domestic country can freely import all it wants at the world price. Thus, P^{World} is the price that prevails in the domestic market. At this price, domestic firms choose to produce $Q^{\text{S(FT)}}$ and consumers choose to consume $Q^{\text{D(FT)}}$. This means that the country imports $Q^{\text{D(FT)}} - Q^{\text{S(FT)}}$ units of the good. Note that we add the PS of the domestic firms, C , into the SS now.

Next consider what happens when we impose a tariff. Now domestic consumers have to pay the world price plus the tariff for imports. However, they can still consume as much as they like at a fixed price of $P^{\text{World}} + \text{Tariff}$, and so this becomes the new domestic market price. At this price, domestic firms choose to produce $Q^{\text{S(TAR)}}$ and consumers choose to consume $Q^{\text{D(TAR)}}$. Imports are reduced to $Q^{\text{D(TAR)}} - Q^{\text{S(TAR)}}$ units of the good. We can see that the tariff policy has served the purpose of expanding domestic production and providing benefits through higher prices and higher PS to domestic manufacturers. The government has also benefited by gaining the tariff revenue, F . However, these gains are more than offset by the losses to consumers. Whether this is a good policy depends upon whether you think consumer welfare is more important than government and producer welfare.

Another policy instrument that is sometimes used to protect domestic industry is called an **import quota**. The idea is that foreign firms are given a fixed number of permits to export goods to the domestic country, but that these goods are *not* subject to tariffs.

To make things comparable, fix the quota at $Q^{\text{D(TAR)}} - Q^{\text{S(TAR)}}$, the quantity imported under the tariff we just analyzed. Note that $P^{\text{World}} + \text{Tariff}$, is also the market clearing price with this quota since domestic producers choose an output of $Q^{\text{S(TAR)}}$ at this price, foreign producers export the maximum quantity permitted under the quota: $Q^{\text{D(TAR)}} - Q^{\text{S(TAR)}}$, and the sum of these two equals the quantity demanded, $Q^{\text{D(TAR)}}$. The consumer and producer surpluses are the same as under the tariff, however, F which used to be tariff revenue, now becomes a surplus that goes to the foreign firms lucky enough to get a permit to sell goods that cost P^{World} on the world markets for the higher price of $P^{\text{World}} + \text{Tariff}$. In other words, this policy simply gives tariff revenue to foreign firms.

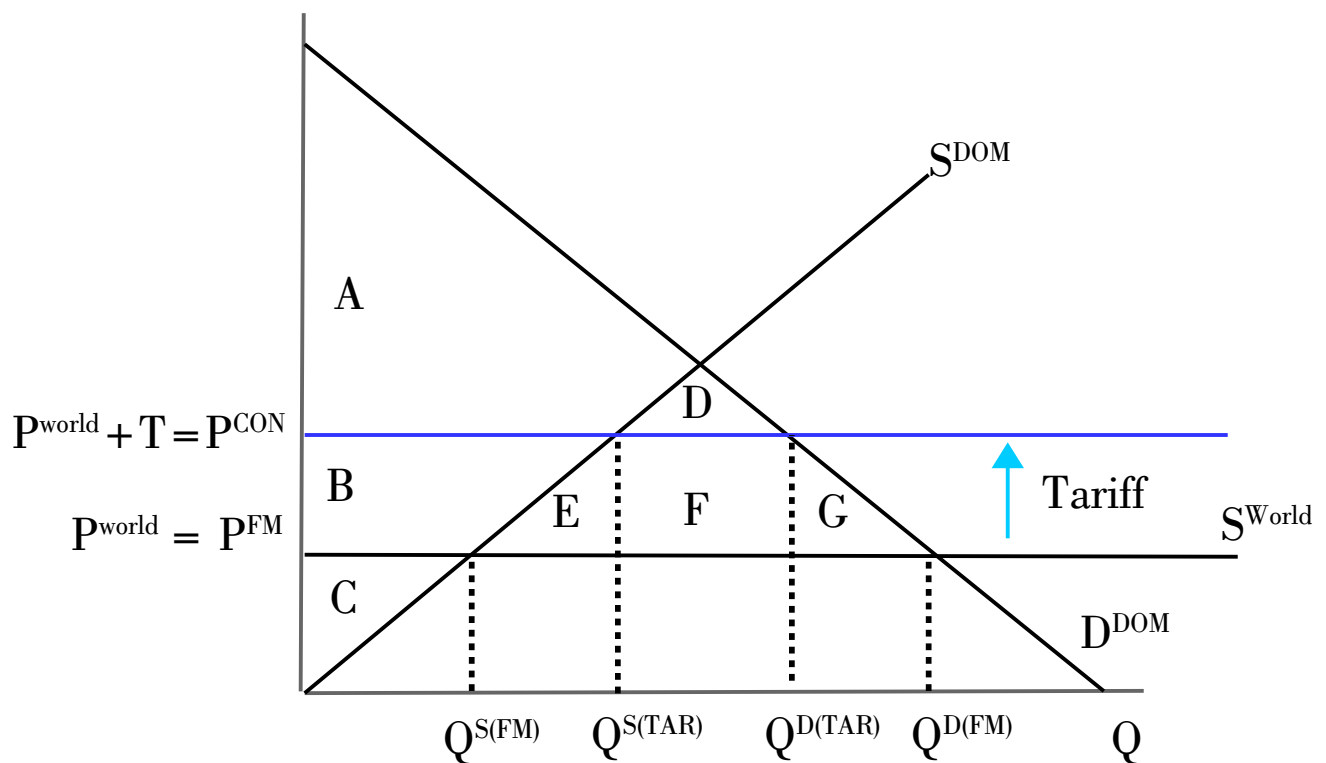


Figure 97: Graphical analysis of tariffs in a small country with domestic production

	FM	Tariff	Quota
CS	ABDEFG	AD	AD
PS	C	CB	CB
TR	-	F	-
SS	ABCDEFGF	ABCF	ABCD
DWL	-	EG	EFG

Table 11: Welfare analysis of tariffs in a small country with domestic production

Why would we give revenue away? It may surprise you to learn that we do so frequently. Examples include sugar, steel, and automobiles at various times in recent history. The reasons are most likely political.

- Without the quota, exporting to the domestic country does not yield much profit for the foreign firms. With the quota, they get a profit equal to the amount of the tariff on every sale. Foreign firms are not likely to object very much, and so will not lobby their government to oppose the quotas or retaliate. With a tariff, they might very well do so.
- Quotas are often given directly or indirectly to foreign governments to distribute to companies that want to export. Thus, Japan might get a quota that allows them to export 500,000 cars to the US, but it is up to the Japanese government to hand these out to Toyota, Nissan, and Honda. In other words, the US government has given the Japanese government a bunch of rents to pass out. Japanese car companies will compete for these rents through campaign contributions, lobbying, moving factories to a key politician's hometown, and even direct bribes. Thus, quotas tend to support corruption in other countries.
- Taking this one step backwards, how does the US government decide how to divide the car import quota between Japan, Korea, and Germany? Unfortunately, you can see that quotas produce rent seeking and corruption in the country that sets them as well!

In effect, quotas convert what would have been domestic tax revenue into rents that are distributed by the importing country to the exporting countries, and then by the exporting countries to local firms. Politicians in both countries, and foreign exporting firms benefit from these rents. Domestic consumers are hurt, but domestic firms are indifferent between quotas and tariffs when seeking protection from overseas competition.

Section 7.5. Bergson-Samuelson Social Welfare Functions

Welfare functions are a method of ranking different allocations that are feasible for an economy or a game. Spoiler alert: they do not make much sense except perhaps as a conceptual tool. However, let's build them up before we tear them down.

More formally a **Bergson-Samuelson social welfare function** (SWF) is a kind of social utility function over distributions of utility for agents.

$$W: \mathbb{R}^I \Rightarrow \mathbb{R}.$$

One can take this idea one step backwards and make the SWF depend directly on allocations. For example: $W(u_1(x_1), \dots, u_I(x_I))$ is a mapping from allocations $x \in \mathbb{R}^{NI}$ into social welfare.

Leading examples of welfare functions are the following:

Utilitarian or Benthamite: $W(u_1, \dots, u_I) = \sum_{i \in I} u_i$

This captures the social value of “the greatest good for the greatest number.” If the social planner cares only about the total utility enjoyed in the society and not at all about who gets utility or how equally or unequally it is distributed, then we call him a Utilitarian. Consider the two agent case. Suppose we put the utility of agent 1 on the x-axis and agent 2 on the y-axis. Now graph the social welfare level sets, that is, the allocations of utility over the two agents that would all give a specific level of social welfare. For example, the level set given by the equation $W = 10 = u_1 + u_2$ is a straight line between (10, 0) and (0, 10). This looks exactly like perfect substitute (linear) indifference curves. In fact, with Benthamite social welfare, the utility of all the agents in the economy are perfect substitutes to the planner. We could also have a weighted Utilitarian SWF such as $W(u_1, \dots, u_I) = \sum_{i \in I} \alpha_i u_i$. The vector of social welfare weights α tells us how much the planner cares about one agent compared to another.

Egalitarian or Rawlsian: $W(u_1, \dots, u_I) = \min \{u_1, \dots, u_I\}$

This captures the social value of “we are no better off than the least of us.” If the social planner is a strict egalitarian, he cares only about the utility of the least well-off agent. Giving more utility to any other agent just makes the allocation more unequal. In the two agent case, you can verify that social indifference curves are right angle shaped and just like indifference curves for perfectly complementary goods. Of course, we could also have a weighted version of this SWF.

Constant Elasticity of Substitution (CES): $W(u_1, \dots, u_I) = \frac{1}{1-\rho} \sum_{i \in I} u_i^{1-\rho}$

This captures the social value of “we care both about total welfare and the distribution of welfare”. This is sometimes presented as balancing equity (like the Rawlsian SWF) with efficiency (like

the Benthamite SWF). You can verify that in general, CES SWFs are curved just like convex utility functions. In fact, we can get both of the examples above as special cases:

- $\rho = 0$ gives a Utilitarian SWF: $W(u_1, \dots, u_I) = \sum_{i \in \mathcal{I}} \ln(u_i) \Leftrightarrow \prod_{i \in \mathcal{I}} (u_i)$.
- $\rho \rightarrow \infty$ gives an Egalitarian SWF: $W(u_1, \dots, u_I) = \min \{u_1, \dots, u_I\}$.
- $\rho = 1$ gives a convex SWF.
- You can verify that the **natural log**, “ $\ln$ ”, and multiplicative versions of the SWF in this case gives the same social indifference curves.

Thus, the CES SWF encompasses both extreme views of equity, and also intermediate cases. We interpret ρ as the social preference for equity. As ρ goes to zero, we care more about efficiency and less about equity. The social indifference curves get flatter and flatter. As ρ goes to ∞ , we care more about equity and less about efficiency. The social indifference curves get pointier and pointier.

Social welfare functions are used in various areas of economics to make judgments about social policy. For example, a social project (a new road, for example) might benefit farmers and merchants a lot, but laborers very little. To evaluate the costs and benefits of this project, therefore, we might use a SWF to determine how much better off society is when the benefits are distributed in this way. In public finance, tax systems have significant distributional effect. How progressive do we want a tax to be, and how do we weigh the redistributive benefits of taxation against the distortions and inefficiencies they produce? If we have a SWF at hand, we can evaluate, compare and even optimize tax and redistribution schemes.

Social welfare functions sound very useful. They appear to allow us to deal with difficult questions involving value judgments over how much we care about the winners and losers from social policies. The difficulty is coming up with a SWF. In short, the problems are these:

- SWF are defined over utility. Utility, however, is ordinal in general. If we doubled the utility values of each consumption bundle for any given agent, we would not change his indifference curves or his demand behavior. To the planner, however, it would look like he had just become twice as good at making utility which is what the social planner values. Thus, to use SWFs, we must have a cardinal measure of utility for all agents that allows us to make both interpersonal comparisons of utility (one util for me has to mean the same as one util for you) and intrapersonal comparisons of utility (two utils is twice as good for me as one util). Unfortunately there is no generally accepted way of cardinalizing utility.
- SWF could be defined over something like money or income instead of utility. The rationale is that dollars can be compared over agents, and two dollars is twice as good as one dollar. This is not really true, however. The first dollar I get I use for a more important purpose than the second dollar. By the same argument, a rich man uses a dollar for a less urgent purpose than a poor man. This aside, some agents may just be better at getting happiness out of dollars than others. It may be that they find bargains, make sophisticated purchases, or are sim-

ply more able to appreciate what they buy. Why should we treat money as something to be equitably and efficiently distributed when the only purpose of money is to gain utility?

- Choosing a SWF is fundamentally the same as stating one's values. SWFs can do no more than give a mathematical description of our ethical, philosophical, or religious feelings. There is no way to choose a SWF using positive economics. They are a normative statement by their very nature. Since the policy results one gets depends on the SWF one chooses, we can see the results as flowing directly from this choice. Thus, we are not making positive, objective, value-free recommendations, but suggesting that policymakers do something in line with a specific, normative world-view. Unless the policymaker gives the economic adviser the SWF as data, the adviser has no basis at all for suggesting an optimal policy in these dimensions.
- SWF are in essence an attempt to aggregate agents' preferences to derive a social ordering over the feasible allocations. Unfortunately, Arrow's impossibility theorem tells us there is no way to do this that satisfies even a minimal set of reasonable axioms. Thus, SWFs are bound to give unsatisfactory results unless one adds domain restrictions over preferences. We go into more detail on this point in the next subsection.

We conclude that the thought experiment of imagining what a benevolent dictator would do, or how one might attempt to make social judgments about the desirability of different allocations (or the policies that imply them) is interesting. However, it is hard to understand any reasonable way to incorporate this methodology into any real world policy problem. Doing so takes us out of the realm of offering technical advice and places us instead in the position of imposing our own world view on outcomes. Awesome as that may sound, the political or ethical views of economists, either individually, or as a profession, are not any more special, and deserve no more privilege, than those of any other citizen.

Section 7.6. Arrow Impossibility Theorem

The social welfare functions described above are actually an example of a (not very successful) attempt to address a more general problem. What we would really like to do find a way to generate a social preference ordering over a feasible set of alternatives by somehow aggregating the preferences of individual agents.

In general, suppose we had a set of feasible alternatives $\mathcal{A}$. This could be almost anything, a set of social proposals (should we build a highway or an airport), or a list of allocations or payments for some or all of the agents in the economy, for example. For simplicity, assume that $\mathcal{A}$ is a finite set.

Suppose that each agent $i \in \mathcal{I}$ has a complete transitive ranking over these alternatives denoted: R_i . We will write:

$$\bar{a} R_i a$$

to indicate that agent i prefers alternative $\bar{a}$ to alternative a . Thus, R_i is essentially just a preference relation like $\succ_i$. For simplicity, we will also assume that each agent is able to *strictly* rank all alternatives and so no pair of alternatives are equally desirable to any given agent: $\forall a, \bar{a} \in \mathcal{A}$ and $\forall i \in \mathcal{I}$, if $a R_i \bar{a}$, then it is not the case that $\bar{a} R_i a$. Let $\mathcal{R}$, denotes the set of all such rankings over $\mathcal{A}$. Thus,

$a \in \mathcal{A}$ social alternatives

$i \in \mathcal{I}$ agents

$R_i \in \mathcal{R}$ strict, complete, and transitive rankings over $\mathcal{A}$ for each agent

Given this, a social welfare function (SWF) is defined as follows:

$$F: (\mathcal{R})^I \Rightarrow \mathcal{R}.$$

That is, a SWP maps a list of rankings for each of the I agents into an aggregate a social ranking over alternatives. The question Arrow asks is: can we establish a SWF that satisfies at least a minimal list of reasonable properties? He proposes the following:

Unanimity/Pareto Efficiency: If all voters prefer a to $\bar{a}$, then the social ranking should place a above $\bar{a}$.

No Dictatorship: There must not be a dictator, that is, an agent whose preference are taken as the social preferences regardless of the preferences of others. In other words, an agent is a dictator if no conceivable set of potential preferences held by any or all of the other agents would ever prevent the social ranking from being identical to the dictator's ranking.

Transitivity: If society prefers a to $\bar{a}$, and $\bar{a}$ to $\hat{a}$, then society prefers a to $\hat{a}$.

Independence of Irrelevant Alternatives (IIA): If a social ranking puts a above $\bar{a}$ when $\hat{a}$ is available, it can never rank $\bar{a}$ above a when $\hat{a}$ is not available. Another way to say this is that the social ranking between a and $\bar{a}$ depends only on the individual preferences between a and $\bar{a}$.

Unrestricted Domain: The SWF must at least consider every logically possible combination of individual rankings over problems with at least three alternatives.

Arrow shows the following:

Arrow Impossibility Theorem: If there are at least two agents in the economy, then there does not exist a SWF satisfying Pareto Efficiency, No Dictatorship, Transitivity, Independence of Irrelevant Alternatives, and Unrestricted domain.

This means that not only is the case that the Bergson-Samuelson approach fails by the criterion listed above, but so do all conceivable approaches of generating a social ranking. For example, voting over proposals as a way of ranking choices will not generally satisfy the four conditions.

This is bad news if we think that the objective of government is to find socially optimal outcomes. On the other hand, it reminds us that there is no value-free way of choosing over outcomes in general. Again, the tools we develop in economics can help inform the decisions of governments and policymakers, but in the end, we cannot escape the need for someone to make noneconomic choices that involve philosophy, religion, ethics, and politics.

Glossary

Benthamite SWF: The Benthamite, or Utilitarian, welfare function asserts that it is socially optimal to maximize the sum of utility over all agents. Thus, it captures the value that all that matters is total wealth in the society, and how this wealth might be distributed over agents is irrelevant. Formally: $W(u_1, \dots, u_I) = \sum_{i \in \mathcal{I}} u_i$,

Bergson-Samuelson Social Welfare Function (SWF): This is a kind of social utility function over utility levels of agents in a society. Equivalently, the SWF can be taken as depending directly on allocations. For example: $W(u_1(x_1), \dots, u_I(x_I))$ is a mapping from allocations into social welfare. The basic purpose of a SWF is to give a social ranking over all possible allocations in order to help find the most socially preferred feasible distribution.

Constant Elasticity of Substitution (CES) SWF: The CES welfare function seeks to balance the total wealth of a society with how it is distributed. Thus, it takes into account both efficiency and equity when ranking allocations over agents. Formally: $W(u_1, \dots, u_I) = \frac{1}{1-\rho} \sum_{i \in \mathcal{I}} u_i^{1-\rho}$.

Consumer Surplus: The area beneath the demand curve, out to the quantity, and above the price. This is the value that agents get in terms of money from buying goods in a given market.

Dead Weight Loss (DWL): This is the differences between the maximum possible social gain and the gain seen in any alternative social policy. Usually the free market maximizes the social gain.

Draft: A special kind of price ceiling in which the government forces agents to supply production up to the quantity demanded.

Egalitarian SWF: The Egalitarian, or Rawlsian, welfare function asserts that it is socially optimal to maximize welfare of the worst off agent. Thus, it captures the value that all that matters is the equality of welfare over agents and that total wealth is in itself, unimportant. Formally: $W(u_1, \dots, u_I) = \min \{u_1, \dots, u_I\}$.

Give it to the Deserving Poor (GDP): A policy thought experiment where a surplus of goods are resolved by giving away for free to those agents who do not buy at the government imposed above free market price floor, but who are the next highest on the *MB* or demand curve. This is a form of a “very smart government” approach since it requires the government to somehow magically determine an agent’s reservation price.

Hire the Best Qualified Worker (HBQW): A policy thought experiment where a shortage of jobs due to mandating an above market wage is solved by giving the remained jobs to the “best qualified” workers, as opposed to allowing rent seeking to decide the allocation. This is a form of a “very smart government” approach since it requires the government to somehow magically determine an a worker’s opportunity cost of working, or perhaps his productivity, depending on what “best qualified” means.

Homophily: An affinity for other agents who are similar to you in some way. This phenomenon tends to result in the segregation of agents across race, class, gender, religious, and other diminutions.

Import Quota: Foreign firms are given a fixed number of permits that allow them to export goods to the domestic country, but does not subject the imported goods to a tariff.

Kaldor-Hicks Criterion of Potential Pareto Improvement: This says that a policy or allocation is socially beneficial if there is a potential for the winners to fully compensate the losers and yet still be better off themselves.

Loss-Leader: A loss-leader is a product or service that is offered at a price below cost with the intention that it will create greater demand for other products and services to those customers that are sold at the same time and at a profit. This strategy can also be sequential in the sense that an entrant to an industry may sell products at a loss at first in hopes of building a customer or user base that will buy products at a higher price later on.

Lottery with Resale (L with RS): A policy thought experiment where a lottery is held to determine who is allowed to buy a good in short supply at a below market price. Allowing resale means that the winners of the lottery can either sell the winning ticket to other agents who value the good more than they do (and therefore the right to buy at a below market price), or sell the good directly at whatever price the market will bear. Note that a lottery can also be used to result in a surplus by only allowing the winners to sell goods at the artificially enforced above market price.

Lottery Without Resale (L without RS): A policy thought experiment where a lottery is held to determine who is allowed to buy a good in short supply at a below market price. Disallowing resale means that the winners of the lottery must consume any goods that they choose to buy and cannot sell them to agents who value the good more than they do.

Price Ceiling: A maximum price that firms are allowed to charge by government mandate.

Price Floor: A legal minimum price set by government mandate. The primary example is agricultural price supports of various kinds.

Producer Surplus: The area beneath the price, out to the quantity, and above the supply curve. This is the value that agents get in terms of money from selling goods in a given market.

Rawlsian SWF: See Egalitarian SWF.

Rent: The difference between what a factor is paid and what its value is in its next most valuable use. Under a minimum wage, for example, the wage necessary to make the number of workers demanded by employers to offer their services is below the minimum wage, and under rent control, the highest willingness to pay for an apartment by an unhoused tenant is above the offi-

cial rent controlled price. This leads to **rent seeking** and sometimes **rent dissipation** which adds to dead weight loss.

Reservation Price: The highest price an agent is willing to pay in exchange for the next unit of a commodity. This price is equal the marginal benefit an agent would receive from buying or consuming the next unit of a commodity. Symmetrically, the lowest price an agent is willing to accept in exchange for the next unit of commodity. This is equal to the marginal cost an agent would incur from producing, or the opportunity cost of selling, the next unit of a commodity.

Social Surplus: The total of the surpluses and losses seen by all agents in a market. This includes consumers, domestic producers, the governments, and any other domestic agent that benefits or is harmed by the market. In the simple case of a market without any outside interference (a free market): $CS + PS = SS$.

Standing In Line (SIL): A policy thought experiment where shortages of goods are resolved by selling them on a first come, first served, basis. This induces rent seeking behaviors like search, bribery, or literal standing in line in order to be allowed to buy the good that is in short supply.

Subsidy: Subsidies are equivalent to negative taxes. When an agent buys or sells a good, the government gives the agent an amount of money equal to the subsidy instead of demanding that the agent pay an amount of money equal to the tax.

Tariff: A tax imposed on the import of foreign goods. Similar domestically produced goods are not taxed at all.

Tax: An amount of money that firms or consumers must pay to the government when they engage in an economic activity. Taxes may be per unit, proportional, lump sum, progressive, regressive, or even more complicated. Responsibility to pay taxes may rest on the consumer, the producer or both.

Throw it in the Ocean (TO): A policy thought experiment where surplus of goods are resolved by simply dumping or destroying any surplus that is not bought by the public at above equilibrium prices.

Utilitarian SWF: See Benthamite SWF.

Very Smart Government (VSG): A policy thought experiment where the government allows only the agents who are on the highest part of the demand curve (and so have the highest *MB* of consumption) to buy the goods. This means the government has to be magically smart enough to know everyone's reservation price.

Welfare Economics: Using economic tools to try to weigh the costs and benefits of various governmental activities. This typically requires interpersonal comparisons of gains and losses and cannot be done in a value-free and purely positive way.

Problems

1. Let the for supply curve for pizzas be given by $S(P) = 4P - 400$, and the demand curve be given by $D(P) = 1000 - 2P$.
 - a. What is the equilibrium price and quantity?
 - b. What is the elasticity of demand for pizza at the equilibrium? (your answer should be a number.)
 - a. What is the consumer surplus at the equilibrium? (your answer should be a number.)
2. Suppose that the demand and supply for iPads is given by the following equations, $D(P) = 4000 - 2P$ and $S(P) = 4P$.
 - a. What is the equilibrium price and quantity?
 - b. What is the elasticity of demand for iPads at the equilibrium? (your answer should be a number.)
 - c. What is the consumer surplus at the equilibrium? (your answer should be a number.)
3. Would workers benefit if the part of social security tax that is currently paid by employees were reduced and the part paid by employers increased by the same amount? Use a picture to prove your answer.
4. The demand for insulin is completely inelastic. The supply ice cubes is completely elastic. Suppose that the government puts a \$1 per unit tax on producers of insulin and consumers of ice cubes.
 - a. In two separate pictures, show the free market price and quantity, and the consumer, producer, and nominal price, and the quantity after the tax is imposed in each market. Now do the welfare analysis of these taxes. Show the consumer surplus, producer surplus, tax revenue, and total social surplus both before and after the tax in each in these two markets. (Thus, you should have four columns of answers: Before the tax in the insulin market, after the tax in the insulin market, before the tax in the ice cube market, and after the tax in the ice cube market)
 - b. Which policy, the ice cube tax or the insulin tax, is better from the standpoint of minimizing dead weight loss?
5. Bruno Mars is in town for one night only and gives a concert in a 1000 seat auditorium. Tickets cost \$80 each, but at that price, 2000 people want to see his show.
 - a. Suppose that tickets are sold on a first come, first served basis. Using a diagram, show the consumer surplus and producer surplus. Compare this to the consumer and producer surplus when tickets are sold at an equilibrium price. Who gains and who loses? Is there a dead weight loss? (Hint: Think carefully about the supply curve. How many tickets are supplied at any given price?)

- b. Suppose instead that tickets are allocated using a lottery, and that people can costlessly resell tickets on the scalper's market. Using a diagram, show the equilibrium price of tickets on the resale market. Show the consumer surplus and producer surplus in this case. Be sure to separate the consumer surplus of lottery winners from that of concert-goers. Compare this to the consumer and producer surplus when tickets are sold on a first come, first served basis. Who gains and who loses? Is there a dead weight loss?
 - c. Briefly discuss how your answer to part (b) would change if resale were prevented.
6. Living wage campaigns are an effort to make sure that all jobs pay enough so that workers can afford a basic standard of living if they work full-time.
- a. Draw a labor supply and demand curve and show the price and quantity in the free market. Now do a complete welfare analysis (consumer surplus, producer surplus, etc.)
 - b. Now show how your results would be affected in a living wage law were passed. Specify how jobs are allocated in your analysis? Justify this as the most reasonable allocation rule for real world situations.
 - c. Is it possible some workers will benefit in some circumstances? Is it possible that all workers will benefit in some circumstances? Is it possible that no workers will benefit in some circumstances? In each case, specify what allocation rule is required for your conclusion.
7. Consider a small country that imports all of its computers. These computers are bought on a competitive world market. Suppose the county is considering imposing a \$100 per computer import duty.
- a. In a picture, show the price, quantity, and consumer surplus both before and after the tariff.
 - b. Does the county as a whole benefit? Does anyone benefit?
8. The French are only producers of Beaujolais wine. This is shipped around the world on the day it released with great fanfare. What is little known is that this is a tremendous joke France is playing on the rest of the world. They won't drink this swill in France at any price or under any circumstances, preferring instead wine from nobler vines. One day, the minister of wine gets an idea to make the joke even funnier. He proposes and an export fee! The idea is to add a fee of one Euro to every bottle shipped overseas to be collected by the central government.
- a. Assume that the international demand for Beaujolais is downward sloping, and the domestic supply upward sloping. Do the welfare analysis for the free market case and export fee case. Of course one is concerned only with the welfare of the French!
 - b. Are French Beaujolais producers better off or worse off? What about French consumers? Is France as a whole better off or worse off? What does it depend on?

Chapter 8. Market Power

Section 8.1. Price Taking and Making

So far we have considered only *competitive* firms and consumers, that is, agents who are *price takers*. In this chapter we begin to consider agents with *market power*, that is, who are **price makers**. Our focus is on the more empirically relevant case of firms with market power. It is worth noting that it is theoretically possible that consumers might have market power as well. In most of these cases, however, such market power is exercised by firms demanding intermediate inputs from other firms and rather than by consumers of end products.

We use the following terms to describe different levels of market power for firms:

Monopoly: A market in which a single firm supplies all the output to the market.

Duopoly: A market in which two firms supply all the output to the market.

Oligopoly: A market in which a small number of firms supply all the output to the market.

In all of these cases, the firms individually face downward sloping demand curves rather than the flat demand curves faced by each firm in a competitive market. This means that the price a firm receives for its output will change in response to how much it chooses to produce. Firms are very creative in how they exploit this trade-off in their efforts to maximize profit, as we shall see.

If market power happens to be on the consumer side we have:

Monopsony: A market in which a single agent consumes all the output.

Logically, there should also be duoposony and oligopsony, but these do not seem to be very much studied by economists.

Section 8.2. Monopoly Pricing

We showed in the previous chapters that maximizing profits for any firm (regardless of competitive or noncompetitive nature of its market) requires setting $MC = MR$. For competitive firms, we also showed that the demand curves they faced were flat. This means that the price a competitive firm can get for its output is fixed and is so not a function of the quantity they choose to produce. Firms with market power, on the other hand, can affect the price they receive for their output by choosing how much to produce. This means that their demand price is a function of their output quantity.

To understand how a producer might take advantage of its ability to influence market price, we begin with the simplest case: a monopoly firm that chooses a profit maximizing quantity of output and sells to all consumers at the same price.

Suppose a monopoly faced the following downward sloping demand curve: $Q(P)$. If we solved for P in terms of Q , we would have what is called the **inverse demand curve**. This allows us to write the general form of the total revenue curve: $TR(Q) = QP(Q)$. This means that the monopoly's profit function takes the following form:

$$\pi(Q) = TR(Q) - TC(Q) = QP(Q) - TC(Q).$$

To illustrate, suppose a firm faced a linear demand curve:

$$Q = a - bP,$$

where a and b are positive constants. Then the inverse demand function is:

$$P(Q) = \frac{a - Q}{b} = \frac{a}{b} - \frac{Q}{b},$$

which implies:

$$TR = Q \left(\frac{a}{b} - \frac{Q}{b} \right) = \frac{aQ}{b} - \frac{Q^2}{b},$$

and so,

$$MR = \frac{\partial TR}{\partial Q} = \frac{a}{b} - \frac{2Q}{b}.$$

We see that the slope of the demand curve is $-b$, while the quantity intercept is a and the price intercept is a/b . The slope of the marginal revenue curve, on the other hand, is $-2b$, while the quantity intercept is $a/2$ and price intercept is a/b . The figure below illustrates this.

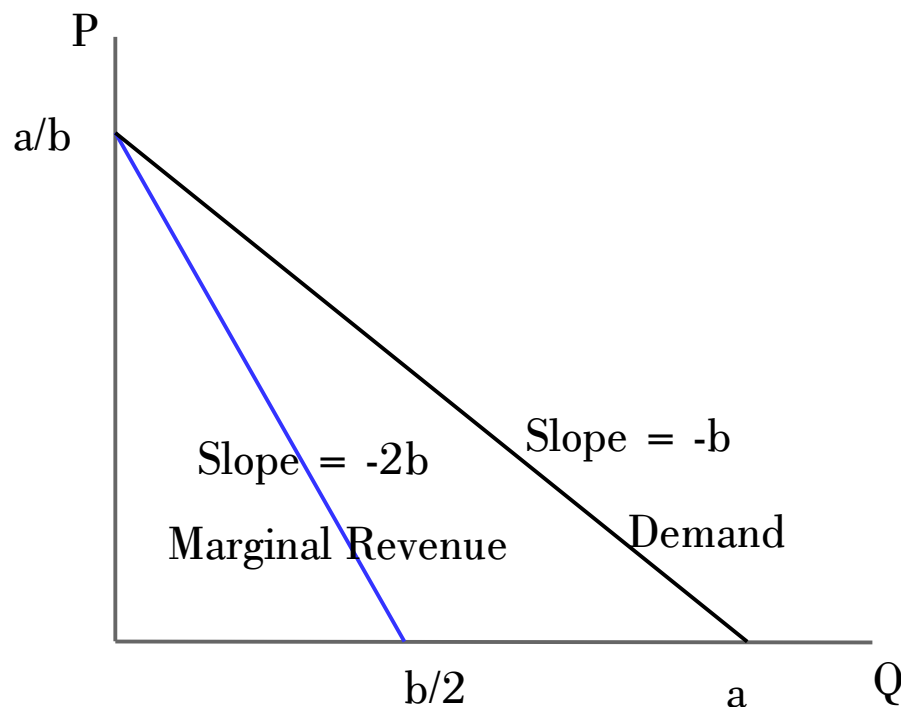


Figure 98: The demand and marginal revenue for a monopolist

The dilemma faced by a monopolist is that to sell more output, it has to lower price not only on the last item sold, *but on all the previous units of output produced as well*. For example, suppose that a monopolist is currently selling 10 units at a price of \$40 but can increase sales to 11 if he reduces the price to \$38. By doing so, it gains \$38 of revenue from selling the 11th unit, however, it is forced to reprice the first 10 units he produced at \$38. This repricing decreases revenue by \$20 = 2 × 10. Thus, the net effect of lowering the price from \$40 to \$38 is that the monopolist sells one more unit and revenue goes up in net by \$18 = 38 − 20. In other words, the marginal revenue of the 11th unit of output is \$18.

The strategy that monopolies follow is therefore to choose a quantity such that the MR of the last unit produced equals the MC , and then set the (monopoly) price at the highest level that the market will bear. In the figure below, the monopolist chooses Q_M , since this is where $MC = MR$, and then goes up to demand curve and finds that P^M is the highest price he can charge for each unit at this level of output.

Note that this means that *monopolists have no supply curve!* They set quantity based on an *interaction* between demand and the MC . We cannot even ask the question: “what would a monopolist supply at a given price” since this would be treating a monopolist as a price taker.

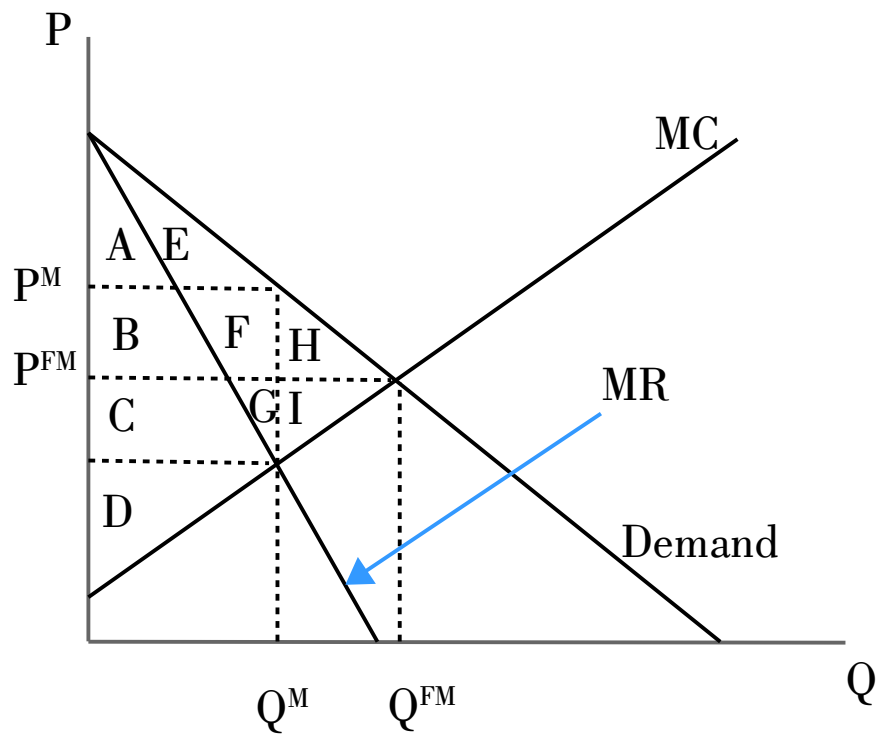


Figure 99: Graphical analysis of a monopoly market

As you can see, the monopoly raises the price to consumers above the free market level. Consumer surplus goes down, but producer surplus goes up. There is a DWL since consumers lose more than the monopolist gains.

	Free Market	Monopoly
CS	ABEFH	AE
PS	CDGI	BCDFG
SS	ABCDEFGHI	ABCDEFGF
DWL	-	HI

Table 12: Welfare analysis of a monopoly market

Section 8.3. Monopsony

If a market has only one demander, we call this consumer a **monopsonist**. The classic example is a coal mining town where a mining company is the only employer, and so, the only consumer of labor. Modern day examples might include towns with a single hospital, which is therefore a monopsonist consumer of skilled nurses, or the US government, which is the only demander of weapons-grade plutonium (we sincerely hope).

A monopsonist has the power to choose its own demand level just as a monopoly can choose its own supply level. The monopsonist holds back on the quantity it demands, and by so doing, artificially drives down the price of the goods it buys. This is just like what a monopolist does in holding back supply in an effort to drive up the price of the goods it sells. Not surprisingly, the problems are extremely similar at a formal level.

Suppose the good the monopsonist consumes is supplied by a competitive set of firms and so the industry supply curve is given by $Q^S(P)$. Solve for the **inverse supply curve** (supply price as a function of quantity: $P^S(Q)$). Then:

$$\text{Total Expenditure} = TE(Q) = Q P^S(Q).$$

is the **total expenditure** by a monopsonist that choose to demand Q units of the good. Just as monopolists do not have supply curves, monopsonists do not have demand curves. Just as monopolists know their marginal cost, monopolists know their marginal benefit. Profit maximization then requires choosing quantity demanded so that the **marginal expenditure** equals to the marginal benefit of consumption:

$$\frac{\partial TE(Q)}{\partial Q} \equiv \frac{\partial Q P^S(Q)}{\partial Q} = P^S(Q) + Q \frac{\partial P^S(Q)}{\partial Q} \equiv ME(Q) = MB(Q).$$

The figure below illustrates:

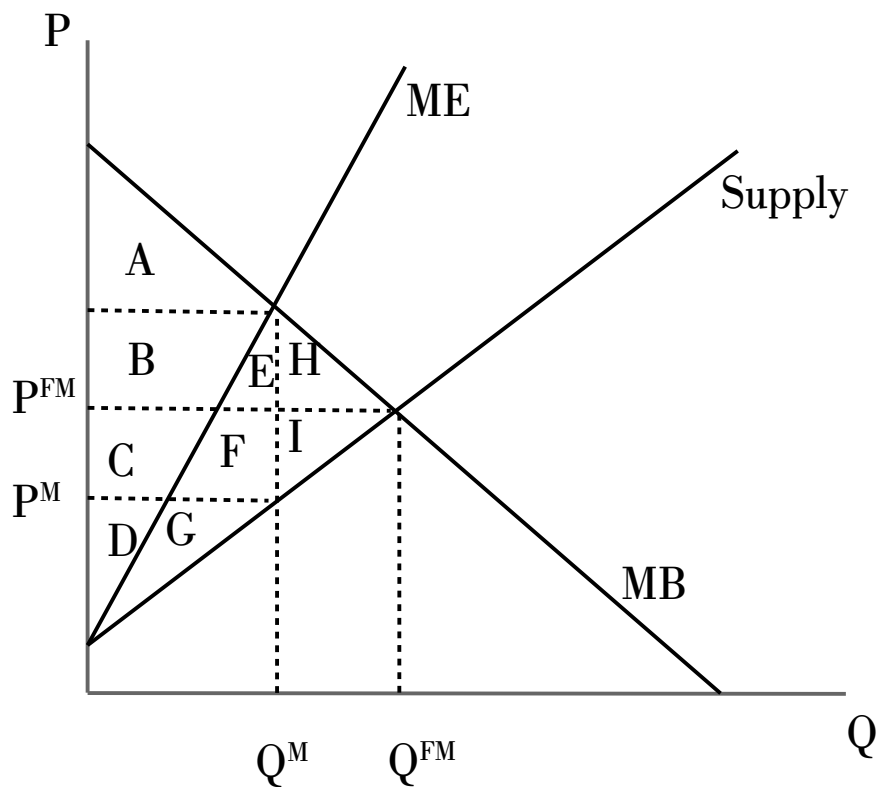


Figure 100: Graphical analysis of a monopsonized market

We leave the welfare analysis to the reader. Notice that the monopsony price (P^M) is now as low as possible at the monopsony quantity (Q^M) in contrast to the case of a monopoly.

Section 8.4. Why Do Monopolies Exist?

You might wonder why monopolies and monopsonies exist at all. After all, monopolies make profits and this should attract entrants to drive down the price. How can these economic profits persist in equilibrium? Several factors can be at work, and it turns out this affects how (or even if) we should make policy to deal with such noncompetitive markets.

There are five leading reasons that monopolies persist. We discuss them below.

Subsection 8.4.1. Natural Monopoly

Natural monopolies occur when an industry is characterized by high costs of initiating production, but lower incremental costs for producing subsequent units of output. For example, to build even a single new passenger aircraft, Boeing has to design it, test the model, get regulatory approval, build an assembly line, retool its machines, train its workers in the new process, and so on. Thus, if it made only one copy of the new design, it would be extremely expensive. The second unit, however, only requires the additional time and materials that go into actually producing the aircraft, so the incremental cost of the second unit is much smaller than the first. This means that the AC of two aircraft is smaller than the AC of one aircraft, and that AC continue to decrease for a long time (maybe forever). Railroads, electric and water utilities, even software and publishing, to some extent, are natural monopolies.

Recall from the chapter on producer theory that over the range of quantities for which the AC curve is downward sloping, the MC is below the AC . Even if the AC starts to slope upwards, it may take a while (that is to say, the firm may have to produce a very high quantity) before the MC finally penetrates the minimum of the AC curve. If demand is not large enough, then the demand curve may cross the AC before this point. Formally:

Natural monopoly: A firm is said to be a natural monopoly if the minimum of its AC curve is above the market demand curve.

The figure below illustrates this case.

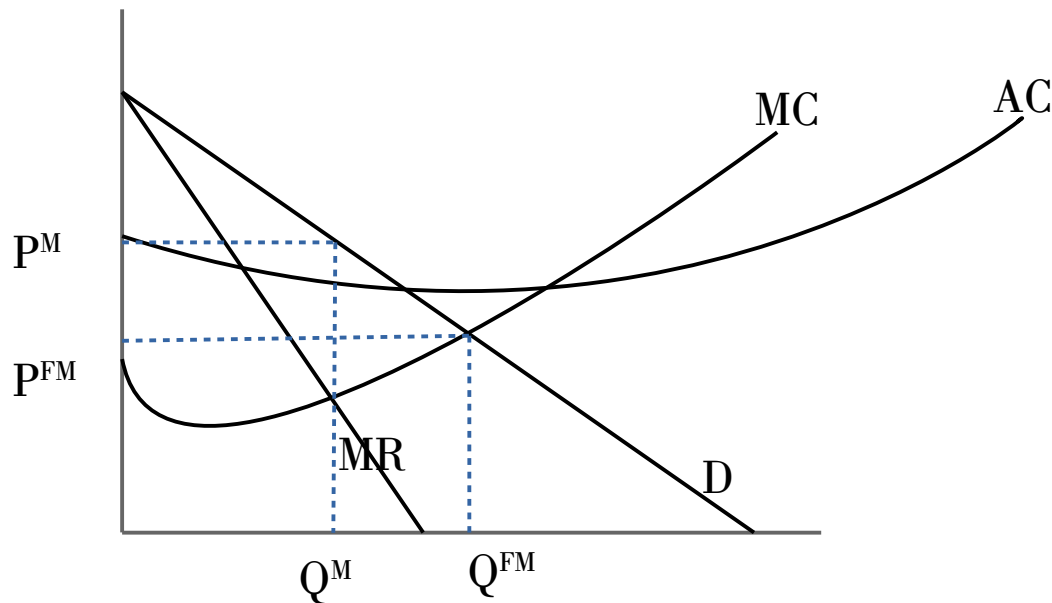


Figure 101: MC, AC and demand for a natural monopoly

The obvious problem is that it is socially optimal not to break up this kind of monopoly. We would not want two parallel sets of under-utilized railroad tracks crisscrossing the country nor to have more than one power, water, gas, or sewer, connections made by different firms to each individual house. This would require additional use of capital (rails, pipes, wires, etc.) but not change the amount of transport, water, electricity, and so on, that are actually produced and consumed. It would simply be a waste of resources that could be better used in other ways.

Entry by a competing firm is difficult in such a market. An entrant will initially be smaller than the incumbent firm and will therefore have higher average costs. To make matters worse, it is extremely difficult for two firms to exist in such a market while making positive profits. Doing so would require a collusive agreement to jointly restrict output. The sum of the profits each firm would make would be less than what the natural monopolists would have made on its own (since AC for each of the two firms splitting output is higher than it would be for one firm producing all of the output).

In some cases, it may be possible to decouple the natural monopoly part from the competitive part of an industry. For example, we might have power transmission lines owned by the public or a regulated monopoly, but allow private competitive firms to supply electricity to the grid. Similarly, we might have the network of rails and stations be publicly owned, but allow private companies to run trains on this network on a competitive basis.

Subsection 8.4.2. Network Externalities

Suppose it is 1876 and Alexander Graham Bell knocks on your door and tries to sell you his amazing new communication device called the “telephone”. You ask him who you would be able reach with this device, and he responds that you are the first door he has knocked on. At this point, the network would contain only you and so the telephone is useless.

If several of the local businesses had already signed up, the phone would be useful. If the whole town had signed up, it would be even more useful. If the whole country or world joined the network, the value would go up even more. Other examples include software, social media and gaming, and technological standards like MP3, PDF, HTTP, and the metric system. More formally:

Network Externality: A situation in which the value of a product to each individual consumer increases with the number of other consumers who buy the good (and thus, join the network).

In many cases, beneficial network externalities continue to increase regardless of the size of the network. When network externalities are extensive in this way, the efficient thing to do is to have only one product in any given market. For example, putting half of the country on a different telephone network or requiring that for every Facebook user, another person would have to join MySpace, only makes communication more difficult and lowers social welfare.

Similarly, if everyone else uses Windows 10 or Microsoft Office, it is better for me to follow suit since (a) there are likely to be more applications available, (b) it is easier to get help and support since more people know how the system works, and (c) it improves my job prospects since employers are likely to require me to use this system. Even if Linux and Apple OS are better in every possible way, the network externality may overwhelm their advantages.

Entry into such a market is difficult. Why would a consumer want to choose a new product with a small network when there is already an incumbent offering a product with a large network? This is especially true for Information and Communications Technology (ICT) products in which the price to the consumer is zero and the firm makes its money by advertising, offering premium subscriptions, or in-app purchases.

If Facebook allows users join for free, any entrant to the market would somehow have to offer a negative price to make it attractive for users to join his smaller network. To add to the difficulty, ICT products are often natural monopolies as well. The cost of writing software, developing platforms, setting up servers, and so on might be large, but the marginal cost of letting one more user enjoy the software or social platform is typically very small. Thus, the AC of allowing new users is downward sloping even as the user base gets very large.

There is still hope for some degree of competition in such markets. For example, if two firms enter a market at close to the same time, both have a strong incentive to grow as quickly as possible at the expense of the other. If either gains a substantial lead, then the network externalities can snow-

ball and cause users of the smaller competing network to switch to the leader. Knowing this, such firms may offer consumers low prices, higher quality products, better support, and other incentives to join. In the build up phase, consumers can be big winners.

Unfortunately, if one firm becomes dominant, then it can use the strength of the network externality to charge higher prices and exploit its users. If you do not like Facebook's privacy policies, its attempts to market to you, or the way it manipulates your feeds, you really have no alternative but to become a social pariah or a hermit. If you think Windows 10 sends too much telemetry back to Microsoft, runs too many background processes without your knowledge or permission, or steers you too much into Microsoft's ecosystem and to its preferred vendors, you can always become one of those weird Linux guys (also social pariahs and hermits).

Even when network firms compete, we are likely to end up with one winner in the end. This might be due to one firm running out of capital or being acquired by the other. An offsetting factor is that technology moves so quickly now, that it by the time that the fight for dominance is over, the product category may be approaching obsolescence.

It is also possible that network externalities are more limited and may even reverse. Consider the following:

Diminishing returns to network size: As networks grow larger, it may be that each additional user provides less value on the margin to other users. For example, think about your network of friends. The first two or three friends were very important. As your friend group grows to ten, twenty, one hundred people, each additional person provides less benefit to you and other members of the group. You simply do not have the time or attention span to get much value from so many people. It might even be that having more friends becomes negative at some point as you find yourself posting happy birthday wishes to so many Facebook pages, and feeling obliged to respond to so many Tweets and Snapchats. Consider SoundCloud, for example. Having 20 or 30 people who share your tastes to suggest new music you might like is great. But how much better would it be if there were 1000 people who all like the same music as you? Surely 20 people can suggest more music than you could possibly listen to anyway. Wading through so many suggestions might even be a negative rather than a benefit.

Differences in tastes: Some people like Apple and others like Windows. Some like to Tweet and others like Vine or Snapchat. In other words, the features of the product are important to many people. While larger networks are better, they may be outweighed by other aspects of the platform or good. In this case, we might find that there is room for two competing network products in a market. This would require both that tastes over features are strong enough and the network externalities diminish quickly enough.

Local monopoly?

Subsection 8.4.3. Patents, Copyrights, and Trademarks

The US Constitution grants the federal government the authority:

To promote the progress of science and useful arts, by securing for limited times to authors and inventors the exclusive right to their respective writings and discoveries.

Patents and Copyrights are monopolies that are established and enforced by law. At the nation's founding, copyrights and patents lasted for a term of 14 years. Copyrights could also be renewed once by a living author for another 14 years. Currently, patents last for either 14 or 20 years and cannot be renewed. Trademarks last 10 years, but can be renewed indefinitely. Copyright law is a bit more complicated. Anything published before 1922 is in the public domain. If a work was published between 1922 and 1978, copyright lasts 95 years from the date of publication. If a work was published after 1978, copyright lasts for the lifetime of the author plus another 70 years.

Society benefits when new ideas and works are created, and when new products are brought to market. Copyrights and patents are meant to allow creators to profit from their efforts by granting them a temporary monopoly. This provides an incentive to innovate and so this kind of monopoly serves a useful purpose. On the other hand, the cost is that for 14, 20, 95 or even more years, society suffers deadweight losses in these monopolized markets.

A balance must be struck, but doing so in practice is difficult for several reasons:

One size does not fit all: It takes a lot more potential profit incentive to get George Lucas to produce Star Wars than it does to get a lonely teenager to write a poem. Some innovations are easy and would take place with little or no reward. Other innovations are very difficult and certain areas of research may be too speculative to payoff even with a 20-year patent.

How broad should the patent be?: Suppose I invented the blow-dryer. Should I get a patent on all hot air drying technologies, or just the specific blow-drying machine I developed? Too broad a patent, and you foreclose innovation and over-compensate inventors. Too narrow a patent and you discourage innovation. For example, what if my competitor could get around my patent by producing essentially the same machine as do but in blue plastic instead of the green I chose? He would get the benefits of my investment in research and development without paying anything for it. Knowing this and that I will face price competition when I start producing blow dryers, I may find that innovating is not worth the effort.

Technical competence: It is difficult for the patent office and judges involved in intellectual property cases to determine what constitutes “prior art”, what is “useful” and what is “nonobvious”. To the extent that this is true, patents can be granted for well known ideas. Not only does the patent system fail to cause the creation of new knowledge in this case, but it also locks old knowledge up in a monopoly for many years. This slows down both economic and scientific advance. On the other hand, the chance that a sensible patent might be rejected by an incompetent patent examiner may discourage some innovators from investing effort in risky but useful technologies.

Subsection 8.4.4. Legal Barriers to Entry

Sometimes the state awards a monopoly by making it illegal for more than one favored firm to enter an industry. Examples include: local cable service, the Hudson's Bay Company, the American Medical Association, and the US Postal Service. Patents and copyrights are a special case of this more general class.

These monopolies are granted by state, local, and national governments to further some public purpose. Consider the case of broadband Internet providers. Laying cable and fiber to residences is capital intensive and also requires digging up streets and upgrading telephone poles. Cable and internet providers must get permission from state or local authorities to undertake these activities, and they commonly argue that they should be granted monopoly rights for some period of time to give them an incentive to invest the required capital.

This last claim is questionable, nevertheless, the state and local authorities in charge have generally agreed to the arrangement. At the same time, these authorities have usually required that local programming (like city council meetings) be carried, that cable or fiber be run even to isolated houses, connections be provided to schools at discounted prices, or that some of the potential profits be spent on other things that the decision makers happen to favor. Sometimes bribes, campaign contributions, and promises of future jobs or current employment to family members of politicians have been offered as well.

The argument that granting these legal monopolies has served to get broadband service deployed more quickly and has produced other public benefits may or may not be true. That the process of granting these monopolies is inherently corrupting to all involved and does not foster innovation, efficiency, or provide incentives for the monopoly to woo customers with good service, however, is certainly true.

Another example is the Hudson's Bay Company which had an exclusive right to trade certain goods with Canada from the time of its founding in 1670 through part of the 19th century. In exchange, they provided troops and other government services to maintain the colony at no expense to the British Crown. This arrangement benefited the Crown as well as the owners of the company. Whether this was a better arrangement than levying taxes and tariffs and using troops from the British army is an open question.

A more modern example is the American Medical Association (AMA), a private professional association that has been granted exclusive rights by the US government to certify doctors and medical schools. It may not surprise you to learn that the AMA contributes millions of dollars to various political campaigns each year. The AMA makes it very difficult and time-consuming to gain certification and new doctors are required to endure long periods of poorly paid internship. In addition, very few new medical schools are allowed to open.

On the one hand, a high bar generates high quality doctors. On the other hand, the high bar also creates less competition for existing doctors, raises their wages, and reduces access to doctors for

everyone. At the very least, the AMA has a conflict of interest. At worst, it is engaging in **rent seeking** by lobbying and donating to congress on behalf of its members.

Subsection 8.4.5. Natural Resource Monopoly

If a country, person, or firm owns all of a particular item that exists in the world, we say it has a **natural resource monopoly**. For example, Alcoa (the Aluminum Company of America) once bought the rights to most of the easily available aluminum ore in the world, OPEC controls a significant part of the world's readily accessible oil, De Beers produces most of the world's gem quality diamonds and buys up most of what other countries and companies dig up. Picasso was the only person in the world who could create a Picasso.

These are the classic bad monopolies. There is no offsetting social gain. It is good if you are the seller of monopolized natural resources, but bad if you are a consumer.

Section 8.5. Government Solutions to Monopoly

We see that monopolies often exist for good economic reasons. In many cases, it is impossible or undesirable to introduce competition into such markets. Nevertheless, we would like to reduce the *DWL* that high monopoly prices cause even if we do not wish to get rid of the monopoly itself and its associated benefits. What strategies might the society or the government employ to get back to zero *DWL*? There are three obvious possibilities.

Subsection 8.5.1. Subsidizing Monopoly

By giving a monopolist a production subsidy the government can lower its *MC* curve. If it reduces *MC* far enough, the subsidized *MC* curve can be made to intersect the *MR* curve at the free market quantity rather than the old monopoly quantity. It is easy to verify that this results in the consumers paying the free market price for the free market quantity, and that consumers are just as well off as if the firm had acted as a competitor. The monopoly gets very large producer surplus at the expense of the government and whoever pays the taxes to support it. This might cause significant *DWL* in other markets where taxes are increased in order to pay for the subsidy, but the *DWL* in the monopoly market would become zero. We would get an analogous outcome if we gave monopolists a consumption subsidy.

There are two major problems with this approach. The first is political. Who wants to give Bill Gates a subsidy to get him to lower the price of Windows? He has enough money already! Even though the subsidy might improve social welfare, giving money to large exploiting corporations is likely to be very unpopular. The second has to do with measurement. In order to set the correct subsidy, we would have to know the shape of the entire *MC* and *MR* curves. This is very unlikely, and thus, the correct level of subsidy is difficult to figure out.

The figure below illustrates this.

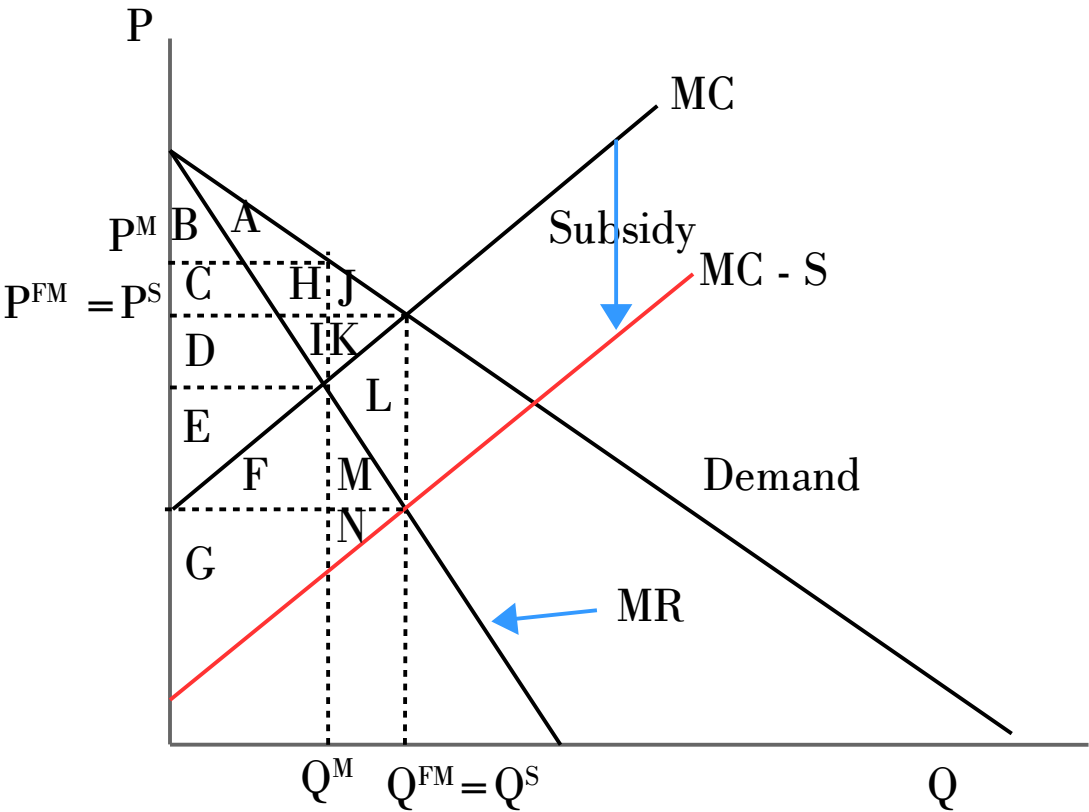


Figure 102: Graphical analysis of a monopoly with a subsidy

As you can see, the subsidy causes the monopoly to lower the price back to the competitive level. Although this gets rid of the *DWL*, the government must come up with money to pay for the subsidy’s cost, which is therefore subtracted from the social surplus.

	FM	Monopoly	Subsidy
CS	ABCHJ	AB	ABCHJ
PS	DEIK	CDEHI	DEFGIKLMN
Subsidy Cost	–	–	DEFIKLM
SS	ABCDEHIJK	ABCDEHI	ABCHJGN = ABCEHIJK
DWL	–	JK	–

Table 13: Welfare analysis of a monopoly with a subsidy

Subsection 8.5.2. Price Ceilings

An alternative approach to addressing the market inefficiency of monopoly is to impose a price ceiling equal to the free market price P^{FM} . Although the demand curve is above P^{FM} up to the free market quantity, Q^{FM} , the monopoly is not allowed to charge anything above P^{FM} . This implies that up to quantity Q^{FM} , it charges P^{FM} for each unit, and this in turn implies that $MR = P^{FM}$ up to quantity Q^{FM} . From there, it continues down to the quantity axis at a slope equal to twice that of the demand curve.⁵

This implies that the profit maximizing intersection of MR and MC takes place at point (Q^{FM}, P^{FM}) . As a result, it is optimal for the monopoly to act exactly like a competitive firm. There is no DWL , and the consumer and producer surpluses return to free market levels. The figure below illustrates this. (In the case of monopsony, a price floor equal to P^{FM} gets us to a similar zero DWL equilibrium.)

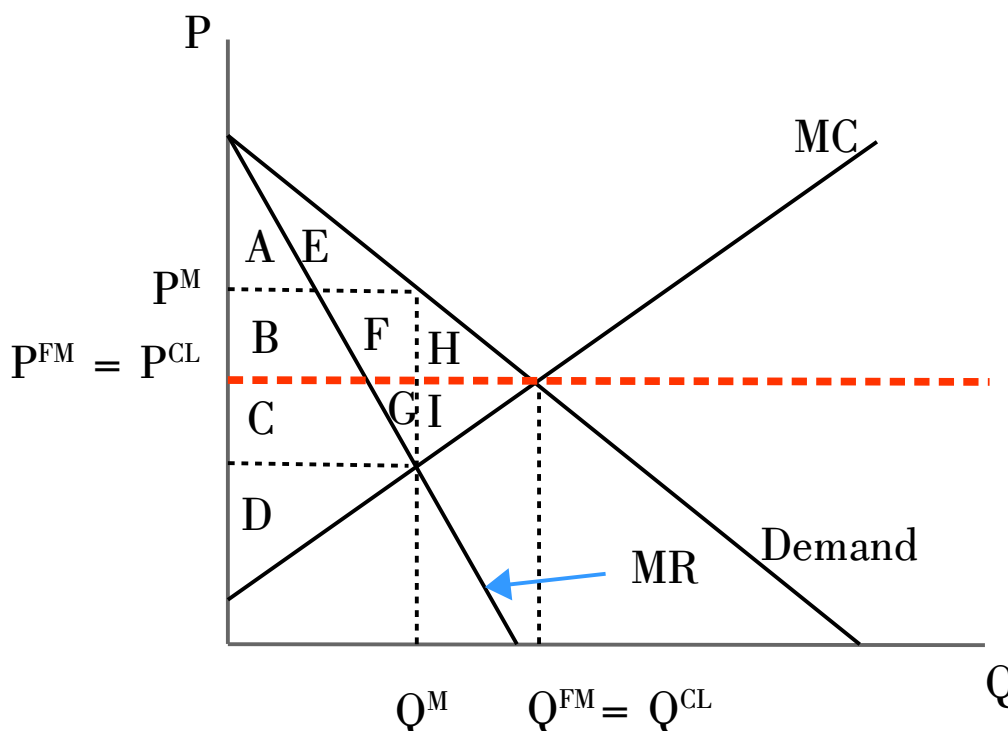


Figure 103: Graphical analysis of a monopoly with a price ceiling

The details of the welfare analysis are left to the reader.

⁵ You can also verify that the MR curve at quantities greater than Q^{FM} starts out at point (Q^{FM}, P^{FM}) .

The major problem with this solution is measurement (again). To find the correct price ceiling we would have to know the shape of the MC and the demand curve. Otherwise, we could not know where they intersect and thus, how to determine the free market price. If we set the price ceiling above free market price but below the monopoly price, consumers benefit, the firm is hurt, and the DWL gets smaller. Errors like this move us closer to the zero DWL equilibrium, but not as close as they could be. Suppose, on the other hand, we happen to set the price ceiling below free market price. The monopoly would choose to supply out to the point where this price ceiling hits its MC curve, while the consumers will demand the higher quantity at which the price ceiling hits the demand curve. This excess demand creates a shortage and may cause an even larger DWL than the unregulated monopoly, depending upon how the shortage is resolved.

Subsection 8.5.3. Rate of Return Regulation

The idea behind Rate of Return Regulation (RRR) is the observation that there are no profits in a free entry competitive equilibrium. Since monopolies raise prices in order to make positive profits, all we should have to do is to force a monopoly to behave like competitive firms. In particular, we could have regulators force monopolies to make zero profits by setting prices such that $P = AC$ like we see in a free entry competitive equilibrium.

Rate of Return Regulation: A regulatory system in which a firm is forced to set price equal to its average cost once the opportunity cost of any capital owned by the firm is calculated at a rate of return determined by the regulatory body.

This approach is frequently used to regulate some government sanctioned monopolies such as water, gas and electric utilities. We consider two cases:

Standard case: The intersection of the MC and AC for the monopoly firm is below the demand curve.

Natural monopoly case: The intersection of the MC and AC for the monopoly firm is above the demand curve.

Standard Case

As before, the profit maximizing monopoly quantity is found by setting $MC = MR$. The free market quantity is found by setting $MC = D$. The RRR quantity is found by setting $AC = D$. Note that $Q_M < Q_{FM} < Q_{RRR}$.

You can see that consumers most prefer the RRR outcome, while the firm most prefers the monopoly outcome. The DWL is zero at the free market price and quantity. The red vertically hatched area is the DWL under monopoly pricing, and the black horizontally hatched area is the DWL under RRR. In the figure the monopoly DWL looks larger, but in general, the reverse could be true. Under monopoly, too little is produced in the sense that $MB > MC$ for the last unit produced. Under RRR, too much is produced in the sense that $MB < MC$ for the last unit produced.

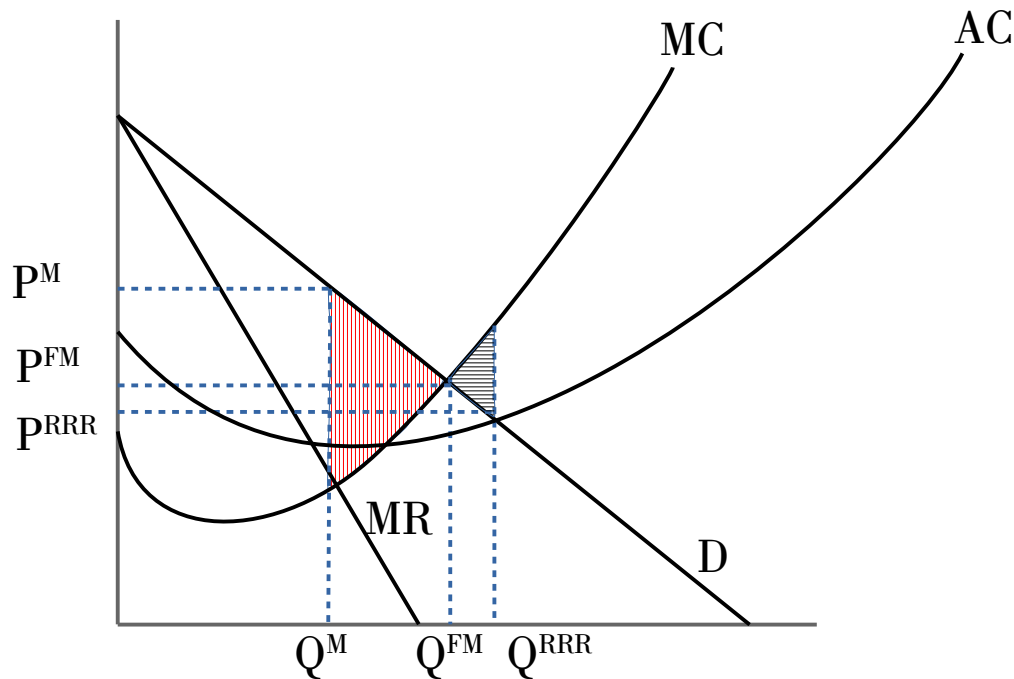


Figure 104: Graphical analysis of RRR for a standard monopoly

Natural Monopoly Case

In contrast to the case above, we now find $Q^M < Q^{RRR} < Q^{FM}$. Thus, the RRR approach lowers DWL as compared to the alternative of leaving the monopoly to set prices on its own. Of course, consumers benefit, and the monopoly is hurt by RRR, but DWL is still positive. Thus, RRR will reduce DWL , but never to zero. Also, note that if we tried to force a natural monopolist to sell the free market quantity at the free market price, the firm would have negative profits since the AC is above P^{FM} at Q^{FM} .

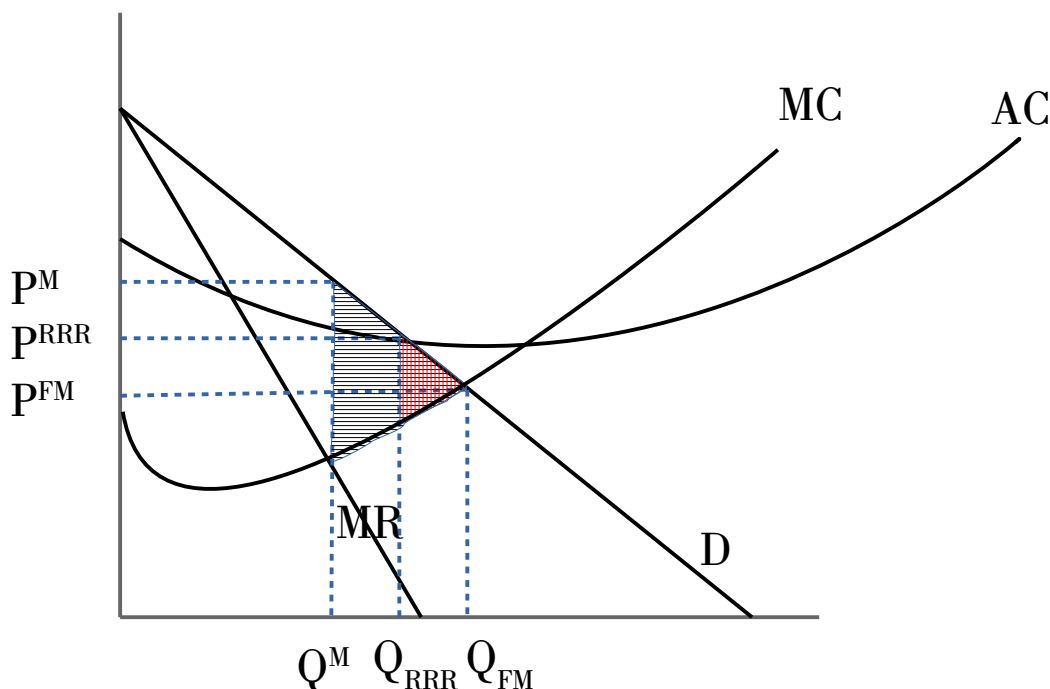


Figure 105: Graphical analysis of RRR for a natural monopoly

One of the great attractions of RRR is that it does not suffer from a measurement problem. Finding MC and demand over quantities that have never been seen in equilibrium is speculative at best. Finding the AC , however, only requires knowing accounting data like TC and Q , things that should be readily available in the annual report of a publicly traded company.

The major problem with RRR is that monopolies can game this regulatory system. Recall that “zero profits” means “zero economic profits”. Firms traded in the stock market (especially public utility companies) own large amounts of capital outright. The services of such capital are therefore not a direct cost to the firm and so public utility commissions are tasked with finding a fair “rate of return” to compensate stockholders for their investment.

It turns out that this rate of return is generally set above the interest rate at which firms can borrow from bond markets for various reasons (such as **regulatory capture**). Thus, a firm could borrow a million dollars at, say, 5%, from the bond market, use it to build a giant concrete statue of the CEO (a use that generates no revenue and has no productive value), but then claim the higher regulated rate of return (say 7%) as the opportunity cost of this capital investment. This effectively “launders” \$20,000 of the unrealized potential monopoly profits which can then be returned as extra profit to the shareholders.

Doing this causes the AC to increase. We conclude that to maximize returns to shareholders, a firm should continue to build CEO statues until the AC has gone up so much that it intersects the demand at the monopoly quantity. Thus, $AC = MC = D = P^{newRRR}$ after the firm responds optimally (from a selfish standpoint) to the RRR system. To put this another way, in the case of natural monopoly, RRR regulation is a practical solution that improves matters only if we assume that costs are exogenous and that firms continue to minimize costs in face of regulation.

The example above is a bit fanciful, but in real life firms may choose to artificially raise costs by building fancy executive offices in nice locations, buying corporate jets, or even by paying managers or union employees more than is required. After all, rate-payers will pick up all these costs under RRR, so why fight a union or be uncomfortable? This means that in practice, RRR can be even worse than the monopoly outcome. The quantity and price end up being the same under both monopoly and RRR, but what once were monopoly profits returned in full to shareholders are now partly wasted on excessive production costs and nonproductive capital expenditures. The figure below illustrates:

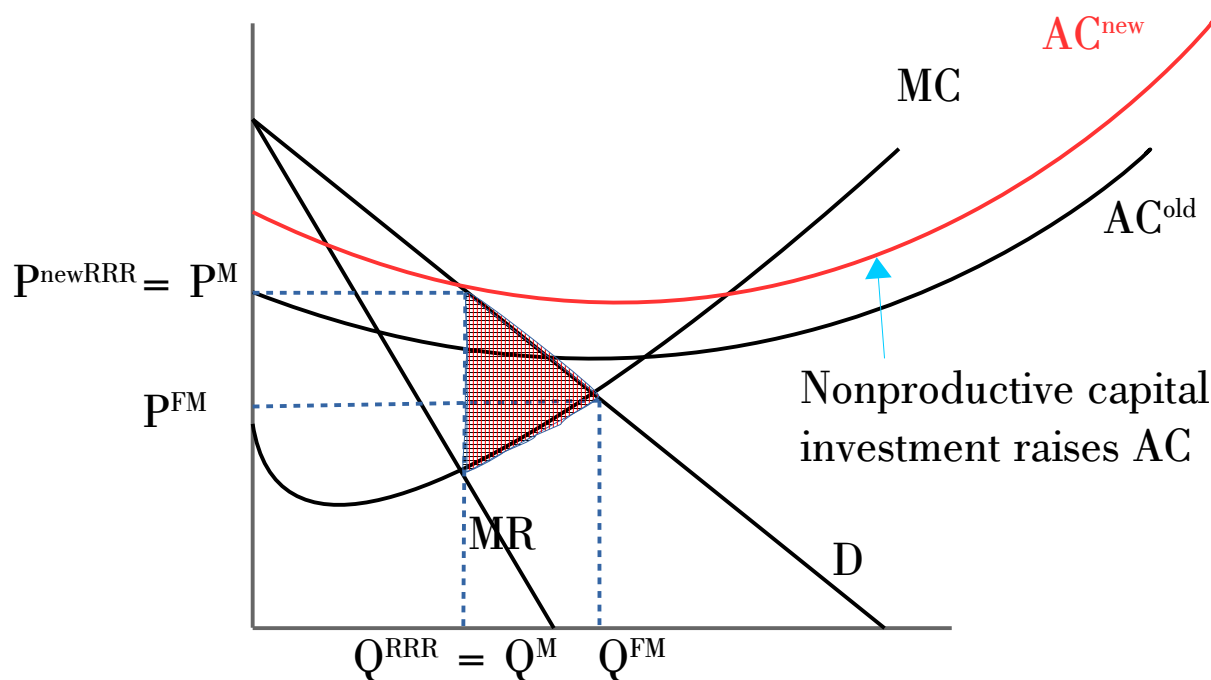


Figure 106: Graphical analysis of RRR with profit maximizing over capitalization

Subsection 8.5.4. Sophisticated Monopoly Pricing Strategies

So far we have considered only the most simple use of monopoly power. In particular, we have assumed that the monopolist sets a single price and sells at this price to everyone. In this section we explore more complicated pricing strategies that might generate even larger profits for a monopoly.

Subsection 8.5.5. Price Discrimination

What if the monopolist could sell to different agents at different prices? Suppose, for example, I owned the only two cars in town worth \$500 on the outside market and I wanted to sell them to two people, one who were willing to pay \$2000, and one who would pay at most \$1500. What do I need in order to be able to charge these consumers different prices?

Market Power: If I try to charge above MC and there are other competitive firms that can enter the market (or who already exist), they can undercut my price and still make profits. Thus, I must have market power if I hope to charge anything besides MC .

Identification: If I cannot tell who the high demand agent is, both will claim to be the low demand agent. Thus, I can either sell one car at \$2000 or two cars at \$1500. If I can tell which agent is which, on the other hand, I can just set the price at their individual reservation levels and tell them to take it or leave it.

No Resale: Even if I can identify an agent's willingness to pay, if the low demand agent can walk in and get the low price and then turn around and sell to the high demand agent, then I will not be able to take advantage of the high demand agent.

If all these conditions are satisfied, exactly how a firm maximizes profit usually depends on the strategy of identification. There are three broad categories of price discrimination. Bear in mind that most of the real world falls somewhere in between these categories.

Subsection 8.5.6. First Degree Price Discrimination

First Degree Price Discrimination: A marketing strategy in which every consumer is charged exactly his reservation price.

The good news is that if a firm can perfectly identify agents' willingnesses to pay and therefore engages in perfect price discrimination by charging each agent exactly his reservation price, there is no DWL . The firm has an incentive to sell goods as long as there is any agent with a MB above the firm's MC . As a result, the firm produces the free market quantity. The bad news is that the firm takes *all* of the surplus, and the consumer surplus is zero. Consumers as a group are worse off than if they faced a single monopoly price.

Examples

Sliding scale fees depend on income: Income is a good proxy for willingness to pay. If a doctor, for example, can accurately determine a patient's income and then charge a price below the patient's estimated reservation price but above the marginal cost of seeing the patient, the doctor benefits on the margin. This is true even if the price he gives the patient is below the average cost (which includes the fixed cost of his overhead). Of course, he must make this up by finding enough wealthier patients who can be charged above average cost to make positive profits overall.

College education with financial aid: Parents hoping to get financial aid to help pay for college have to fill out forms and include recent tax returns. Universities are therefore able to verify the resources available to the parents and form an estimate of their willingness to pay for college. The marginal cost of adding another student to a class already on the schedule is small. The average cost, on the other hand is high. After all, we have to support at least two administrators for every professor, a private aircraft for the Dean, not to mention feeding the gargoyles and rotating the squirrels. The same principle as above applies.

Big data and analytics: Every time you use the Internet, you leave a trail of information about yourself. Your browser stores cookies that allow Google, Microsoft, and others to see what sites you have visited and, sometimes, what you have done there. If you log in to your Chromebook, iPad, or Windows 10 machine, you are sending all kinds of data about your browsing and other habits to the respective mother ships. Creepy, but why should you care? Well, Amazon, among others, employs a team of economists to analyze the browsing, purchasing, and other behaviors of its users. It experiments with different prices for goods in order to determine what someone with a given data profile would be willing to pay. Putting together thousands of experiments with data from millions of users and billions of interactions, Amazon is working hard to figure out your personal reservation price for various goods as well as how presenting and steering your search results might affect your eventual purchases. Users going to Amazon from different zip codes and countries, with different collections of cookies, get different offers, prices, and experiences. Even worse, once you log in, the data set that allows Amazon to do these things is enriched by the entire record of everything you have ever done on the site. On behalf of economists, at least, I apologize for this.

Subsection 8.5.7. Second Degree Price Discrimination

Second Degree Price Discrimination: A marketing strategy in which different units of the good are priced differently.

Such pricing schemes may take the form of quantity discounts or what is called “versioning” of products. Both of these are examples of a more general strategy called **bundling**.

Bundling: A marketing strategy in which specific amounts of one or more goods are sold in a package at a fixed price instead of individually at per unit prices.

Versioning: A marketing strategy in which a firm offers related products with different bundles of features.

The point of versioning is to extract the highest possible price from high demand consumers while still being able to sell goods to lower demand consumers at a lower price. For example, a firm might offer a cell phone without a camera for a low price, but also a higher priced version with a camera for the premium market. Low demand and high demand agents identify themselves by the buying choices they make. Resale is irrelevant since no high demand consumer wants the cheaper phone.

Quantity Discounts: A marketing strategy in which a firm charges high per unit prices to consumers who wish to buy small amounts of a good, but lower per unit prices to consumers who are willing to buy large amounts.

For example, you can buy a can of Coke from a vending machine for \$2. You can also buy a case of Coke at the grocery store with 24 cans for \$6. Clearly, the per-unit price is much lower in the store, but you have to buy bundles of 24 cans at once to take advantage of this. People who use vending machines are thirsty and lacked the foresight or energy to bring a Coke with them. People who buy cases typically have cars, can shop around for the best price, and have many alternative drinks right in front of them to choose from when they buy. Thus, we have people with both high and low willingness to pay for a can of Coke.

If Coke tried to charge \$2 a can in the grocery store, many people would buy Pepsi, the store brand of soda, or some other drink. Coke would lose these sales. There is no need to offer discounts at a vending machine if the day is hot and there is no other vending machine in sight. Those who came unprepared have a high willingness to pay. Again, people identify themselves by their buying choices and resale is generally impractical, although some leakage may occur.

Although second degree price discrimination is a very common strategy for firms with market power, it is difficult for firms to work out exactly what the profit maximizing bundles are. Each bundle competes with the others and so there is always some seepage across bundles of consumers with different willingnesses to pay. Without knowing exactly how many consumers are willing to pay how much for each possible quantity or version, a firm can not fully maximize profits. From a welfare standpoint, some consumers win, and others lose. The total welfare effect depends on how the bundles are set up and how consumer demand is distributed. There is no general result here.

Examples

Bulk packages of food and other groceries: Suburbanites with large families, larger houses, and access to transportation are both willing and able to shop around for good values, go to more remote locations, and store quantities of a good for future use. Hipsters who live in apartments with roommates may find it difficult to buy large quantities of a good since storage is limited and their rate of consumption is relatively low. If they do not have easy access to a car, they must purchase from nearby convenience stores, farmers' markets, and bodegas. Therefore, suburbanites identify themselves as price sensitive shoppers both by their willingness to drive out to Costco, and to buy twelve packs of pineapples. Hipsters are willing to do neither of these and so the small number of stores in their neighborhoods have market power and can sell small packages of the various goods at high per unit prices (can we say that the Hot Pockets and Fritos they buy are at least locally sourced?). Resale does not take place because hipsters want nothing to do with suburbanites and inversely.

Building supplies: Ace Hardware and Home Depot stock a wide variety of things like glue, paint, nails, and plumbing fixtures in packages in small quantities. These are generally purchased by homeowners and cost a lot more per unit than the large packages purchased by builders. Builders also get bulk discounts of large purchases of lumber, dry wall, and so on. Builders are both willing to shop around and can quickly use up large quantities of supplies. It would not make sense for a homeowner to drive to four of five places in hopes of saving a few cents on the package of 50 nails he wants (of which he only plans to use six.)

Group discounts for movies, shows, or admissions: Groups that go to events together are often as interested in one another's company and the event. Many events might serve as an equally good place to enjoy one another's fellowship. Purchasers of single or small numbers of tickets are more likely to be interested in the event itself. Thus, groups are more price sensitive than individuals and identify themselves with their willingness to purchase many tickets at once.

Subsection 8.5.8. Third Degree Price Discrimination

Third degree price discrimination: A marketing strategy in which a firm charges different prices for the same good in different markets.

Firms can benefit if they can identify markets or market segments with different elasticities of demand. Each market is offered a different price and does not have access to prices offered in other markets. Preventing resale is sometimes easy but is often very tricky.

Examples:

Bargain matinées, child and senior discounts at movies: Junior and senior movie goers generally have less money than working age people and so a lower willingness to pay. The fact that you are willing to go to a bargain matinée means that you are probably a loser trying to drown your sorrow in butter and Jujubes. You cannot take a date to a movie at 2:00pm on a Tuesday, it has to be Friday, Saturday or at least in the evening. Thus, there is a group of people (the dateable) who are willing to pay a premium for a movie at prime time. Enjoy your sad, expensive popcorn.

Airline tickets: Travelers pay different prices for equivalent seats depending on how close it is to flight time, if there is an overnight or weekend stay, how many stops are made, and so on. Business and personal travelers have different willingness to pay. Business travelers are using someone else's money and have little incentive to be price sensitive, plan in advance, or take inconvenient flights. They end up choosing tickets with higher prices as a result.

Textbooks: Prices for textbooks are different in different countries. The willingness to pay for a textbook in a poor country like India is much lower than it is in the US. The demand curves are quite different. Thus, publishers often put out an Indian edition that sells for a fraction of price of the US version. This allows them to sell to and profit from both markets. What prevents resale? The cost of shipping a book from India to the US is a deterrent as well as the fact that the overseas editions are often printed in paperback and on cheaper paper. Publishers have also made the claim that such books are sold under the condition that it only be used in the country they were intended for, and export out of the country is illegal. Fortunately, this claim turns out to be a violation of a legal principle in copyright law called the "first sale doctrine". This doctrine holds that when you purchase a physical object, you can sell it, rent it, lend it, destroy it, and do anything you wish to it, besides copying it. As a result, a secondary market has arisen in which entrepreneurs in India and other countries buy local editions in bulk and ship them by container back to the US, Europe, and other more expensive markets.

Academic journals: Publishers charge libraries much higher prices for scientific journals than they do individual subscribers. Sometimes a journal costs a library \$40,000 or more per year. Different countries and types of universities also get charged different prices. You can work out for yourself how these markets differ, and how resale is prevented.

CDs and DVDs: Each disk has an embedded code which makes it unplayable on a machine manufactured for sale in a different region. A European CD will not play on a Japanese machine, for example (unless it was made and encoded for sale in Europe). Similarly, some digital content purchases at Amazon or through iTunes cannot be accessed if you travel to certain locations. Some of the content you have purchased may only be licensed for use in the US.

As above, the welfare results depend on how many markets there are and what prices are set. Some consumers will win and others will lose.

Subsection 8.5.9. A Numerical Example of Third Degree Price Discrimination:

Consider an industry with two types of agents, A and B , with separate demand curves. Note that if profits are being maximized, it must be the case that:

$$MR_A = MR_B = MC.$$

Otherwise, a firm could keep the overall output the same, but shift some units from the low MR market to the high MR market and thereby increase profits. Let $P(Q_i)$ be the inverse demand curve for a given market $i \in \{A, B\}$. The total revenue in one of these markets is:

$$TR(Q_i) = Q_i P(Q_i).$$

We take the derivative and set it equal to zero to find the profit maximizing output level. With a little algebra, we get this:

$$\frac{\partial TR}{\partial Q_i} = MR(Q_i) = P_i(Q_i) + Q_i \frac{\partial P_i}{\partial Q_i} = P(Q_i) + Q_i \frac{\partial P_i}{\partial Q_i} \frac{P_i}{P_i} = P_i - P_i \frac{1}{|\epsilon_i|}.$$

From the argument above, $MR_A = MR_B$ so:

$$P_A - P_A \frac{1}{|\epsilon_A|} = P_B - P_B \frac{1}{|\epsilon_B|} \Rightarrow P_A \left(1 - \frac{1}{|\epsilon_A|}\right) = P_B \left(1 - \frac{1}{|\epsilon_B|}\right).$$

Therefore, if $|\epsilon_A| > |\epsilon_B|$ then $P_A < P_B$ and inversely.

We conclude that the more elastic market gets a lower price. This makes sense. If you buy regardless of price, (that is, you are an inelastic demander) you get charged a high price. If you try to find an alternative if the price increases (that is, you are an elastic demander), you get low prices as the firm tries to draw you in.

A final word of caution: costs of providing a good or service may differ across groups and this may result in different pricing. For example, consider discount prices for kids at all-you-can-eat buffets, and family discounts at gyms. Children eat less, and it would be unusual if all members of a family used the gym membership as intensely as someone who signs up as an individual. Thus, the different prices charged in these different markets may reflect the different costs of serving them and may have nothing to do with price discrimination.

Subsection 8.5.10. Two-Part Tariffs

A **two-part tariff** is a marketing strategy in which a firm sets two separate prices, a fixed price that a consumer must pay regardless of how much of a good he consumes, and a per-unit price the consumer pays for each unit of good.

To see how this works, consider the case of a single consumer and a firm with a constant MC . The profit maximizing strategy for the firm is to set the per-unit price equal to MC , and the fixed price equal to the consumer surplus (or just less than this). In the figure below, the per-unit price would be $P_{TPT} = MC$ and the fixed price would be a small amount less than A . At this per-unit price, the consumer would buy Q_{TPT} units and get a consumer surplus of A . Thus, the right to buy at price of P_{TPT} is worth A to him, and he would be willing to pay anything less than A for this privilege. Note that essentially all the consumer surplus goes to the producer in this case, but that there is no DWL .

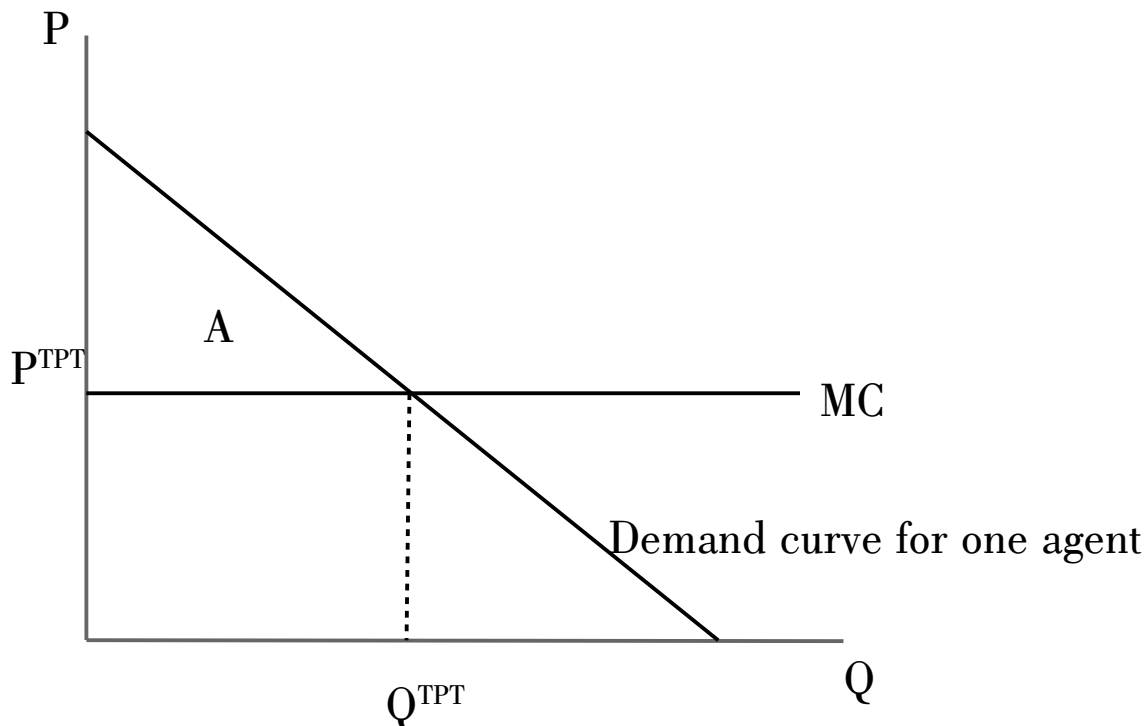


Figure 107: Optimal two-part tariff

We can generalize this a bit. Suppose that there were many identical agents with this demand curve. A profit maximizing firm should set a fixed price just less than A and a per-unit price equal to MC . All agents would accept this offer, almost all the surplus would go to the firm, and DWL would be zero.

What would happen if there were several types of agents. For example, it is a little known fact that each year on groundhog's day, as many English professors and economists as can find the time get together in Moose Jaw, Saskatchewan for a one-day wienie roast, beer blast, and mutual admiration event. English professors are cultured and enlightened people driven by art and unburdened by the mundane things that surround us. Economists study the dismal science and feel collective guilt for our role in facilitating price discrimination. As a result, economists have much higher demand for beer than English professors.

Consider the figure below. If the monopoly supplier of beer in Moose Jaw (Beer, Eh? Inc. or BEI) sets up a tent at the event, how should he choose the tariffs? Of course, it is always optimal to set the per-unit price equal to the MC , but what should the admission price to the tent be?

If BEI chooses A as the admission price, both economists and English professors would be willing to pay. If the price were set at $A+B$, only the economists would drink. Let $\#EP$ and $\#E$ denote the number of English professors and economists, respectively, who are at the event. The choice is therefore between a profit of $(\#EP + \#E) \times A$ and $\#E \times (A+B)$. That is, between losing some potential revenue from high demanding economists and losing all revenue from low demanding English professors while exploiting economists fully.

Which is best? Clearly it depends on how many economists and English professors are at the event. If the crowd is mostly English professors, $\#E \times B$ is small compared to $\#EP \times A$ and it is better to charge the lower admission price. If the crowd is mostly economists, the opposite is true. In the real world, of course, each consumer has a different demand curve, and balancing these gains and losses is difficult and requires a lot of information about how demand is distributed.

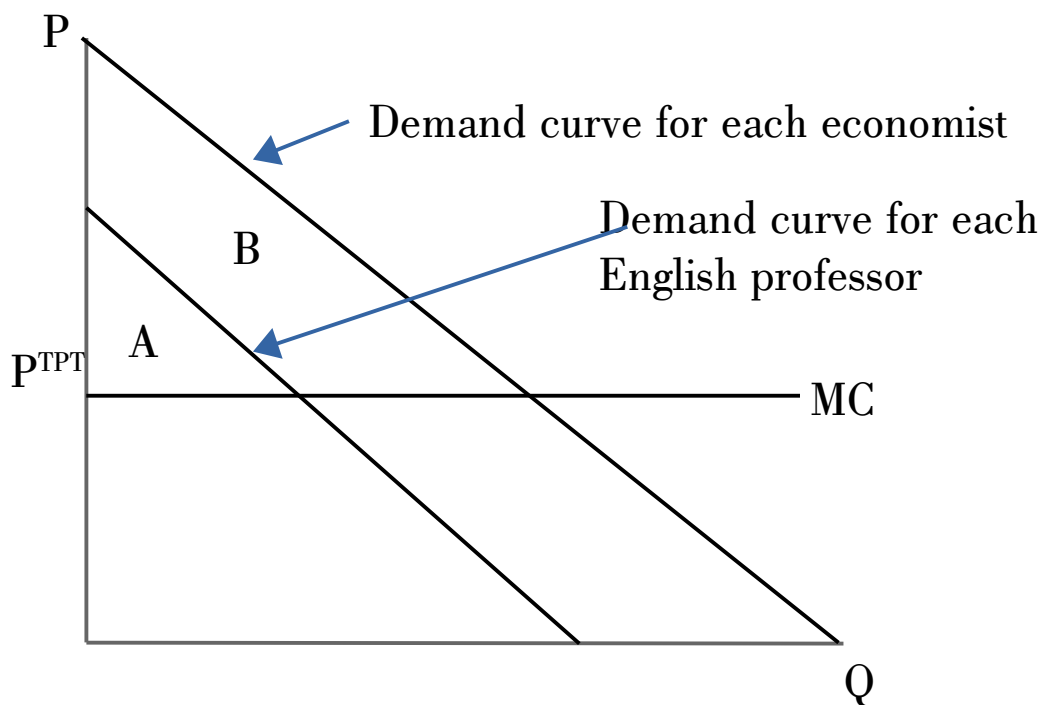


Figure 108: Two-part tariff with two types of agents

Subsection 8.5.11. Resale Price Maintenance

Resale Price Maintenance (RPM): A marketing strategy in which a firm supplies its goods to retailers on the condition that they sell the goods at a specified price or above.

It might seem odd that a firm would do this. After all, a monopoly producer of a good controls the wholesale price, and if retailers lower the price to consumers, then the firm only sells more while retailers suffer. Why would a firm try to protect retailer's profits by preventing them from competing on price?

Consider this from the standpoint of a retailer. If selling a good means guaranteed profits, the retailer would like to sell as much as possible. Since cutting prices is not an option, the only thing he can do is try to convince consumers to buy more of the good and, in particular, buy from him instead of his competitors.

Retailers typically solve this problem by offering better product placement (the goods are put on more attractive displays near store entrances), offering services (think about how cosmetics are sold with offers of free make-overs and advice about what to buy and how to use the products), expertise, and a chance to experience the goods (going to a stereo store and listening to the speakers you want to buy), or by a nice showroom or sales experience (what do you think Ferrari dealerships look like?). Stores may also steer customers to the RPM goods in preference to those sold at zero profit competitive prices.

Thus, the argument is that firms offering high-end goods, goods that are complicated to use or understand, or goods that have to be experienced to be appreciated, may find that sales are higher when retailers compete for customers by offering various kinds of customer service than by lowering the price. One irony here is that if retailers are competitive, they will end up spending all the profit that they might have made due to the higher profit margin on customer service. If they offer less service than other retailers, customers will go to the more pleasant stores. As a result, any potential profits to retailers end up being competed away.

Whether consumers benefit or not is debatable. Retailers provide expertise, help and a pleasant buying experience, all things that customers value. On the other hand, consumers pay higher prices as a result of RPM. It is unclear how these balance out.

Subsection 8.5.12. Durable Goods Monopolists

What pricing strategy should a firm with monopoly power employ if it produces a **durable good** such as automobiles, computers, aircraft, or aluminum? Durable goods last over a period of time and continue to provide services until they wear out. This contrasts with nondurable goods which are consumed and provide value all at once. Producing a durable good creates two significant tensions that lower a monopoly's profit potential.

First, goods sold earlier compete with goods sold later. Consider aluminum for example. Aluminum is fairly easy to recycle. If Alcoa buys up all the bauxite deposits in the world, it has a natural resource monopoly in supplying aluminum. It can charge a very high price at first. However, the aluminum it sells at first creates secondary suppliers in the form of recycling firms.

This recycled aluminum competes with what Alcoa supplies each year and reduces Alcoa's monopoly hold on the market. As a result, the price of aluminum follows a downward trajectory over time, as do Alcoa's profits. Alcoa therefore has a choice. It can sell only small amounts of aluminum each period and take high per unit profits on the small output produced, or produce larger amounts and make smaller per unit profits each period but for a larger number of units. How to balance these factors depends upon the interest rate.

Second, consumers may choose to strategically hold back demand. For example, consumers may anticipate future price drops and therefore forego buying computers in early periods. They may also delay replacing computers for the same reason. Finally, if consumers anticipate not only that prices will drop, but that quality will increase as technology advances, they may hold back demand even more. All of this further cuts into the monopolist's profits and makes it difficult to get large numbers of consumers to be early adopters of new technologies.

Section 8.6. Acquiring Market Power

Now that we know something about what market power is, why it persists, how it is exploited, and how we might make policy to counter its effects. Next we discuss how market power is acquired and retained.

Subsection 8.6.1. Mergers

We often see two firms deciding to merge, or one firm acquiring or being acquired by another. There are many reasons firms might do this. On the positive side, mergers allow firms to share technologies and expertise. Becoming one large company may also make it possible to take advantage of economies of scale or scope. The company that emerges from a such merger may be worth more than the sum of its parts. It may even be able to offer a better product at a cheaper price to consumers. On the negative side, if two companies merge, they cease to be competitors. If too many companies merge in an industry, prices may rise due to the exercise of market power on the part of the surviving firms. Often we see both effects going on at once.

Subsection 8.6.2. Horizontal Integration

Horizontal Integration: A merger between firms that produce the same good.

Typically, we see increases in both market power and efficiency with horizontal mergers. Antitrust regulators may get involved if the merger results in very significant market concentration (i.e. one, two, or three firms producing almost all the output in a market), but sometimes not even then. Examples include airlines and cellphone providers buying up competing firms.

Subsection 8.6.3. Vertical Integration

Vertical Integration: A merger between a firm and one of its suppliers or customers.

Vertical mergers do not in themselves increase a firm's market power. Both of the markets (the input good and the output good) are just as competitive or noncompetitive as they were before the merger. However, such mergers have the effect of causing the firm that produces the input good to internalize any harm or benefit its actions might have on the new firm with which it merged.

Suppose both input and output good markets were monopolized before the merger. Then the input good firm would have been charging a monopoly price to the output firm, and the output firm would have incorporated this into its marginal costs when it set the price of its good for consumers. If these two firms merge, the single firm uses the true marginal cost of the input good to help determine the marginal cost of the output good. Thus, MC will be lower for the merged firm, the quan-

tity of the output good goes up, and the price goes down for consumers. At the same time, it can be shown that the profits made by the merged company are higher than the sum of profits of the original two companies. Thus, both consumers and producers win in this case.

On the other hand, suppose that the input good is a monopoly, but the output market has many firms. If the input company buys one of the output companies, it may be able to exploit its monopoly position. In particular, it could give the output company advance knowledge of the technical specifications of upcoming products or new pricing policies. The input company also knows what is being ordered by the competing output companies, which may give them a competitive advantage.

For example, computer companies became concerned when Microsoft entered the tablet market for all of these reasons. The overall impact of this kind of merger is hard to determine. In the longer run, the acquired output company may have a competitive advantage and thus may get large enough to have market power, which would hurt consumers. In the short run, the price, or quality, or availability of the input good to competing firms may move unfavorably and hurt both these firms and consumers. Moreover, unless the input company plans on doing something like what is outlined above, it is hard to see why it would want to engage in vertical integration in the first place. Regulators are rightly suspicious of such mergers.

Subsection 8.6.4. Predatory Pricing

Predatory Pricing: A marketing strategy in which a firm sets an artificially low price for its product in an effort to drive out competitors and monopolize a market.

We sometimes see “price wars” between airlines, phone companies, and internet service providers. This is an interesting phenomenon to analyze from a game theoretic standpoint, but we will not do so here. In broad strokes, firms set prices that may be below costs and hope to wait out their competitors. Once the competitor leaves it raises the price to monopoly level and makes money until a new competitor arises. As a result, a lot depends on a firm being able to credibly signal that it is willing to lose money for as long as it takes to drive out the opposition.

Of course the opposition has the same incentive, and there are many interesting ways that a firm might try to convey that they are tougher (or stupider) than any competitor or potential entrant. On the other hand, price wars are sometimes just competitive behavior, and it is difficult for antitrust agencies to tell exactly what is going on.

Subsection 8.6.5. Duopoly and Oligopoly

We close this chapter with a short examination of less extreme cases of market power. As we mentioned, a duopoly is a market with two firms and an oligopoly is a market with more than two, but not a great many firms. In both cases, the firms in the market have market power and so have to consider the actions of other firms when deciding on their own. In other words, firms behave strategically. We will use game theory to explore their behavior.

There are two basic approaches to modeling this type of market:

Bertrand Oligopoly: A market with a small number of firms in which each chooses a price taking the prices of the other firms as given.

Cournot Oligopoly: A market with a small number of firms in which each chooses a quantity taking the quantities of the other firms as given.

The Bertrand case is very easy to analyze. Suppose there are two firms with identical and constant marginal costs MC_0 . Suppose firm 1 sets a price of $P_1 > MC_0$. What is the best-response of firm 2 to this? We see immediately that firm 2 should set a price of $P_2 = P_1 - \varepsilon$ where ε is as small as possible. The reason is that as long as firm 2's price is even slightly below firm 1's price, all consumers will choose to buy from firm 2. Thus, firm 2 takes over the entire market while charging the highest price it can, taking firm 1's price as fixed.

This does not describe an equilibrium, of course. Now, firm 1 sees that firm 2 has set a price of $P_2 = P_1 - \varepsilon$. Using the same logic as above, the best-response is to set a price equal to $P_2 - \varepsilon = P_1 - 2\varepsilon$. That is, take the entire market by setting the new price only slightly lower than firm 2's. This continues until $P_1 = P_2 = MC_0$. No firm has an incentive to lower prices further as this would give negative profits. On the other hand, raising prices would also be fruitless, since it would result in zero sales and profits. Thus, we arrive at the competitive price and quantity even if only two firms are in the market under the Bertrand assumptions.

Things are slightly more complicated under Cournot competition. Consider the simple case of a duopoly in which two firms have identical and constant MC , and face a linear demand curve. For example:

$$TC(Q) = 100Q \Rightarrow MC(Q) = 100$$

and

$$D(P) = 200 - P \Rightarrow D^{-1}(Q) = 200 - Q \Rightarrow TR(Q) = Q(200 - Q).$$

Suppose that firm 1 enters the market first as a monopolist. Profits as a function of Q_1 are

$$\pi(Q_1) = Q(200 - Q_1) - 100Q_1.$$

Taking the derivative to maximize profits gives us:

$$200 - 2Q_1 - 100 = 0 \Rightarrow Q_1^* = 50.$$

Suppose now the second firm enters. In the Cournot model, firm 2 takes firm 1's output as given. This means that if firm 1 produces 50 firm 2's TR is the following:

$$TR(Q_2) = Q_2(200 - 50 - Q_2),$$

which implies:

$$\pi(Q_2) = Q_2(200 - 50 - Q_2) - 100Q_2.$$

Taking the derivative to maximize profits gives us:

$$150 - 2Q_2 - 100 = 0 \Rightarrow Q_2^* = 25.$$

We could go back and forth with each firm finding a best-response to the other firm's most recent output decision, but we can find the equilibrium directly by writing a firm's best-response problem in a more general form. Suppose firm j produces $\bar{Q}_j$. Then the profits of firm i taking this quantity as fixed are:

$$\pi(Q_i) = Q_i(200 - \bar{Q}_j - Q_i) - 100Q_i.$$

Taking the derivative:

$$200 - \bar{Q}_j - 2Q_i - 100 = 0 \Rightarrow Q_i^*(\bar{Q}_j) = \frac{100 - \bar{Q}_j}{2}.$$

We call this the firm's **best-response function**.

Since the firms are identical, we know that at an equilibrium, their output must also be the same $\bar{Q}_1 = \bar{Q}_2 \equiv \bar{Q}$. Thus, we can find the equilibrium by substituting and solving: $\bar{Q} = \frac{100 - \bar{Q}}{2}$, which gives us $3\bar{Q} = 100$, or $\bar{Q} = 33.3$. Since there are two firms, the total output goes up from the monopoly level of 50 to the Cournot duopoly level of 66.7, while the price drops from 150 to 133.3.

The table below shows these best-response quantity choices when firms alternate acting. Thus, firm 1 chooses quantity assuming firm 2's quantity is fixed, and then firm 2 chooses a new quantity assuming that firm 1's quantity is fixed. You can see how the quantity that each firm chooses converges to 33.3.

	Firm 1 quantity	Firm 2 quantity
Round 1 Firm 1 chooses Q	100	-
Round 2 Firm 2 chooses Q	50	25
Round 3 Firm 1 chooses Q	37.5	25
Round 4 Firm 2 chooses Q	37.4	31.25
⋮	⋮	⋮
Limit (round ∞)	33.3	33.3

Table 14: Cournot duopoly firms' best-responses with alternating choices

We can generalize even more to consider a Cournot oligopoly problem with F firms. Now, each firm i takes the output level of all the other $F-1$ firms combined, denoted $\bar{Q}_{-i}$, as fixed when finding a best-response. Profits become:

$$\pi(Q_i) = Q_i(200 - \bar{Q}_{-i} - Q_i) - 100Q_i.$$

Taking the derivative with respect to its own output level Q_i , firm i finds:

$$200 - \bar{Q}_{-i} - 2Q_i - 100 = 0 \Rightarrow Q_i^*(\bar{Q}_{-i}) = \frac{100 - \bar{Q}_{-i}}{2}.$$

Again, the firms are symmetric, so we can find the equilibrium by substituting and solving we find,

$$\bar{Q} = \frac{100 - (F-1)\bar{Q}}{2},$$

which gives us

$$(F+1)\bar{Q} = 100, \text{ or } \bar{Q} = \frac{100}{(F+1)}.$$

This in turn implies that the total Cournot oligopoly output level added up over all F firms is:

$$\frac{F}{(F+1)}100,$$

so the price becomes:

$$200 - \frac{F}{(F+1)}100.$$

Notice that as F gets large, both the total Cournot oligopoly quantity and price converge to 100. You can easily verify that if firms behave as competitors and set output where $MC = D(P)$, price and quantity would also be 100. Thus, as the number of firms in an oligopoly increases, the outcome gets farther away from the monopoly and closer to the competitive equilibrium outcome. The table below illustrates this.

Number of firms	Price	Total Quantity
1	150	50
2	133.3	66.7
3	125	75
4	120	80
$\vdots$	$\vdots$	$\vdots$
Limit (∞ firms)	100	100

Table 15: Equilibrium price and quantity in a Cournot oligopoly with several firms

Glossary

Bertrand Oligopoly: A market with a relatively small number of strategic firms, each of which chooses a price for their output, taking the pricing strategies of the other firms as fixed. Each firm commits to producing any quantity demanded at the price they choose.

Best-Response Function: Given a game and a choice of strategies for all the other players, the remaining player can find the element of his own strategy set that maximizes his payoff. Doing this maximization exercise for all possible lists of strategy choices for the other players gives a mapping from the strategy space of all other players into a best-response strategy for the remaining player.

Bundling: A marketing strategy for firms in which specific amounts of one or more goods are sold in a package at a fixed price. This contrasts with the more conventional **per-unit** pricing which allows consumers to choose how much of each good to buy at a stated price per unit. Versioning and quantity discounts are variations of bundling strategies.

Copyright: A legal protection that gives an agent the exclusive right to sell or license a creative work for a set period of time. This monopoly covers a wide variety of mainly intellectual property such as books, songs, recordings, movies, video games, certain kinds of computer code, paintings, photographs and other types of art. The motivation to grant this monopoly is to reward and encourage the creation of such works in order to enrich culture and advance knowledge. As such, new works that are transformative or which are new expressions of ideas contained in copyrighted works are specifically not covered by this protection.

Cournot Oligopoly: A market with a relatively small number of strategic firms, each of which chooses an output level taking the output strategies of the other firms as fixed. Each firm commits to selling this quantity at whatever price emerges in equilibrium.

Duopoly: A market in which two firms supply all the output to the market.

Durable Goods: Goods which last over a period of time and continue to provide services until they wear out. This contrasts with nondurable goods which are consumed and provide their value all at once.

First Degree Price Discrimination: A pricing strategy in which a firm charges every agent exactly his reservation price for a good. This is also known as **perfect price discrimination** and leaves the consumer with no surplus but also produces zero overall dead weight loss.

Horizontal Integration: A merger between firms that produce the same product. This increases the market power of the firm that results.

Identification: In the context of price discrimination, the ability to distinguish between higher and lower demand consumers either directly or through various indirect marking strategies.

Inverse Demand Curve: A demand curve that gives the quantity demanded by consumers in a market as a function of price: $Q(P)$. The inverse demand curve simply requires solving the demand curve for P in terms of Q : $P(Q)$.

Legal Barriers to Entry: A law or regulation that makes it difficult or impossible for new firms to enter an industry and compete with the incumbent firms.

Marginal Expenditure: The total expenditure by a monopsonist on the quantity of goods it chooses to buy is quantity it buys times the price as given by the supply curve: $TE(Q) = QP^S(Q)$. The marginal revenue is found by taking the derivative of total expenditure with respect to quantity:

$$ME(Q) \equiv \frac{\partial TE(Q)}{\partial Q} = \frac{\partial QP^S(Q)}{\partial Q} = P^S(Q) + Q \frac{\partial P^S(Q)}{\partial Q}.$$

Marginal Revenue: Total revenue for any firm is equal to the quantity produced times the price at which the good is sold, that is, quantity times the price given by the inverse demand curve: $TR(Q) = QP^D(Q)$. The marginal revenue is found by taking the derivative of total revenue with respect to quantity:

$$MR(Q) \equiv \frac{\partial TR(Q)}{\partial Q} = \frac{\partial QP^D(Q)}{\partial Q} = P^D(Q) + Q \frac{\partial P^D(Q)}{\partial Q}.$$

Merger: Two or more firms becoming a single firm through buyout, takeover, or a vote of the stockholders or other governing body. The new unified firm holds ownership over the assets of all the firms in question.

Monopoly Power: The ability of a firm, consumer or other agent to affect the price in a market through his actions. This is also called **market power**.

Monopoly: A market in which a single firm supplies all the output.

Monopsony: A market in which a single agent consumes all the output.

Natural Monopoly: An industry is said to be a natural monopoly if the minimum of the AC curve is above the demand curve. Natural monopolies are typically observed when an industry is characterized by high costs of initiating production, but lower incremental costs for subsequent units of output.

Natural Resource Monopoly: An industry in which a single person or firm owns all of a particular item that exists in the world and is therefore the only one with the power to supply the good to the market.

Network Externality: A situation in which the value of the product to each individual consumer increases with the number of other consumers who buy the good (and thus, join the network).

Resale: The act of an original purchaser of a good selling the good to a different consumer rather than consuming the good himself. If resale is possible, it is difficult for firms to price discriminate between high and low demand consumers.

Oligopoly: A market in which a small number of firms supply all the output.

Patent: A legal protection that gives an agent the exclusive right to sell or license an invention. This monopoly covers mainly items that have or could have a material existence such as machines, medicines, chemicals, crops, and materials, and has more recently been extended to business methods and certain types of computer algorithms. Similar to copyrights, the motivation to grant this monopoly is to reward and encourage the invention of new and useful products. Unlike copyrights, patents cover both the expression of, and idea behind, the innovation. The trade-off for this broader protection is that patents are of much shorter duration.

Predatory Pricing: A pricing strategy in which a firm sets an artificially low price in an effort to drive out competitors and monopolize a market.

Price Discrimination: A pricing strategy in which firms charge different prices to different consumers in an effort to increase profits by extracting more revenue from higher demand consumers without losing profitable sales to lower demand consumers. If firms had to sell at one price, as in the standard monopoly case, they would have to choose between losing revenues from the high demanders or losing sales to the low demanders. In order for a firm to engage in this practice, three conditions must be satisfied: market power, identification, and no resale.

Quantity Discounts: A pricing strategy in which a firm offers goods at a lower average cost per unit provided that the consumer agrees to buy at least a certain quantity of the good at once. This can be accomplished by offering a menu of quantity bundles (a six-pack, a case, a pallet) or by specifying price break points (rent is \$100 per day, \$2500 per month or \$20,000 per year).

Rate of Return Regulation: A regulatory strategy used by governments mainly for utilities such as power, water, and telephone service. Regulated firms are required to report all of their direct and indirect costs and then set prices such that economic profits are apparently zero. More formally, firms must set a price such that $P = AC$ where the AC includes all the opportunity costs, especially capital invested by stockholders.

Resale Price Maintenance (RPM): A pricing strategy in which a monopoly firm sells to retailers on the condition that they retail the goods at a specified price or above.

Second Degree Price Discrimination: A pricing strategy in which a firm charges prices which depend on the quantity that a given consumer demands. Quantity discounts are the leading example of this strategy.

Third Degree Price Discrimination: A pricing strategy in which a firm charges different prices in different markets. Markets may be distinguished by location, time, type of consumer or any other observable trait that allows a firm to distinguish high demand from low demand consumers. Within a given market, however, only a single price is charged.

Trademark: A legal protection that gives an agent the exclusive right to use a name, slogan, logo, or other identifying item. Trademarks last 10 years, but can be renewed indefinitely. The protection is limited to a given market or line of business for which it is registered with a view to preventing any possibility of consumer confusion. Thus, if Moe's bar opens in Springfield and trademarks its name, it does not prevent another firm from opening a Moe's bar in Shelbyville as these are distinct markets. If Moe's were to be franchised all over the country, the bar in Shelbyville would have to choose a different name, however, it would still be legal to open Moe's Motel, or Moe's Tire and Auto.

Two-Part Tariff: A pricing strategy in which a firm sets two separate prices: a fixed price consumers must pay regardless of how much of a good they consume, and a per-unit price they must pay for each unit of consumption.

Versioning: A type of bundling in which a firm produces slightly different products with different bundles of features in an effort to extract the highest possible price from high demand consumers while still being able to sell goods to lower demand consumers at a lower price.

Vertical Integration: A merger between firms that produce or consume one another's output, for example, when a firm merges with one of its suppliers.

Problems

1. Competition and Monopoly

- a. During Prohibition, Peoria had only one speakeasy. One year the anti-alcohol candidate for mayor was elected. As a result, the protection money that the speakeasy was forced to pay to the chief of police each month doubled. Predict what happened to the price of beer in Peoria.
- b. Suppose that a candidate from the same party was elected in Chicago, where speakeasies are a competitive industry. What would happen to the price of beer in Chicago?

2. Consider a firm producing left-handed lug-nuts. The total cost of production is:

$$TC(Q) = \frac{3Q^2}{2}$$

and the company faces a demand curve given by:

$$D(P) = 100 - P.$$

- a. What is marginal cost function?
 - b. What marginal revenue function?
 - c. What is the monopoly quality and price?
 - d. Suppose this firm behaved like a competitor and set price equal to marginal cost. What is the competitive price and quantity?
3. True false or uncertain. When the variable costs go up in a competitive industry, the long run effect is for price to go up and quantity to go down. But since there is no entry or exit in a monopolized industry, the monopoly simply has to absorb rises in variable costs, and there are no long run effects on price or quantity.
 4. Consider an industry with a linear demand curve. Suppose that a \$1 per unit tax is placed on producers. True, false or uncertain: Since monopolists don't have to worry about other firms undercutting their price, more of the tax will be passed along to the consumers if the industry is monopolized than if it is competitive.
 5. The South African diamond cartel is a monopolistic supplier of raw uncut diamonds. In a graph, show the equilibrium price and quantity of diamonds. What is the consumer, producer and social surplus in this case. Suppose that an optimal subsidy was placed on the diamond market. What is the price, and quantity, and the consumer, producer, and social surplus in this case? What is the dead weight loss from the monopoly?
 6. Competitive firms make zero profits in the long run. In some industries like public utilities, there are such strong economies of scale that the only efficient equilibrium is to have only one firm in the market. To prevent public utilities from exploiting their market power, regulatory boards have been established to make sure that these firms also make zero profits.

- a. What do we call the economic situation in which it is efficient to have only one firm in the market? What do we call the type of regulation that forces such firms to make zero profits?
 - b. In a picture, illustrate the situation described above. Be sure to label the unregulated, the regulated, and the efficient price and quantity.
 - c. Using your picture, show either that the unregulated situation is always worse or only sometimes worse than the regulated situation.
7. The Federal government is considering regulating the cable television monopoly. They will do this by rate of return regulation. In a graph, show the consumer and producer surplus, and the dead weight loss when the monopoly is unregulated. In the same graph, show the consumer and producer surplus, and the dead weight loss when the cable company is regulated. Can you say for certain that consumers benefit from this regulation? Can you say for certain whether the dead weight loss is larger or smaller under regulation?
 8. Health clubs often have pricing schemes like the following: A non-racquetball membership that includes aerobics but excludes the use of the racquet ball courts costs \$50 per month. A non-aerobics membership that includes the use of the racquetball courts but excludes aerobics also costs \$50 per month. However, an all-inclusive membership costs only \$60 per month. This sort of pricing structure is like a quantity discount and suggests price discrimination. For example, if you already do aerobics, you probably have a low demand for racquetball, so you are offered a chance to buy the right to use the courts at a reduced rate. People with higher demands, in particular non-aerobics doers, are charged a higher price for racquetball. Is this the only explanation possible, or can you think of another reason why different types of people are charged different prices for the same good?
 9. Given the huge popularity of Buffalo wings (hot chicken wings), you get a brilliant idea for a new restaurant. You will sell hot turkey wings, which you will call "Prehistoric Buffalo wings". Your idea is to sell these at marginal cost (\$1 each), but to charge admission to the restaurant. Suppose there are two types of consumers: frat boys and geeks. Frat boys love wings and the individual demand curve of each frat boy for PHBW's is: $D^{FB}(P) = 15 - 3P$. Geeks, on the other hand, don't have as big an appetite. The demand of each geek is given by: $D^G(P) = 8 - 2P$.
 - a. Draw a picture with both demand curves. Now calculate the highest admission price that frat boys would pay, and the highest admission price that geeks would pay to enter your restaurant. (Your answers should be numbers.)
 - b. Suppose that you can only set a single admissions price (discrimination between geeks and frat boys is not allowed). Suppose also that the town contains 200 frat boys and 400 geeks. What is the profit maximizing admission price in this case?

- c. How many geeks would there have to be in the town for you to be indifferent between having the higher and lower of the two admission prices you calculated in part (a) from a profit maximization point of view?
10. There are only two makers of fraternity paddles in the world: Smith Paddle And Novelty Company (SPANC), and Williams Amalgamated Corporations (WAC). Suppose that these companies produce paddles using the same constant marginal cost functions.
- a. Assume that the SPANC and WAC behave as Bertrand duopolists. Explain what this means, and predict the price and quantity in equilibrium. (A simple picture may help illustrate your answer.)
 - b. Suppose that the companies behaved as Cournot duopolists instead. Explain what this means.
 - c. Suppose you are a freshman considering pleading at a fraternity. Remembering that the market demand for paddles is downward sloping, would you prefer that paddles be supplied by Cournot duopolists or Bertrand duopolists? Be sure to explain why.
11. Suppose that banking is deregulated in Illinois. As a result, all the banks in the state merge into one mega-bank. This mega-bank has a monopoly, does it follow that consumers will be worse off?
12. Suppose that the American Hot dog Company buys the National Hot dog Bun Corporation. What is the name for this sort of merger? Under what conditions would you expect consumers to gain from this merger, and why?

Chapter 9. Noncooperative Game Theory

Section 9.1. What is Noncooperative Game Theory?

Economics uses two main theoretical tools to analyze the world. The first revolves around markets and focuses on how agents behave on the average in large groups. This allows us to assume that (a) agents are rational utility maximizers, and (b) agents take their environment (prices, for example) as fixed and not liable to be influenced by their own market choices.

These assumptions are both justified as reasonable approximations of reality when agents are “small” relative to the size of the market. This is because, even though individual agents may not be rational maximizers, in large groups, agents seem to behave on average as if they were. Similarly, even though the market choices of each individual do have some effect on prices, the effect is small enough in a large economy that it does not influence the welfare of individuals enough to make it worthwhile to take it into account in their decision-making.

The second major tool is game theory, especially noncooperative game theory. In contrast to market economics, game theory focuses on interactions between and within small groups of agents. These agents must take into account that their actions may affect the strategic choices of other agents, and therefore, their own payoffs.

Game theory is especially good at modeling situations with incomplete or asymmetric information, with agents who may be less than fully rational, and with agents whose beliefs or expectations about their environment or other agents’ beliefs may affect their strategic choices. Although games with many agents can also be treated, as games become larger, the outcomes tend to look more and more like those that come from markets.

Here are some examples of strategic situations that people might find themselves in:

- Should I donate to PBS?
- Should I speed on the way to work?
- Should I cheat on my income tax?
- Should I accept the low-ball bid just placed on my house?
- Should I seek revenge or turn the other cheek when someone does me wrong?
- Should I bid my true reservation price on Ebay?

In each of these cases, agents must make a strategic choice in hopes of getting the best possible outcome. Agents may play against one, a very small number, or even an infinity of adversaries. Agents may be fully informed about the strategic choices of their adversaries before they choose their own, or they may have to move first or move without knowing what other agents have chosen. Agents may have complete information about the environment, the motivations of their adversaries, and so on, or may have to make their best guesses about these.

Some games are played once and only once, some are played repeatedly, and others are played sequentially with one or several agents making strategic choices in succession. Games can last for a fixed number of periods, an unknown number of periods, or an infinite number of periods. In some cases, agents know the name, type, or objective functions of their opponents as well as the underlying structure of the game including the sequence of moves, the payoffs, and the probability of any random events or states of nature that might affect any of these. In other cases, a subset of these things, or even nothing, is known. We will discuss some of these complications below.

A large literature has grown around exploring many of these details. Here, we provide a brief introduction.

At the most general level, three elements are needed to define a **Noncooperative Game**.

A Set of Agents or Players: Participants in a game might be consumers, producers, workers, politicians, or even animals or nature. All that is required is that they have well-defined objectives they seek to maximize. These agents might be identical, divided into classes (male and female, or voters, lobbyists, and candidates), or may each be unique. Agents might find themselves playing against all the other agents at once, or being matched or divided into smaller groups in some way to form games (marriages or startup companies). Typically, noncooperative games are best used to analyze situations in which the actions of each agent are likely to be noticed and reacted to by the rest of the players. For this reason, most noncooperative games involve a relatively small number of players, although any number of players could be included from a formal standpoint.

A List of Strategies: A **Strategy Set** is a list set of actions that each agent can choose to employ in a game. Strategy sets can be finite (turn right or turn left, or choose one of five possible advertising campaigns), countably infinite (name an integer, or how many times a target server should be sent a page request in a distributed denial of service (DDoS) attack) or uncountably infinite (choose a point in time to fire your gun in a duel, or a probability with which to bluff at poker). In some games, all agents have the same strategy sets (members of a board or directors can vote yes or no on a proposal). In others, strategy sets might differ by agent type, or might even be different for each individual agent.

A Set of Payoff Functions: The payoffs that agents get from playing a game depend on the collective strategy choices of all the players. A **Strategy Profile** is a list that includes a single such choice for each agent. For example, a strategy profile for a board of directors might be (yes, no, no, yes, yes), meaning that the first board member voted yes, the second one voted no, etc. A **Payoff Function** for any given player is a mapping from strategy sets into outcomes.

For example, given the strategy profile (yes, no, no, yes, yes), the proposal passes, and the first board member might get a payoff of \$1000, or a promotion to CEO, or a new office, as a result. Thus, each agent in the game has a payoff function that specifies the payoff he gets for every possible strategy profile.

A Formal Definition of a Noncooperative Game

One-Shot, Simultaneous Move Game: $(\mathcal{I}, S, F)$ where:

Players or Agents: $i \in \{1, \dots, I\} \equiv \mathcal{I}$

Strategies: $s = \{s_1, \dots, s_I\} \in S_1 \times \dots \times S_I \equiv S$ where $s_i \in S_i$

Payoff Functions: $F \equiv (F_1, \dots, F_I)$ where $F_i: S \Rightarrow P$

We denote by P some **payoff space**. This could be any sort of finite or infinite set where the elements represent amounts of money or utility, market shares, discrete objects like houses or art, possible grades of a test, jobs, binaries such as winning or losing a game or a war, etc. To give some vocabulary:

Strategy Set for an Agent: S_i is called a strategy set for an agent i . This is also sometimes call the **Action Space**.

Strategy for an Agent: $s_i \in S_i$ is called a strategy for agent i .

Strategy Profile: $s = \{s_1, \dots, s_I\} \in S_1 \times \dots \times S_I \equiv S$ is called a strategy profile. In other words, a specific strategy choice for each agent.

Deviation Strategy Profile for Agent i : $(\bar{s}_i, s_{-i}) \equiv (s_1, \dots, s_{i-1}, \bar{s}_i, s_{i+1}, \dots, s_I)$ denotes the strategy profile $s \in S$ with the i^{th} element $s_i \in S_i$ deleted, and replaced with an alternative strategy $\bar{s}_i \in S_i$.

Section 9.2. Normal Form Games

The simplest case of a noncooperative game is a one-shot, simultaneous move game. This is also called a **Normal Form Game**. If there are only two players, we can represent this in bimatrix form.

Selfish Gene Game		Bubba Peacock's Genes	
		Big Tail	Little Tail
Skeeter Peacock's Genes	Big Tail	NE = DSE SP: Find a mate then die BP: Find a mate then die	SP: Find two mates and die BP: Find no mates and live
	Little Tail	SP: Find a no mate and live BP: Find two mates and die	SP: Find a made and live BP: Find a mate and live

Table 16: Equilibrium price

The matrix lays out the players (Skeeter and Bubba Peacock's Genes), the strategies sets available to each (have a big tail or have a little tail) and the payoffs to each of the four possible strategy profiles. For example:

$$F_{SP}(BT, LT) = \text{Find two mates and die.}$$

Having a big tail uses a lot of energy and makes a male peacock more vulnerable to predators. However, to carry it off, a male has to be in top physical condition. Thus, while having a big tail will kill a male by the end of the breeding season, chicks dig it. Females will rationally choose the most vigorous male, and so only those with big tails will be able to breed and pass on their genes. Small tailed males will live long, but unhappy lives, and will end up being selected out of the gene-pool. This is an example of a game in which nature or biology is the player. Genes do not make rational choices, and Peacocks do not control their genetic makeup. Nevertheless, genes act as if they maximize survival and replication, and so we can use game theory as a modeling tool.

Subsection 9.2.1. Nash Equilibrium

Now that we know what a non-cooperative game is, the next question is what is likely to happen when one is played. That is, what outcome would we expect to see when agents are faced with one of the strategic situations we model. This really amounts to a behavioral question about the way agents think and so what strategic choices are likely be stable in the real world. Game theorists

spend a lot of time worrying about what properties equilibrium concept should have, and many alternatives have been proposed. The most important of these is Nash equilibrium (NE).

Informally, Nash equilibrium⁶ requires that all agents choose a strategy that is a best-response to the strategies chosen by the other agents in the game. We have already discussed the rationale behind Nash equilibrium, Here is a formal definition:

Nash Equilibrium (NE): A strategy profile $s \in S$ is a Nash equilibrium if $\forall i \in \mathcal{I}$ and $\forall \bar{s}_i \in S_i$, $F_i(s_i, s_{-i}) \geq F_i(\bar{s}_i, s_{-i})$.

For the selfish gene game, the only NE is (BT, BT) . That is, if one peacock has a big tail, the best-response is for the other peacock to have a big tail as well. The ultimate objective of the gene is to replicate. Genes do not care about the cost it imposes on the organism it defines. Although genes do not really choose anything, of course, having the code for a big tail generates maximal replication opportunities for the gene regardless of the genetic coding of the other peacocks. In this sense, big tails are mutually best-responsive under the rules of nature.

Subsection 9.2.2. Dominant Strategy Equilibrium

A second very important type of equilibrium is called dominant strategy equilibrium (DSE).

Informally, a strategy is dominant for a player if it gives him the highest possible payoff in every possible situation. If all players have a dominant strategy in a game, then these collectively form a DSE. More formally:

Dominant Strategy: We say $s_i \in S_i$ is a dominant strategy if $\forall \bar{s}_i \in S_i$ and $\forall s_{-i} \in S_{-i}$, $F_i(s_i, s_{-i}) \geq F_i(\bar{s}_i, s_{-i})$.

Dominant Strategy Equilibrium (DSE): We say $s \in S$ is a dominant strategy equilibrium if $\forall i \in \mathcal{I}$, $s_i \in S_i$ is a dominant strategy.

You can easily verify that (BT, BT) is a DSE as well as a NE for the selfish gene game. This means that each peacock's genes are better off with a big tail *no matter what the kind of genes the other peacock has*. In other words, not only is having a big tail the best thing to do when the other males have big tails, but it is best in all cases.

In general, a strategy profile is a DSE, it must also be a NE (try to prove this). That is:

$$DSE \subseteq NE.$$

Now consider the following game between professors and students. Students can either cram for the test or drink, and professors can test or not test. Students hate to cram (-10) , but hate to get

⁶ Not to be confused with the :Nash bargaining solution, which is an equilibrium concept for cooperative games.

bad grades as well (-20). Professors hate to give exams (-10) but also hate it when students drink instead of cramming (-20). Below, we see the exam game in bimatrix form.

Exam Game		Professor	
		Test	No Test
Student	Cram	S: -10 P: -10	S: -10 P: 0
	Drink	S: -30 P: -30	NE S: 0 P: -20

Table 17: Equilibrium price

You can verify that (D, NT) is the only NE strategy profile of this game. That is, if the student drinks instead of cramming, the professor should not give a test (it is depressing). If the student crams on the other hand, there is still no reason to test (it is a waste of the professor's time.). This means that not giving tests is a dominant strategy for the professor. However, there is no dominant strategy for the students. If the professor decides to give a test, the student's best-response is to cram, if he does not, the best-response is to drink. Thus, there is no DSE.

Section 9.3. Examples of One Shot Games

Subsection 9.3.1. Coordination Games

Many situations require that agents coordinate their actions. For example, if we want to chat or Skype we have to be at our computers at the same time, companies benefit if they use the same or compatible standards in their products, and close friends may want to join a club or chose a certain vacation spot if and only if the other does at the same.

The simplest version of this is a symmetric game in which agents do not care about which outcome is chosen as long as they agree. For example, suppose that there a sudden acceleration of continental drift and Australia crashes into California. Americans, such as Joy Ryder, drive on the right-side of the road, while Australians, such as Helen Weals, drive on the left-side. Now that the Australian and US road networks are connected, we can model the choice of which side of the road to drive on as a game.

As you can see, choosing different sides results in crashes. A bad outcome for both Joy and Helen. However, if they coordinate on either the right-side or left-side, there are no crashes, and they can drive safely. It does not matter which of these two standards is adopted. Thus, we have two NE, (RS, RS) and (LS, LS) . There are no DSE.

Driving Game		Helen Weals	
		Right Side	Left Side
Joy Ryder	Right Side	NE JR: Drive Safely HW: Drive Safely	JR: Crash! HW: Crash!
	Left Side	JR: Crash! HW: Crash!	NE JR: Drive Safely HW: Drive Safely

Table 18: Equilibrium price

It is often the case that agents have preferences over the possible NE. For example, Helen might find it easier to continue to drive on the left, but would prefer to switch to right-side if Joy refuses to

switch. Similarly, technology companies benefit from using standards, but may prefer one standard over another. For example, Sony pushed Blu-ray over HD-DVD as a video disk standard because it held some of the key patents involved and would be able to charge license fees to other companies using their IP.

These asymmetric coordination games are refereed to as **Battle of the Sexes Games**. Consider the case of Kent Waite and Polly Amorous, two deeply caring environmentalists who just started dating. Kent is devoted to saving the whales, while Polly is dedicated to saving the Sidewinder (a snake found in the Arizona desert). They love being together even more than they love saving the environment. Of course, if they could do both at the same time, it would be truly awesome. As you can see, Kent gets the most utility (100) when both he and Polly are on-board ship saving whales. He gets less if they are saving sidewinders together, and even less if he is saving whales by himself. The worst thing would be to find himself out in the desert saving sidewinders without Polly. Again there are two NE: (STW, STW) and (STS, STS) , however Kent prefers the first and Polly prefers the second.

Dating Game		Polly Amorous	
		Save the Whales	Save the Sidewinder
Kent Waite	Save the Whales	NE KW: 100 PA: 50	KW: 25 PA: 25
	Save the Sidewinder	KW: 0 PA: 0	NE KW: 50 PA: 100

Table 19: Equilibrium price

Subsection 9.3.2. The Hotelling Location Game

An interesting variation of a game in which agents end up coordinating is attributed to Harold Hotelling. He considered a locational model in which two firms had to decide where to locate on a street. Specifically, suppose that there was a one-mile-long main street in a town and two agents had decided they were going to open bars. The strategy space is where on the one-mile street to locate. All the people in the town work on this street and they always go to the nearest bar after work for a beer. Assume that business and workers are uniformly distributed on the main street. Given that profits are proportional to the number of customers a bar gets each day, where would the bars be located in a Nash equilibrium?

Suppose that the first bar locates to the left of the center of town. Then by locating next door on the right side, the second bar gets all the customers from the right side since it is the closest bar. But this means the first bar gets less than 50% and the second bar more than 50% of the customers. If the first bar moves just the right of the second bar, the market shares are reversed. Thus, locating to the left or right of the center of town cannot be a NE. In fact, only when the bars locate at the exact center of town next door to one another do we have a stable locational equilibrium.

You can see that the bars end up at the same location. This is unfortunate from the standpoint of customers. It means that the average worker must walk $1/4$ mile to get a beer after work. However, if the bars located at the $1/4$ mile and $3/4$ mile points on the main street, the average distance to a bar for workers would be only $1/8$ mile. Both bars would still get half the customers, and the same profit. Co-locating is therefore socially inefficient.

This is an important model because it applies to more than where businesses choose to locate. For example, companies might find that the only viable strategy is release products that are almost identical to those of their competitors. This reduces customer choice and means that customers are generally forced to choose products that are more distant for the one they would most prefer. This model has also been widely applied in political science. Any political candidate that strays from the center opens up space for a competitor to take positions just a little more moderate and win the election (since he appeals to more voters). The result is that politicians will steer for the middle, or at least say that they are centrists.

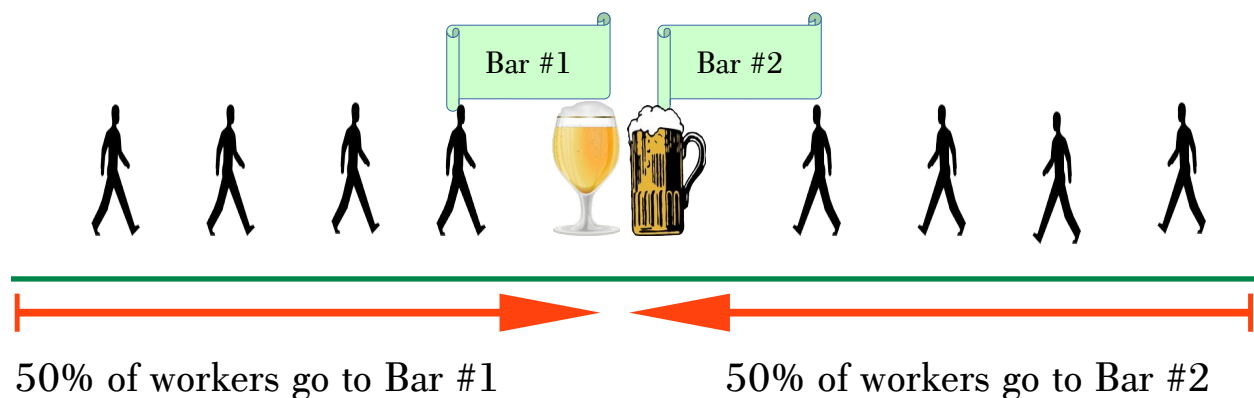


Figure 109: Hotelling Game

Subsection 9.3.3. Prisoners' Dilemma

The classic **Prisoners' Dilemma Game** involves two thieves who are caught by the police. They are taken to separate interrogation rooms and offered a deal. If they confess and their partner does not, they will be allowed to turn state's witness and go free while their noncooperative partner will go to jail for ten years. If both confess, they will both be convicted but given leniency in the form of a five-year sentence. If neither confesses, then they will only get one year in jail each since then they can only be convicted of lesser offenses.

You can verify that the only NE is for both to confess. In fact, this is also a DSE. The reason that this is a dilemma is that if the two prisoners could collude, perhaps by signing binding agreements before they were arrested, they could both commit not to confess. As a result, each would get a one-year sentence instead of going away for five years. The structure of the simultaneous move game the police have set up, however, makes it a dominant strategy for each to confess regardless of what the other does, and so the prisoners are unable to get to the more desirable outcome.

Below is an example of a prisoner's dilemma in the form of a voluntary contribution game. Two programmers, Justin Credible, and Shirley knot, have a choice of contributing or not contributing to a FOSS project. The software is a nonrival good, but its quality depends on the level of effort put into its writing. If both contribute \$75 worth of effort, they each get \$100 worth of benefit from project (a net benefit of \$25 each). If only one of them contributes, they each get \$50 worth of benefit. If neither contributes, the project is abandoned, and no benefits are realized.

Voluntary Contribution Game		Shirley Knot	
		Contribute	Free Ride
Justin Credible	Contribute	JC: \$25 SK: \$25	JC: \$-25 SK: \$50
	Free Ride	JC: \$50 SK: \$-25	NE = DSE JC: \$0 SK: \$0

Table 20: Equilibrium price

Subsection 9.3.4. Zero-Sum Games

Consider the following normal form game between a hunter and a stag. The deer has two choices each day: go graze in the valley or go graze near the pond. The hunter also has two choices: hunt the stag in the valley or hunt the stag at the pond. If they choose the same location, the hunter shoots the stag. If they choose different locations, the hunter goes home empty-handed.

Hunter-Stag game		Stag	
		Valley	Pond
Hunter	Valley	H: 1 D: -1	H: -1 D: 1
	Pond	H: -1 D: 1	H: -1 D: -1

Table 21: Equilibrium price

This is an example of a **zero-sum game** since any gain by one agent is exactly offset by a loss from the other agent or agents. In contrast, the voluntary contribution game, given above, is a positive sum game since all agents can benefit at the same time by choosing certain strategy combinations.

This game is also special in that one agent wants to coordinate, and the other wants to dis-coordinate. This is not essential to zero-sum games in general. For example, in winner take all type games, the strategies may have nothing to do with coordinating or coordinating.

What are the equilibrium of the hunter-stag game? It looks like there are none! If the hunter is one place, the best-response of the deer is to choose a different place. If the deer is one place, the best-response of the hunter is to join the deer. Thus, no pair of strategies is a Nash Equilibrium.

Another classic example of this phenomenon is the Rock-Paper-Scissors game.

Rock Paper Scissors Game		Brittany		
		Rock	Paper	Scissors
Madonna	Rock	M: 0 B: 0	M: 1 B: -1	M: -1 B: 1
	Paper	M: -1 B: 1	M: 0 B: 0	M: 1 B: -1
	Scissors	M: 1 B: -1	M: -1 B: 1	M: 0 B: 0

Table 22: Equilibrium price

You can verify that this game does not have a Nash equilibrium pair of strategies. How do people play this game in real life?

Section 9.4. Mixed Strategies and the Folk Theorem

Subsection 9.4.1. Mixed and Pure Strategies

So far, we have only considered what are called **pure strategies** which require agents to play a given strategy with certainty. We could also allow agents to choose **mixed strategies** in which agents randomize over the pure strategies available to them.

A Formal Statement of Mixed Strategies

Suppose that for each agent $i \in \mathcal{I}$ in a one-shot game, there are finite number $A_i \in \mathbb{N}$ of pure strategies available in his strategy set: $\{s_i^1, \dots, s_i^{A_i}\} \equiv \mathcal{S}_i$. Note that:

$$a_i \in \{1, \dots, A_i\} \equiv \mathcal{A}_i$$

is the index set for agent i 's strategy set.

Mixed Strategy: $p_i = (p_i^1, \dots, p_i^{a_i}, \dots, p_i^{A_i}) \in \Delta^{A_i-1}$ where $p_i^{a_i}$ is interpreted as the probability that agent i chooses strategy $s_i^{a_i} \in \mathcal{S}_i$.

Pure Strategy: A pure strategy is a degenerate mixed strategy where $p_i^{a_i} = 1$, for some $s_i^{a_i} \in \mathcal{S}_i$. and so, $p_i^{\hat{a}_i} = 0$, $\forall \hat{a}_i \neq a_i$.

Mixed Strategy Profile: $p = (p_1, \dots, p_i, \dots, p_I) \in \Delta^{A_1-1} \times \dots \times \Delta^{A_i-1} \times \dots \times \Delta^{A_I-1}$.

Note that given a mixed strategy profile p , the probability that agents will end up jointly playing any given pure strategy profile $(s_1^{\bar{a}_1}, \dots, s_I^{\bar{a}_I}) \in \mathcal{S}_1 \times \dots \times \mathcal{S}_I$ is:

$$\prod_{i \in \mathcal{I}} p_i^{\bar{a}_i}.$$

Expected Value: The expected payoff or expected value to agent i of participating in any given a mixed strategy profile p is:

$$EV_i(p) = \sum_{(s_1^{\bar{a}_1}, \dots, s_I^{\bar{a}_I}) \in \mathcal{S}} \left(\prod_{i \in \mathcal{I}} p_i^{\bar{a}_i} \right) F_i(s_1^{\bar{a}_1}, \dots, s_I^{\bar{a}_I}).$$

In words, for any given pure strategy profile: $(s_1^{a_1}, \dots, s_I^{a_I}) \in \mathcal{S}$, multiply the payoff that agent i receives, $F_i(s_1^{a_1}, \dots, s_I^{a_I})$, by the probability that $(s_1^{a_1}, \dots, s_I^{a_I}) \in \mathcal{S}$, ends up being played given the mixed strategy profile p , then add this expected payoff up over all possible pure strategy profiles in the game. Note that there are at total of $\prod_{i \in \mathcal{I}} A_i$ such pure strategy profiles.

Consider the hunter game again. What if the deer decided to flip a coin between V and P ? In other words, his strategy is to choose V with probability .5 and P with probability .5 (thus, $(p_D^V, p_D^P) = (.5, .5)$)? What would be a best-response for the hunter? The expected payoff from choosing V with certainty is $.5(-1) + .5(1) = 0$. The expected payoff from choosing P is the same: $.5(1) + .5(-1) = 0$. Thus, hunter is indifferent between V and P since both have the same expected payoff. This means that P is a best-response, V is a best-response, flipping a coin between V and P is a best-response, in fact, any mixture at all is a best-response since he is randomizing between strategies that have identical payoffs (zero, in this case).

One of the best-responses to the deer flipping a coin is for hunter to flip a coin. You can easily verify that this means that the expected payoff to the deer of V and P are identical (zero) and so he is also just as happy flipping a coin as choosing a pure strategy in response. We conclude that each agent choosing a 50/50 mix over their two strategies are mutually best-responsive and are therefore a NE. So, while we have no pure strategy NE, we do have a NE in mixed strategies. It turns out that this is a general principle.

Nash Equilibrium Existence Theorem: If a game has a finite set of players, each of whom (a) has a finite set of strategies, (b) has a payoff function that maps his strategy set into $\mathbb{R}$, and (c) satisfies the expected utility axioms in payoffs, then there exists at least one NE in mixed strategies.

This is an elementary version of the existence theorem. More general versions are available in the literature.

Subsection 9.4.2. Repeated Games

Suppose that agents play a one-shot game several times with each other. This is a special case of a game in extensive form (discussed in the next section) which is called a **repeated game**. Such games can continue for either a finite or an infinite number of rounds of play. In each round of the repeated game, players are faced with the same normal form game which is called the **stage game**.

Agents generally are assumed to know the final round of play, T , if there is one, but some variants introduce uncertainty and allow each period to have some known probability of being the final round. Since play takes place over time, payoffs received in later periods of play are usually discounted by the interest rate or internal rate of time preference. Most frequently, agents play each round of the stage game against the same opponents, however, in some cases, agents are matched randomly from a pool of possible opponents.

As an example, suppose that AT&T and Comcast were the only two providers of high-speed Internet connection in Nashville. If they find a way to cooperate, they can both keep prices of Internet data plans high and make large profits each year. If one of the two companies sets a high price, however, the other is tempted to defect from the agreement. By offering a slightly lower price, the

defecting company can get everyone in Nashville to switch providers. Thus, the defecting company ends up with the entire market and leaves his competitor with no customers at all. If both defect from the agreement to keep prices high, then both earn a much lower competitive profit. The matrix below summarizes this game.

One such punishment strategy is called **tit-for-tat**. Very simply, a tit-for-tat strategy means that each agent plays whatever his opponent played in the previous period. For example, suppose both Comcast and AT&T start out cooperating and use tit-for-tat. Then AT&T would see that Comcast played *C* in the previous period and would therefore play *C* in the current period. Comcast would do the same. As a result, AT&T and Comcast would cooperate and get a payoff of 100 each period of play.

Suppose instead that at some period, t , Comcast decided that it had had enough of playing nice and chose to defect from the collusive agreement. Doing so allows Comcast to get a big payoff of 200 in period t , but in period $t+1$, AT&T would respond by choosing *D* as well. As a result, Comcast's payoff would drop to 10 from period $t+1$ until the end time. For AT&T, the outcome is even worse. AT&T would get only 0 in period t and then 10 each period for the rest of time.

An even more severe punishment strategy is called the **grim trigger**. In a grim trigger strategy, some stage game strategy profile (pure or mixed) is identified. An agent using the grim trigger plays his part of this strategy profile as long as all the other agents do likewise. However, the first time any other agent defects from this profile, the rest of the agents play a punishment strategy from the next period on until the end of time. Generally, this punishment strategy gives the worst payoff possible to the defecting agent. The reason this trigger strategy is called "grim" is that the punishment is as strong as possible and lasts forever.

In our example, the worst punishment that one player can impose on the other is the **minimax payoff**. The minimax is found by first considering the maximum payoff a player could get for any strategy choice(s) of his opponent(s). Second, the minimum over all of these maximal payoffs is found and used as the punishment payoff.

The idea is that the minimax is the worst payoff that the rest of the players can be certain that they can impose on any defecting agent even when the agent chooses a best-response to minimize the damage. In our example, if Comcast chooses *C* the best that AT&T can do is to choose *D* and get a payoff of 200. If Comcast chooses *D* the best that AT&T can do is to choose *D* and get a payoff of 10. The smallest of these two is 10, so the grim trigger minimax punishment that Comcast would impose on AT&T for defecting would be to play *D* for all future rounds.

Subsection 9.4.3. The Folk Theorem

The next question is: do these strategies make it possible to cooperate in a repeated prisoner's dilemma game? It will depend on how many times the game is repeated. Suppose the agents played the game a finite number of times, say, T . Could we enforce the cooperative strategies as an equi-

librium by using the grim trigger? One might think so. In period 1, defecting gives a one-time payoff of 200 followed by payoffs of 10 from period 2 to T . On the other hand, cooperating each period gives a payoff of 100 in each period. Surely this should be enough to make playing cooperatively beneficial!

Unfortunately, this logic breaks down as we approach the final period. Suppose that the agents cooperated each period and finally arrive at period T . There is no future at this point, and so any threat of punishment is meaningless. Therefore, it is a dominant strategy for both agents to defect in period T .

Now go back to period $T-1$. Both agents know that they will defect in the next period. In other words, the punishment strategy will be used regardless of whether they cooperate or defect in period $T-1$. It is therefore optimal to defect in this period in hopes of getting the big payoff since both players will certainly choose D next period in any event.

This shows that cooperation breaks down in period $T-1$. But then in period $T-2$, all agents know that D will be played in the next two periods regardless of their actions. Again, the grim trigger threat is empty. This logic is called **backward induction**, and shows that when games are played a finite number of times, any possible cooperation **unravels**. We see that if agents can't cooperate in period t , they also cannot cooperate in period $t-1$. Then since we just argued that cooperation is impossible in period T , it is impossible in period $T-1$, and therefore in period $T-2$, and therefore, cooperation is impossible in period 1.

Suppose instead that there is no final period and the game is played forever. The grim trigger strategies do support cooperative play as a NE in this case. To see this, consider the expected value of cooperating as compared to defecting when you know your opponent is playing grim trigger. Suppose we are in period t and we have cooperated up until period $t-1$. If we defect in period t we get a payoff of 200, however, in all subsequent periods, we get 10 instead of 100. However, these payoffs come in the future, so we must discount them by the interest rate, $r > 0$. This makes the present value of following the cooperation and defection strategies from the standpoint of period t the following:

$$PV(\text{coop}) = 100 + \frac{100}{(1+r)^1} + \frac{100}{(1+r)^1} + \dots + \frac{100}{(1+r)^t} + \dots$$

$$PV(\text{defect}) = 200 + \frac{10}{(1+r)^1} + \frac{10}{(1+r)^2} + \dots + \frac{10}{(1+r)^t} + \dots$$

Clearly, unless r is very large (which would mean agents are extremely impatient and care very little about the future compared to the present), the present value of defecting must be smaller than the present value of cooperating. Therefore, if there is no final period and the interest rate is low enough, cooperating each period can be a NE.

In fact, we can support many other strategy pairs as NE with grim triggers. For example, suppose we wanted Comcast to cooperate each period, but we wanted to allow AT&T to flip a coin between cooperating and defecting. For Comcast, the expected present value is:

$$E(PV(coop)) = \frac{(.5(100)+.5(0))}{(1+r)^0} + \frac{(.5(100)+.5(0))}{(1+r)^1} + \dots + \frac{(.5(100)+.5(0))}{(1+r)^t} + \dots$$

$$E(PV(defect)) = \frac{(.5(200)+.5(10))}{(1+r)^0} + \frac{10}{(1+r)^1} + \dots + \frac{10}{(1+r)^t} + \dots$$

Clearly, if r is small enough, it is better to get an expected payoff of 50 each period than 105 for a single period followed by 10 for the rest of time. It is therefore a NE for Comcast to put up with AT&T defecting half the time rather than have the certainty that AT&T will defect every period.

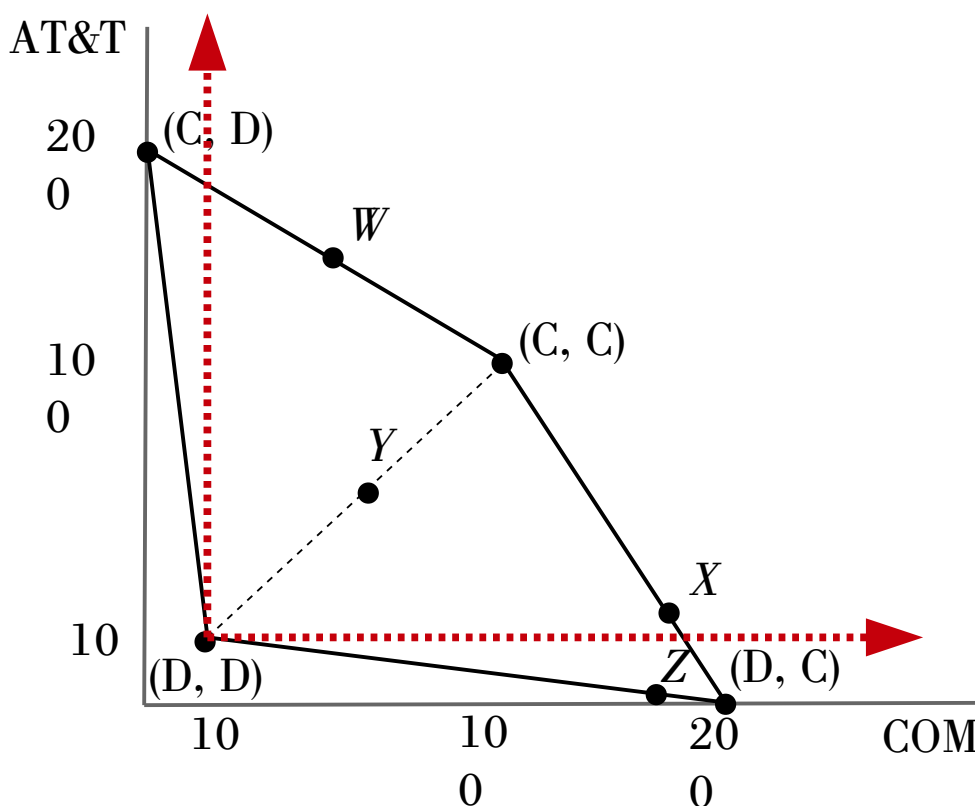


Figure 110: Mixed Strategies in a Prisoner's Dilemma

The figure above can help us visualize this. The two axes are the payoffs of Comcast and AT&T. We have graphed the payoffs that the agents get from each of the pure strategy profiles, and they form the corners of the quadrilateral above.

Now, suppose that Comcast continues to cooperate with probability 1, but AT&T decides to flip a coin between cooperating and defecting. Then half of the time the payoff is $(100, 100)$ and half of the time it is $(0, 200)$. The expected or average payoff is therefore:

$$.5(100, 100) + .5(0, 200) = (50, 150) = W.$$

Suppose that AT&T continues to cooperate with probability 1, but Comcast decides to defect with probability $p_{COM}^D = .8$. The expected payoff is therefore:

$$.2(100, 100) + .8(200, 0) = (180, 20) = X.$$

In the same way, you can verify that the expected payoff of $p_{COM}^D = 1$, and $p_{ATT}^D = .1$ is $(1, 181) = Z$, the expected payoff of $p_{COM}^D = .5$, and $p_{ATT}^D = .5$ is $(77.5, 77.5) = Y$. More generally, the edges of the quadrilateral in the figure are expected payoff that can be achieved with one agent playing one of his pure strategies with certainty, while the other agent plays a mixed strategy. In addition, any point in the interior of the quadrilateral can be achieved as an expected payoff for some pair of strictly mixed strategies on the part of the agents.

Recall that the minimax payoff for each agent in this game is 10. Thus, any expected payoff strictly above this can be enforced as a NE with a grim trigger strategy (if the interest rate is small enough). In the figure, this means that any set of mixed strategies that give expected stage game payoffs above the dashed red lines (such as W, X, Y) can be NE small enough r , while those with less expected payoff such as Z could never be NE regardless of the value of r .

This intuition generalizes beyond prisoners' dilemma games and in fact, is another one of the most important findings in game theory:

Folk Theorem: Consider an infinitely repeated game. Let p be a mixed strategy profile for the one-shot, stage game such that the expected payoff for each agent is strictly larger than his minimax payoff. Then there exists a grim trigger strategy for the repeated game for which playing p in every period is a subgame perfect equilibrium provided the discount rate is small enough.

This is bad news in a way. It means that game theory may not give sharp predictions about likely outcomes. Almost anything can be rationalized as a NE in infinitely repeated games.

Section 9.5. Extensive Form Games

One complication we should consider is that many games are played sequentially. For example, one player may move first and then be followed by a second player who gets to see what the first agent did before deciding on his own action.

We can represent this as an **extensive form game**. This looks like a kind of tree formed of **decision nodes** and **actions**. In the example below, the **initial node** belongs to the student since he decides to follow one of two **decision branches** (cram or drink) before any other player has a chance to take any action.

Once he chooses a branch, the game is reduced to a **subgame**. For example, if the student decides to drink, we get the bottom subgame with the professor's decision node having two branches (test and no test). By choosing to drink, the student has reduced the set of possible outcomes from four to two **terminal nodes** (or **payoff nodes**). If the professor chooses to test, the payoffs are $(-30, -30)$, if he chooses no test, the payoffs are: $(0, -20)$.

Note that the convention is that the payoffs are given in the order that the players move, so the first number indicates the payoff for the student, and the second is the payoff for the professor. Clearly, the professor should choose no test in this case. We see that if the student can make the first move, he can force the professor into a position where it is in his best interest to choose the student's most preferred outcome: drinking and no testing. This is a very sad outcome for professors. It seems there is no way to get students to study. What can be done about this?

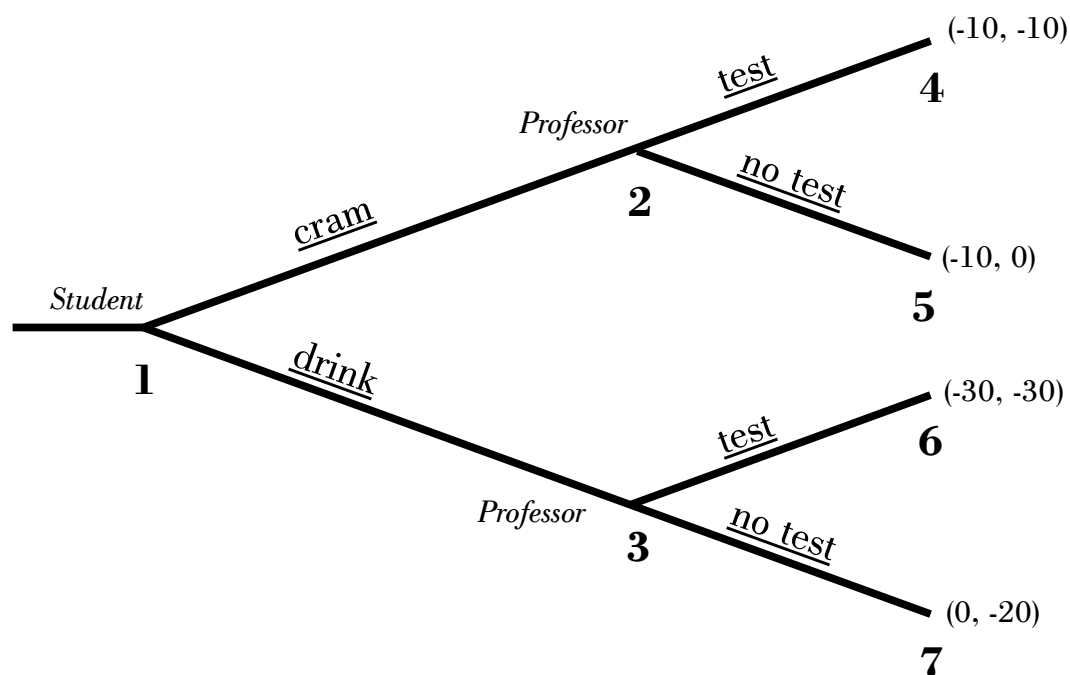


Figure 111: An Extensive Form Game

In an extensive form game, the student has the same choices as in the normal form version of the game. However, the professor's strategy set is a bit different. He can observe what the student does before he has to choose an action. Once the student chooses a strategy, he is committed and cannot change. The professor, therefore, can make his decision *contingent* on the student's decision. That is, the professor (or any agent with a decision node that is the initial node of a subgame) must specify the strategy he will choose in the event that any particular subgame happens to come up as result of decisions made by agents farther back in the game tree.

We will use a simple notation to indicate strategies here of the form $h \Rightarrow a$. By h , we mean a history of play that leads to a specific decision node for some agent, and by a , we mean an action that the agent chooses at this node. In the games above, there are three decision nodes. There is no history of play leading to the initial node, and so the students strategies are not contingent. If the students choose C , then we arrive at node 2, which is decision node for the professor. If the student chooses D , then we arrive at decision node 3. Thus, there are eight possible strategy profiles:

Student's Strategy	Professor's Strategy
C	$C \Rightarrow T, D \Rightarrow T$
C	$C \Rightarrow T, D \Rightarrow NT$
C	$C \Rightarrow NT, D \Rightarrow T$
C	$C \Rightarrow NT, D \Rightarrow NT$
D	$C \Rightarrow T, D \Rightarrow T$
D	$C \Rightarrow T, D \Rightarrow NT$
D	$C \Rightarrow NT, D \Rightarrow T$
D	$C \Rightarrow T, D \Rightarrow NT$

Table 23: Equilibrium price

Given the sequential nature of the game, what can we say about the likely equilibrium? Suppose the student chooses D . The professor has two best-responses:

$$C \Rightarrow T, D \Rightarrow NT \text{ and } C \Rightarrow T, D \Rightarrow NT.$$

In either case, drinking is a best-response from the student, so we have two NE strategy profiles already:

$$(D, C \Rightarrow NT, D \Rightarrow NT) \text{ and } (D, C \Rightarrow T, D \Rightarrow NT),$$

both of this lead to terminal node 7 with a payoff of $(0, -20)$. Suppose instead the student chooses C . The professor has two best-responses:

$$C \Rightarrow NT, D \Rightarrow T \text{ and } C \Rightarrow NT, D \Rightarrow NT.$$

However, cramming is a best-response only to $C \Rightarrow NT, D \Rightarrow T$. If the professor chooses $C \Rightarrow NT, D \Rightarrow NT$, the student should choose to drink. Thus, we had only one additional NE strategy profile:

$$(C, C \Rightarrow NT, D \Rightarrow T),$$

which leads to terminal node 5 with a payoff of $(-10, 0)$.

In words, if the student drinks, then the professor should not test, since this lowers everyone's payoff. What the professor says he will do out of equilibrium at node 3 does not make any difference to either payoffs or the best-response of the student. Thus, both $(D, C \Rightarrow T, D \Rightarrow NT)$ and $(D, C \Rightarrow T, D \Rightarrow NT)$ are NE strategy profiles. Suppose instead that the professor announces that if the student comes into class drunk, he will test, but if the student is sober and has studied, he will not test.

Under the Nash thought experiment, the student takes this strategy as given and chooses a best-response. If the student studies, then he is not tested and gets a payoff of -10 . If he drinks, however, he is tested and gets a payoff of -30 . His best-response is to study! Even better, if the student studies, the professor does not have to give the exam! The professor's payoff is therefore 0 , the best outcome possible for the professor. Thus, the strategy profile $(C, C \Rightarrow NT, D \Rightarrow T)$ is also a NE. We see that if play is sequential instead of simultaneous, we get three instead of one NE, and support two terminal nodes instead of one as possible equilibrium payoffs.

A Formal Definition of an Extensive Form Game

Extensive Form Game: $(\mathcal{I}, \mathcal{A}, \mathcal{D}, T, \text{Player}, \text{Action}, \text{Next}, S, F)$ where:

Players: $i \in \{1, \dots, I\} \equiv \mathcal{I}$

In the example: $\mathcal{I} \equiv \{ \text{Student}, \text{Professor} \}$

Actions: $a \in \{1, \dots, A\} \equiv \mathcal{A}$. This is the set of every action available to any agent at any node.

In the example: $\mathcal{A} \equiv \{ \text{drink}, \text{cram}, \text{test}, \text{no test} \}$

Decision Nodes: $d \in \{1, \dots, D\} \equiv \mathcal{D}$

In the example: there are $D = 3$ three decision nodes: $\mathcal{D} \equiv \{1, 2, 3\}$

Initial Node: Decision node 1, by definition.

In the example: the initial node belongs to the student who decides to either drink, or study.

Terminal Nodes: $t \in \{1+D, \dots, T+D\} \equiv \mathcal{T}$

In the example: there are $T = 4$ terminal nodes at the end of the game tree: $\mathcal{T} \equiv \{4, 5, 6, 7\}$. Terminal nodes have no successor nodes.

Player-Node Mapping: $\text{Player}: \mathcal{D} \Rightarrow \mathcal{I}$

In the example:

$\text{Player}(1) = \text{Student}$

$\text{Player}(2) = \text{Player}(3) = \text{Professor}$

Action-Node Mapping: $\text{Action}: \mathcal{D} \Rightarrow \mathcal{A}$

In the example:

$\text{Action}(1) = \{ \text{drinking}, \text{cram} \}$

$\text{Action}(2) = \text{Action}(3) = \{ \text{test}, \text{no test} \}.$

A Formal Definition of an Extensive Form Game (Continued)

Next Node Mapping: $Next: \mathcal{D} \times \mathcal{A} \Rightarrow \mathcal{D}$

In the example:

$$Next(1, \text{cram}) = 2, \quad Next(1, \text{drink}) = 3$$

$$Next(2, \text{test}) = 4, \quad Next(2, \text{no test}) = 5$$

$$Next(3, \text{test}) = 6, \quad Next(3, \text{no test}) = 7$$

(Note that each of these four strategies must specify a feasible action for the professor for each of his decision nodes, even if those nodes are not ever seen in equilibrium.)

Strategies: $s = \{s_1, \dots, s_I\} \in S_1 \times \dots \times S_I \equiv S$, such that:

$$(a) \mathcal{D}_i \equiv \{d \in \mathcal{D} \mid Player(d) = i\}$$

$$(b) \mathcal{A}_i \equiv \{a \in \mathcal{A} \mid \exists d \in \mathcal{D}_i \text{ with } a \in Action(d)\}$$

$$(c) s_i: \mathcal{D}_i \Rightarrow \mathcal{A}_i \quad \forall d \in \mathcal{D}_i, s_i(d) \in Action(d)$$

In our example: $S_{Student}$ contains two strategies:

$$(i) s_{Student}(1) = \text{drink}$$

$$(ii) \bar{s}_{Student}(1) = \text{cram}$$

On the other hand, $S_{Professor}$ contains four strategies:

$$(i) s_{Professor}(2) = \text{test}, \quad s_{Professor}(3) = \text{test}$$

$$(ii) \hat{s}_{Professor}(2) = \text{test}, \quad \hat{s}_{Professor}(3) = \text{no test}$$

$$(iii) \bar{s}_{Professor}(2) = \text{no test}, \quad \bar{s}_{Professor}(3) = \text{test}$$

$$(iv) \tilde{s}_{Professor}(2) = \text{no test}, \quad \tilde{s}_{Professor}(3) = \text{no test}$$

(Note that each of these four strategies must specify a feasible action for the professor for each of his decision nodes, even if those nodes are not ever seen in equilibrium.)

Payoff Functions: $F \equiv (F_1, \dots, F_I)$ where $\forall i \in \mathcal{I}, F_i: S \Rightarrow \mathbb{R}$

In our example:

$$F_{student}(\bar{s}_{student}, \tilde{s}_{professor}) = -10$$

$$F_{professor}(\bar{s}_{student}, \tilde{s}_{professor}) = 0.$$

Subsection 9.5.1. Subgame Perfect Equilibrium

The next question is, which of these NE do we find more likely? By announcing that he will test if he sees that the student has been drinking, the professor is committing to punish students if they misbehave. However, to do so, he must also punish himself since then he must then give a test. Fortunately, the professor never has to carry out this threat in a NE since the best-response for the students is to cram.

Is the professor's threat credible? Suppose the student drank anyway. Would the professor really give a test? Remember that the student gets to move first. If the student shows up to class drunk the professor gets a payoff -20 if he gives up and does not test, but -30 if he carries through on his threat and tests. Unfortunately, if the student does not cram, giving him an exam will make him neither sober nor knowledgeable. It only decreases the welfare of everyone. Thus, the only rational response on the part of the professor is to regret the state of the nation's youth and not give the test (and perhaps get a drink himself).

Put another way, the contingent strategy that promises to test the students if they drink is not a **credible threat**. If we ever got to the **subgame** in which the students drank instead of cramming, it would not be a best-response for the professor to give an exam. The student, knowing this, would not believe the threat and would assume the professor would play rationally in every subgame. The professor, in turn, would not waste his time making pointless threats that the students would not believe in any event.

A Formal Definition of a Subgame

Subgame: Let $(\mathcal{I}, \mathcal{A}, \mathcal{D}, \mathcal{T}, \text{Player}, \text{Action}, \text{Next}, S, F)$ be an extensive form game with perfect information. For any $d \in \mathcal{D}$, $(\mathcal{I}^d, \mathcal{A}^d, \mathcal{D}^d, \mathcal{T}^d, \text{Player}^d, \text{Action}^d, \text{Next}^d, S^d, F^d)$ denotes the **subgame beginning at node d** where:

Players: $\mathcal{I}^d \equiv \{i \in \mathcal{I} \mid \exists d \in \mathcal{D}^d \text{ such that } i = \text{Player}(d)\}$

Actions: $\mathcal{A}^d \equiv \{a \in \mathcal{A} \mid \exists \bar{d} \in \mathcal{D}^d \text{ such that } a \in \text{Action}(\bar{d})\}$

Decision Nodes: $\mathcal{D}^d \equiv \{\bar{d} \in \mathcal{D} \mid \bar{d} \notin \mathcal{T} \text{ and } \exists a_1, \dots, a_K \in \mathcal{A}, \text{ for some } K \in \mathbb{N}, \text{ such that } \bar{d} = \text{Next}(a_K \dots \text{Next}(a_2, \text{Next}(a_1, d)) \dots) \cup d\}$

Terminal Nodes: $\mathcal{T}^d \equiv \{t \in \mathcal{T} \mid \exists \bar{d} \in \mathcal{D}^d, \bar{a} \in \text{Action}(\bar{d}) \text{ such that } t = \text{Next}(\bar{a}, \bar{d})\}$

A Formal Definition of a Subgame (continued)

Player-Node Mapping: $Player^d(\bar{d}) = Player(\bar{d}) \forall \bar{d} \in \mathcal{D}^d$

Action-Node Mapping: $Action^d(\bar{d}) = Action(\bar{d}) \forall \bar{d} \in \mathcal{D}^d$

Next Node Mapping: $Next^d(\bar{a}, \bar{d}) = Next(\bar{a}, \bar{d}) \forall \bar{a} \in \mathcal{A}^d, \bar{d} \in \mathcal{D}^d$

Strategies : $s^d \in \prod_{i \in \mathcal{I}^d} S_i^d \equiv S$ such that $\forall i \in \mathcal{I}^d$:

(a) $\mathcal{D}_i^d \equiv \{d \in \mathcal{D}^d \mid Player^d(d) = i\}$

(b) $\mathcal{A}_i \equiv \{a \in \mathcal{A} \mid \exists d \in \mathcal{D} \text{ with } a \in Action(d)\}$

(c) $s_i^d : \mathcal{D}_i^d \Rightarrow \mathcal{A}_i^d$ where $\forall \bar{d} \in \mathcal{D}_i^d, s_i^d(\bar{d}) \in Action^d(\bar{d})$

Payoff Function: $F^d(\bar{s}) = F(\bar{s}) \forall \bar{s} \in \mathcal{S}^d$

The idea is to use the additional structure of sequential games to **refine** the set of equilibrium outcomes by applying stronger rationality requirements. To be a Nash equilibrium, we required only that taking all the other agents' strategies for the game as a whole as fixed, each agent individually follows a best-response strategy.

A stronger behavioral requirement would be that taking the other agents' strategies for the game as a whole as fixed, each agent individually follows a best-response strategy in every subgame (including the game as a whole). In general, the purpose of **equilibrium refinements** is to get rid of Nash equilibria that are somehow noncredible or otherwise undesirable. The refinement discussed above is probably the most important. More concisely:

Subgame Perfect Equilibrium: A strategy profile $s \in S$ is a subgame perfect equilibrium if for every agent, taking the action required by s_i is a best-response in every game or subgame that starts for any of decision node owned by agent i .

You can see that SPE adds the behavioral assumption that agents do not believe that others will ever choose strategies that are not in their own best interests. In other words, agents think that other players will always choose a best-response and every stage of the game and do not find any threat to deviate from this to be creditable. Note that subgame perfection only makes sense in sequential games. There are no subgames in normal form games. Other refinements, however, are possible.

Subsection 9.5.2. Games with Imperfect or Incomplete Information

Above, we considered games with perfect information. By **perfect information**, we mean that when any agent is called upon to make a decision, he knows precisely where he is in the game tree. Put another way, agents know the exact history of play and the identity of the decision node at which the game has arrived.

Games with **imperfect information** have the property that agents may not be able to observe all the actions of agents higher in the game tree. Put another way, they may be uncertain about the history of play when called upon to make a decision. They will know what choices are available to them at the decision node, which may narrow things down, but this will typically not allow the agent to back-out the complete history of play.

Thus, for each agent we need to partition his decision nodes into **information sets**. When called upon to make a decision, agents will know which set the decision node belongs to, but not the precise node at which the player happens to be. To be logically consistent, the information partition needs the following properties: for any given agent, the information set structure must put each of the agent's decision nodes into one and only one information set, each information set can contain decision nodes of one and only player, and all the decision nodes in any given information set must have the same actions sets.

A Formal Definition of an Information Partition

Information Partition: $Info: \mathcal{D} \Rightarrow \{ \text{subsets of } \mathcal{D} \}$ such that:

- (a) $\forall i \in \mathcal{I}, \forall d \in \mathcal{D}_i$ it holds that $d \subseteq Info(d) \subseteq \mathcal{D}_i$
- (b) $\forall i \in \mathcal{I}, \forall d, \bar{d} \in \mathcal{D}_i, Info(d) \cap Info(\bar{d}) = \emptyset$ or
 $Info(d) = Info(d) \cap Info(\bar{d}) = Info(\bar{d})$
- (c) $\forall i \in \mathcal{I}, \cup_{d \in \mathcal{D}_i} Info(d) = \mathcal{D}_i$
- (d) $\forall i \in \mathcal{I}, \forall d, \bar{d} \in \mathcal{D}_i$, if $\bar{d} \in Info(d)$, then $A(d) = A(\bar{d})$

In our previous example of the student-professor game, the information partition is trivial since we have perfect information: $Info(1) = 1$, $Info(2) = 2$, $Info(3) = 3$.

Suppose that, even though the student moves first at the initial node, the professor cannot observe the choice the student ends up making before he must choose his own strategy (perhaps the student shows up to class wearing sunglasses). This changes the extensive form game a bit since the professor can not tell which subgame he is in. Graphically, we represent the information sets by dotted connections or enclosures. Below, node 2 and 3 are in the same information set for the professor. Formally, we represent this information partition as follows:

$$Info(1) = \{1\}, \quad Info(2) = \{2, 3\}, \quad Info(3) = \{2, 3\}.$$

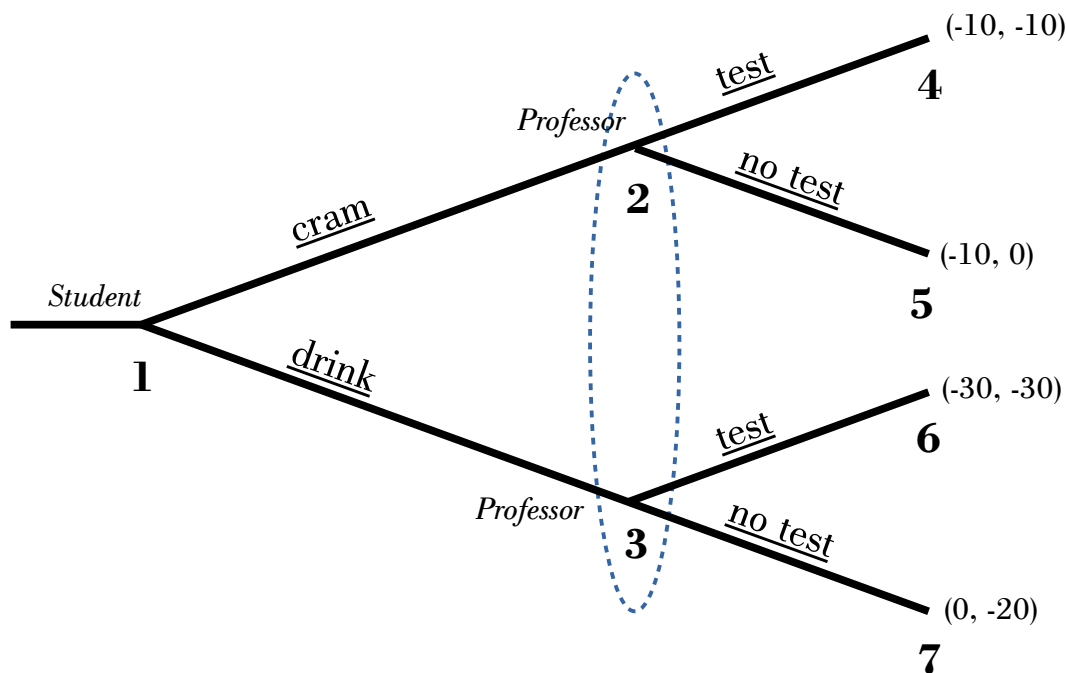


Figure 112: An Information Partition

We now must redefine the strategy space since agents may not know which decision node they have arrived at. Given this, strategies can only map *elements of the information partition* into actions instead of *decision nodes* into actions. Formally we add an extra condition to the definition of strategies.

A Formal Definition of Strategies with Information Sets

Strategies: $s = \{s_1, \dots, s_I\} \in S_1 \times \dots \times S_I \equiv S$ where

$$\forall i \in I, \mathcal{D}_i \equiv \{d \in \mathcal{D} \mid \text{Player}(d) = i\} :$$

- (a) $\mathcal{A}_i \equiv \{a \in \mathcal{A} \mid \exists d \in \mathcal{D}_i \text{ such that } a \in \text{Action}(d)\}$
- (b) $s_i: \mathcal{D}_i \Rightarrow \mathcal{A}_i \quad \forall d \in \mathcal{D}_i, s_i(d) \in \text{Action}(d)$
- (c) $\forall d, \bar{d} \in \mathcal{D}_i, \text{ if } \text{Info}(d) = \text{Info}(\bar{d}) \text{ then } s_i(d) = s_i(\bar{d})$

In our new example with an information partition, S_{Student} contains two strategies:

- (i) $s_{\text{Student}}(1) = \text{drink}$
- (ii) $\bar{s}_{\text{Student}}(1) = \text{cram}$

On the other hand, $S_{\text{Professor}}$ now contains only two strategies as well:

- (i) $s_{\text{Professor}}(2) = \text{test}, \quad s_{\text{Professor}}(3) = \text{test}$
- (ii) $\tilde{s}_{\text{Professor}}(2) = \text{no test}, \quad \tilde{s}_{\text{Professor}}(3) = \text{no test}$

Now there is only one NE which is also a SPE: $(D, C \Rightarrow NT, D \Rightarrow NT)$. The information partition has rendered it impossible for the professor to threaten to give a test only if the student drinks. He cannot tell what the student has done the night before, and so he can no longer make his strategy contingent on the student's (unobservable) actions. Given this information partition, his only best-response is no test.

Imperfect information, by definition is, limited to player's knowledge of the history of play. However, agents might also be uncertain about other aspects of the game. For example, I might not know what type of agent I am playing with. Perhaps "nature" or random chance chooses an opponent for me, and I am therefore not certain about my opponent's strategies, payoffs, or preferences.

We call this an **incomplete information game**. To see how this might affect equilibrium outcomes, suppose that I believe that my rival is a crazy serial killer. In this case my threat to shoot my opponent if he moves even slightly from a cooperative strategy might be credible. Of course, it would not be credible if I thought my opponent was a pregnant soccer mom holding two puppies. However, what if I claimed that I thought the soccer mom was actually a serial killer in disguise? My threat might then be credible, and I could force the soccer mom to be cooperative. Thus, a very important question is: what are my beliefs, and are they reasonable, consistent, and credible? What strategy profiles can be rationalized as equilibria depend on how this question is answered.

In a similar spirit, what if I believe that my opponent's beliefs about me affect the rationality, and therefore the credibility, of my strategies? For example, suppose I believe that my opponent

believes that if I ever agree to accept less than 70% of the surplus, I am a weak player who can be forced to accept 10% if I am punished for at least three periods. A best-response for me in this case is to refuse any offer of less than 70% since I “know” that if I ever do, I will get only 10% in the future. This is recursive, since it may also matter what I believe about what you believe about what I believe, about what you believe ...

Section 9.6. Examples of Extensive Form Games

The examples given in the previous section are extensive form versions of the prisoners' dilemma game. Of course the order in which agents move and the information they have at time of their own move affects the outcome. In this section, we show a several more examples of standard extensive for games in this literature.

Subsection 9.6.1. The Ultimatum Game

The **ultimatum game** has a very simple structure, but has been widely used in experimental economics. The game has two players. The first (Dick Tator) chooses an integer between 0 and 100 called his Bid. Note that this means that his strategy set consists of 101 possible actions. If we allowed fractional bids, his strategy set could be infinite or a continuum. The second player (Owen Cash) then can either accept or reject the bid. If he accepts, Owen gets an amount equal to the bid, while Dick keeps what left after the bid is subtracted from 100. If Owen refuses, both get nothing.

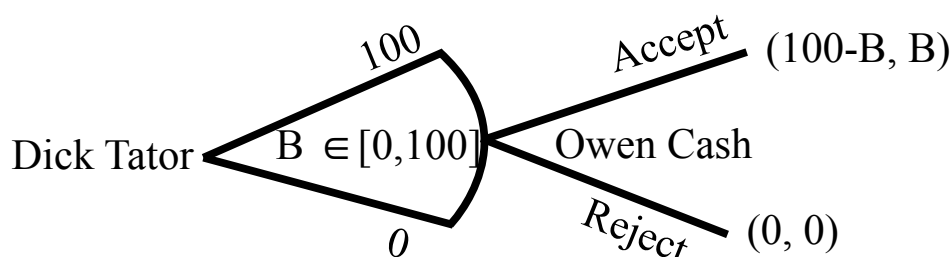


Figure 113: Ultimatum Game

To solve this, apply the logic of backwards induction. Suppose Owen has received a bid from Dick of B . Rejecting gets Owen nothing, but accepting gets him a payoff of B . Thus, if $B > 0$, accepting is a dominant strategy for Owen. Knowing this, the best-response for Dick is to make the lowest acceptable bid, $B = 1$. Thus, the unique SPE is $B = 1$ giving payoffs of $(99, 1)$. (There are other Nash equilibria. Can you find them?)

Experimentalists have tested this in real-world with many types of players. The finding is that in practice, the bid seems to average around 30, and this is typically accepted. Bids that are much below this are rejected. This is surprising since there is no benefit from rejecting a bid of 20, for example. The explanation is that while agents like money, they also like being treated fairly. If the bidder offers them too small a share of the payoff, they feel badly used, and so would prefer to spend 20 in order to punish the bidder then to accept 20 and have to put up with injustice. In other

words, the game as written above may not capture the payoffs correctly because of preferences for fairness and justice. Bidders, on the other hand, know receivers are often willing to give up a small amount of payoff to punish a greedy opponent. Thus, they offer enough so that punishment will not be an attractive option.

Subsection 9.6.2. The Cake Eating Game

Suppose there is a cake with N slices. Each period t , one slice is eaten by a mouse leaving $N - t$ to divide in period $t + 1$. Agents make alternating offers for how to divide the cake. If the offer is accepted, the division is made, and each agent walks home with this share. If the offer is rejected, then one slice disappears, and the other agent makes an offer.

Thus, in period t , one agent offers S slices to this opponent, and if it is accepted, the offer agents get a payoff of $R \equiv N - t + 1 - S$. Below we show the very last part of the game where all offers have been rejected. At the last subgame, there is one slice left and agent A can offer either 1 or 0 to agent B . If $S = 1$, then the residual cake goes to agent A so the payoff is $R = 0$. If agent A offers $S = 0$ and this is accepted by B , then $R = 1$ is the payoff to A . It is a weakly dominated strategy for A to offer $S = 0$. Agent B gets $R = 0$ if he accepts but also gets 0 if he refuses since then the last slice is eaten and the game is over. Assume that he rejects this offer.

Based on this we can use backwards induction. If there are two slices of cake, then B can offer 0, 1, or 2 to A . Offering 1 or 2 will be accepted since otherwise we get to the subgame in which all agents get 0. Clearly $S = 1$ is the best-response. Going back one period to when there were 3 slices, A should offer B , $S = 2$, which is accepted since it is better than he gets in the 2 slice subgame. Going back one more period to when there were 4 slices, B should offer A , $S = 2$, which is better than he gets in the 3 slice subgame. Finally, going back to the beginning when there are N slices, if N is even, then each agent gets $\frac{1}{2}N$ slices, and if N is odd, A gets $\frac{1}{2}N - \frac{1}{2}$ and B gets $\frac{1}{2}N + \frac{1}{2}$ slices. Thus, as N gets large, we converge to equal division of the cake.

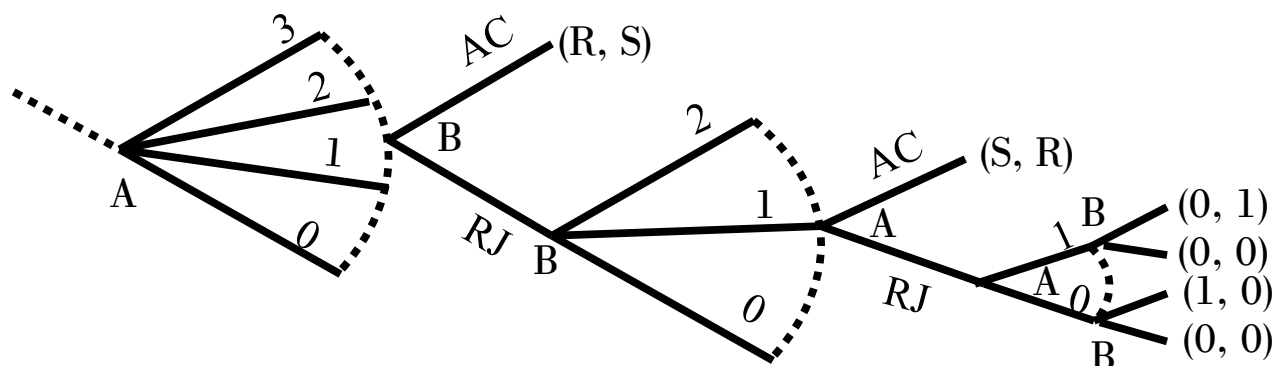


Figure 114: Cake Eating Game

This game can also be set in continuous time, or with an amount of money or capital that depreciates at a certain rate r each period.

Subsection 9.6.3. Centipede Game

Below, we have a finite period **centipede game**. It is called this because of how it looks, a long body and many legs. An example of this is a **war of attrition**. Suppose that there is a prize worth 100 that two agents competed over. This prize lasts six periods and then disappears. Agents decide to fight in the war of attrition or to quit.

The game for below shows these moves taking place sequentially, but we could also write this as a finitely repeated game with simultaneous moves. Fighting costs 10, and if you fight in one period, you must remain for the next period to receive the counterattack which also costs you 10. If any agent quits instead of fighting, he loses the prize and is out the cost of fighting the war. The other agent wins the prize, but has to pay the cost of his own war efforts out of his winnings.

In the figure below, A moves first. If he quits, no battle costs are accrued by either player, but player B gets the entire prize. If A fights, B then has a chance to decide if he wants to counterattack or quit. If he quits, he pays 10 for the cost of the first battle, and ends up with $100 - 10 = 90$. There is no reason to fight past the round six since the prize disappears and fighting then just generates symmetric losses of 10 for each round of continued play.

As before, start from the last round and work backwards. Agent B has a dominant strategy. quit. Fighting just leads to additional losses. Given this, agent A should fight in round five since he gets 50 in round six instead of the -40 he gets from quitting round five. In round four, agent B should quit since then he gets -30 instead of -50 . Thus, agent A should fight in round three. Agent B should therefore quit in round two, and agent A should fight in round 1. The payoff is therefore $(90, -10)$. This the only subgame perfect equilibrium. There are more NE, however. What would happen if the prize disappeared after seven instead of six rounds of play and so agent A had the final move? What would happen if the prize never disappeared?

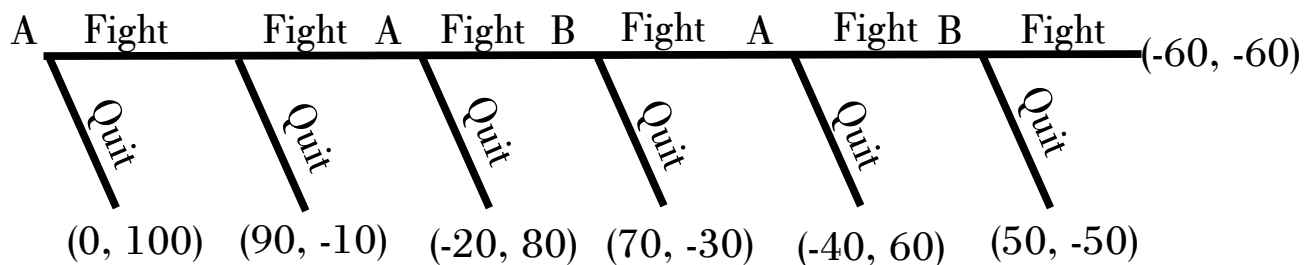


Figure 115: Centipede Game

Subsection 9.6.4. All Pay Auction: Dollar Bidding Game

This is similar to a centipede game. Two players bid successively in increments of $5¢$ for a one dollar bill. The catch is that all agents, winners and have to pay the most recent bid when the auction ends. This is called an **all pay auction**. There is no final period, so we cannot backward induct. What are the equilibrium?

There are no subgame perfect equilibria here. To see this, suppose that A promises to bid as long as B bids. Then it would not be a best-response of B to promise to bid as long as A does (why?). Thus, the “bid forever” pair is not a Nash equilibrium. Suppose that A promised to bid up to some $B > 100$ and then stop. Then it would not be a best-response for B to choosing he same strategy since it would be better to not bid at all, so it is not a NE.

What if A promised to bid up to $B < 95$ and then stop? Then B should bid 95 and win and get a payoff of 5. In this case, A should never have bid at all. What if A promised to bid 95 then stop. Then B could bid 100, and it then it would be a best-response in the subgame for A to bid 105 since if he loses he has to pay 95, while if he wins he as to pay 105 while gaining 100. In short, Not bidding at all is not an NE, since then one agent should bid once. Bidding forever is not a NE since then each agent would be better off not bidding at all. Finally, once an agent starts bidding, it is never subgame perfect to stop. What remains is the empty set.

There are, however, two Nash equilibrium. The first is for A to bid 5 and promise to bid forever, and for B to quit in response. The other is for B to promise to bid forever and for A to quit. Thus, $(0, 100)$ and $(95, 0)$ are both NE payoffs.

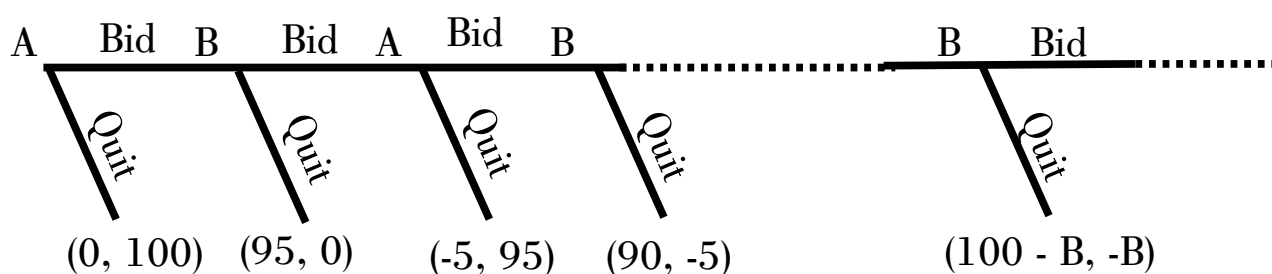


Figure 116: All Pay Auction

Section 9.7. Applications of Noncooperative Games

Subsection 9.7.1. Standards

A **standard** is a statement of commonly agreed upon properties for goods, services, procedures, measurement, formats, procedures, protocols, methods, and so on. Benefits include:

- Increased interoperability of manufactured goods.
- Increased ability to communicate information across providers and users of different goods and services.
- Increases confidence about and understanding of the quality of goods and services.
- Increases in network externalities and economies of scale.
- On the other hand, possible problems include:
- Reductions in product variety and consumer choice.
- Reductions in the flexibility in innovation of new products and services.

Standards are often created by technical or government bodies such as the ISO (International Organization for Standardization), ANSI (American National Standards Institute), IETF (Internet Engineering Task Force), AMA (American Medical Society), FCC (Federal Communications Commission) and the EU (European Union). They can also emerge as a result of a private firm dominating a sector and thereby setting a *de facto* standard. Standards exist in many different areas. For example:

- **Hardware:** Nuts and bolts, plugs and fittings, lumber dimensions
- **Infrastructure:** Railroad gauge, canal width, highway lane width, voltage, water pressure.
- **Professional:** Accounting practices, medical procedures, teacher certification.
- **Protocols:** FTP, TCP/IP, Ethernet.
- **Encoding Schema:** ASCII, Morse code, ZIP
- **Formats:** DVD, PDF, VHS, beta-max, .doc files,
- **Software:** DOS, OSX, Linux, function key and touch-screen gesture conventions.
- **Applications:** Social networks, Oracle, SQL
- **Ecosystems:** Google, Apple, AOL, Microsoft

In all of these cases, we see a greater or lesser degree of network externality. At some level, the more users on a given standard, the more valuable the standard becomes. Naturally, there is signifi-

cant effort spent by firms to monetize as big a share of this value as possible. In general, only a small number of agents participate in this competition, and so the tools of noncooperative game theory can be used for analysis.

Subsection 9.7.2. Standards Wars

There are many examples of contests between firms to dominate a product sector. MySpace Started in 2003 followed in 2004 by Facebook. MySpace maintained a strong market share advantage until 2008 when Facebook surpassed it. Since then, Facebook membership has dropped to tiny levels. Each of these companies has followed a strategy of giving their services away to users in hopes of gaining market share.

The first step in such a market is to gain **critical mass**. If a sufficient number of agents join a network or standard, the network can become self-sustaining. In other words, the network is large enough that when prices are set at a level that at least allow costs to be covered, the network continues to grow or reaches a steady state. It is able to avoid going into a death spiral where users leave, prices have to go up for a less valuable network product, causing more users to leave, and is on. Depending on the market and the nature of demand, one, several or no products may achieve critical mass. Other examples of standard wars are: Windows/Apple OS, Blu-ray/HD-DVD, Metric/US Customary Unit System, and Ethernet/Token Ring.

The implicit game has two or a small number of players with competing proprietary standards, and a set of customers who need to be convinced to join. This takes place over time, and customers can move from standard to standard according to their interests. To “win” such standards war, one player must gain critical mass and force the other players below critical mass so that they cannot make profits and must leave the industry.

This is a winner take all kind of sequential game. However, if the network externalities are quickly exhausted with only a fraction of the total market within a given standard, or if the tastes of agents over product characteristics are strongly differentiated and a network can cover its costs with only a fraction of the market, then the game may have equilibria in which a stalemate occurs, and no standard is able to drive all the others out.

For example, consider the Hotelling approach and suppose that the unit interval was a space of product characteristics. Let’s label the left side “usability” and the right side “flexibility”. Each agent has his most preferred compromise between these two elements. The decision to buy depends on the price of the good, how close the good is to the agent’s most preferred characteristic, and how many people are on the network.

Thus, suppose that we have 1,000,000 agents spread equally on the unit interval such that agent i ’s most preferred good has index number $i/1,000,000$. Then the reservation price (the marginal benefit an agent gets from consuming good $j \in \mathcal{I}$ might look like this for example:

$$MB_i(j, N_j) = C - (i - j)^2 + \left(\frac{N_j}{K}\right)^.5$$

where C and K are positive constants, and N_j is the number of agents in network j . As you can see, the value of a good decreases with the square of the distance from the most preferred good, but the network externality only goes up with the square root.

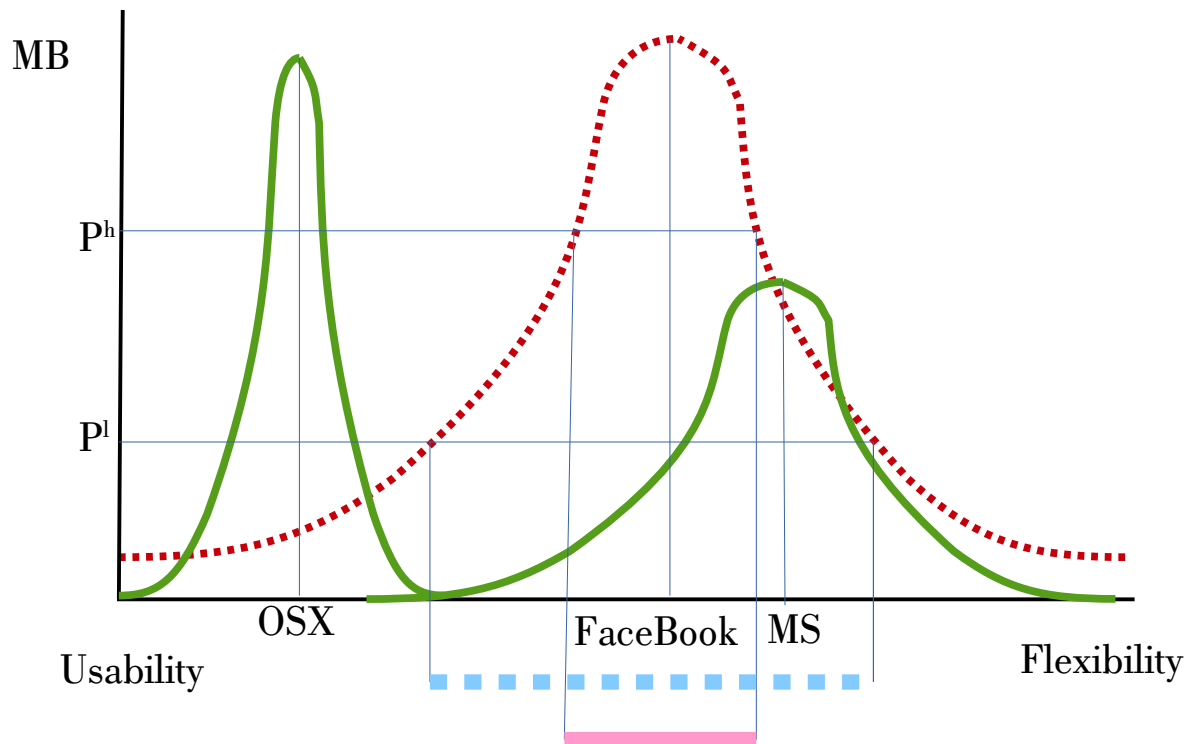


Figure 117: Network Externalities

In such a market, we could find a situation where some users valued usability very highly, but find that network externalities are quickly exhausted. Thus, we have the case illustrated of the two solid green curves.

Consumers on the usability side get a lot of benefit from consuming exactly what they want and are willing to pay a lot for it. Consumers who value flexibility turn out to value the network externalities more strongly as well. Thus, OSX users are fewer but are willing to pay lot for what they want. This gives Apple critical mass with small but loyal customer base. MS Windows users pay less since they get less benefit, but are larger in number. This allows Microsoft to obtain critical mass with lower prices charged to a broader customer base.

On the other hand, if agents value the size of the network more than its characteristics, we could have a situation illustrated by the red dashed curve in which one company like Facebook plays to

the average taste and leaves insufficient room on the edges of the characteristic space to allow a second network to exist with critical mass. Facebook wins the standards war in this case.

This figure also shows how prices affect the size of the product network. If Facebook sets a high price like P_h , only the smaller pink interval of agents still get enough benefit to be willing to join, while if it charges the lower price of P_l , the larger will be the dashed blue interval of agents who are willing to join the network. Also, note that if the price is set too high, additional room at the edges in the form of agents who choose not to join Facebook is created, and this might allow competing networks to gain critical mass.

The standards war discussed above takes place when there are two or more competing network products using **closed standards**. Such standards are generally based on propriety intellectual property. The rights holders may allow others to use their standard in exchange for licensing fees, or may encourage or even subsidize their use in order to gain market share at the expense of their competitors.

Open standards, in contrast, are created by government or professional groups. Such standards are free for all to use without permission or license and do not use proprietary IP. In some cases, private companies will form a consortium to develop a common standard for all to use. This creates concern if the standard needs to incorporate IP owned by individual consortium members.

To prevent hold-up and lock-in in the future, members are generally required to disclose any claim they might have to relevant IP to commit to license in on a **fair, reasonable, and non-discriminatory (FRAND)** basis. If a firm has a chance of winning a war with a closed standard, this is preferable. However, such commercial conflicts are costly since they drive prices down in the early network building phase, slow innovation in a sector as firms develop redundant IP to escape the patents and copyrights of their competitors, and cause consumers to hold back demand as they wait to see which standard they should buy into.

Thus, when there is one dominant firm with a closed standard, the rest of the competitors may benefit from jointly agreeing on common open standard so that all may share in the resulting network externalities and thereby increase their chances of beating the dominant player. Even if there is no dominant firm, all firms may see that it is advantageous to come together and agree on an open standard. This allows the entire sector to enjoy any network externalities and to grow that much more quickly.

Strategies:

There are three main categories of strategies that firms use in standards wars.

Pricing: an early mover in a market may sell their product for a very low price or even give it away. Firms may make introductory offers, send out coupons, or give complementary services to important users. The idea, of course, is to get the ball rolling. Once users start to join the network, its value grows. The earliest adopters bring significant value to the product and to the

firm. As the product becomes more known and more useful, prices can be raised. Firms must be careful in choose their pricing strategy since there is a tension between faster growth, and earlier profits.

Entrants in a network market with an incumbent dominant player may also try a low pricing strategy. Later entrants often have a better product since they can learn from the mistakes and successes of the incumbent firm, can use newer approaches and cutting-edge technology, and are not locked into legacy features and interfaces that may be costly to maintain and are demanded by only a minority of users.

A better product in combination with lower prices may break users away from the existing network, and with luck, displace the incumbent. This kind of low initial pricing is called **penetration pricing**. The opposite of this is called **skim pricing**. Here, an entrant or smaller incumbent may choose to offer a prestige product with a high price. This creates a smaller, but still viable network that can coexist with the larger one. Android used penetration pricing to enter the cell phone market, while Apple uses skim pricing and has become one of the most profitable companies on the planet.

Controlling the Narrative: Sophisticated users and potential users may worry about a number of things that might happen in the future as they decide which standard to adopt. In particular:

- Which standard is likely to win in the end?
- What will happen to prices when one or the other standard prevails?
- What will happen to quality of service or the nature of the product when one or the other standard prevails?
- How soon will the standard or the product be upgraded, and will the upgrade be backwards compatible with existing hardware or systems?

With these questions in mind, companies will claim that they are winning the market share war, that the other firm is on its last legs, that the standard of the other firm infringes on IP belonging to others, that the competing standard is technologically inferior, has a limited capacity to scale or upgrade, or is more likely to lock-in users.

Companies may also try to keep secret the timing of new releases or upgrades in hopes of preventing users from holding off purchase until the upgrade is released (which slows network growth in the interim). On the positive side, companies may decide to make commitments such as promising certain software will always be free, opening up the source code for inspection, building in paths that make it easier for users to switch vendors so that users do not fear lock-in, or signing long term, low price, contracts with larger users and then broadcasting this to other users suggesting that they get on the bandwagon before it is too late.

In a more formal sense, firms in a standards war attempt to alter the beliefs that other firms and consumers have about all aspects of the game including the payoffs that they and their opponent get

from playing certain strategies, the true state of the world, and also to close off or precommit to certain strategies in the future in order to make nonsubgame perfect play credible.

Influencing Other Players: Firms may benefit if they can alter the actions of agents besides their opponents in a standards war. For example, bribing governments, agencies, and other standards setting bodies to impose a standard the firm prefers, offering free software and hardware to universities to get students familiar with their products in hopes that continue to demand them in the future, subsidizing complementary industries such as content producers or app writers to grow the ecosystem and thus the network externalities, or trying to get their own proprietary IP incorporated into a standard.

Subsection 9.7.3. Auctions

Auctions play a role in many aspects of ICT. The FCC uses them to allocate new bandwidth to telecom companies, eBay uses them to find prices in thin markets, Amazon uses them to allocate “AdWords” to competing advertisers, companies, and government agencies use them to find lowest cost providers of goods and services, just to name a few examples.

Auctions are typically used in situations where markets are thin (only a few buyers and sellers) and where the goods being sold are unique, or are in fixed and low quantity. In these cases, buyers and sellers should not be price-takers. It is also difficult to figure out what market clearing prices might be in such markets. Auctions are a way of at least partially eliciting reservation prices which would otherwise be private information held by buyers.

In most cases, the sponsor of an auction is interested in getting the most advantageous price. If the sponsor is a seller, this means the highest price possible, and if the sponsor is a buyer (allowing companies to bid to provide goods and services), then finding the lowest possible price is the objective. A secondary objective, especially in auctions sponsored by government agencies, is social efficiency. The FCC might be happy to make less revenue for the cellular bandwidth it auctions if the buyers make the best social use of the spectrum. (Of course, what is “best” in this context is a matter of policy choice at least as much as economics.)

Auctions come in many forms. Below suppose there are two players A and B , who have reservation prices P_A and P_B . Suppose the minimum bid increment in the auction is b .

First Price Ascending Bid Auction

In a **first price ascending bid auction**, agents bid in turn and highest bidder wins and pays whatever he offered. One problem is that the bidders with the highest value only need to bid just a little more than the agent with the second-highest value (reservation price). Thus, the maximum price may not be obtained by the seller. We can write this kind of auction down as a variation on an unbounded centipede game. The dominant strategy is to keep bidding until you reach your reserva-

tion price, and then drop out. Such auctions are commonly used to sell works of art, used cars, cattle, foreclosed houses etc. These are “open out-cry” auctions in which agents hear the bids of their competitors and must respond immediately. These are not well suited to online or virtual situations, but one might wonder why they are so common in real-space.

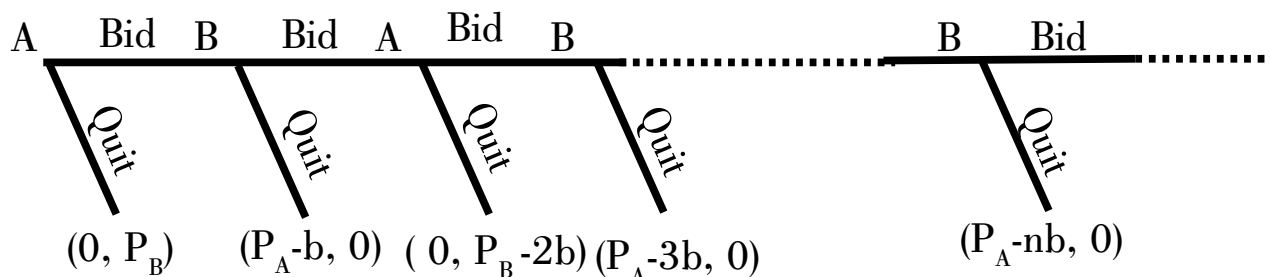


Figure 118: First Price Sealed Bid Auction

First Price Sealed Bid Auction

In a **first price sealed bid auction**, bids are not public (in contrast to the open out-cry auction discussed above). Agents place bids in envelopes (real or virtual) to be opened by the auctioneer at some set time. Charity auctions often use this technique. The most common use, however, is probably an inverted version in which a firm or government agency solicits sealed bids from firms to provide some set of specified goods and services. In this case, the lowest bid is the awarded the contract.

The hope is that since the firm with the lowest costs does not get to see the second-lowest bid, he may bid closer to his true value. In fact, the strategies for the bidders are quite difficult to calculate and depend on what the bidders guess about one another.

For example, I might shade my bid so that it is a little higher than my real costs thinking that my competitors have higher cost structures. If all firms think this, the bids may all come in substantially above costs. Of course, this is not the desired outcome for the agency soliciting bids. In the case where goods are being sold (think of leases for oil drilling rights on federal land) we have another problem. If all firms think that the others have bad technology and high costs, they may speculate that their competitors will submit low bids or not bid at all. The result is a set of sealed bids that are below the true value that the firms place on the leases. This reduces revenue, but even worse, it may misallocate the rights.

It is possible that a higher cost firm wins the auctions since low cost firms bid too low. This means that oil is extracted, but at a higher social cost, which is not in the national interests. Formally, the game is as follows for the two agent case:

Players or Agents: A, B

Strategies: $S_i \equiv \mathbb{N}$

Payoff Functions: $F_i(s_A, s_B) = \{ 0 \text{ if } s_i \leq s_j \text{ and } P_i - s_i \text{ if } s_i > s_j \}$

Although this is a one shot normal form game, it can't be expressed in a bimatrix format. This is because the strategies are a continuum, and so we would need an infinite number of rows and columns. Of course, there are usually more than two bidders which further increases the number of dimensions required. Nevertheless, it is easy to express this game with a simple payoff function.

Second Price Sealed Bid Auction

In a **second-price sealed bid auction**, the highest bidder wins, but pays only the second-highest bid plus one bid increment. This implies that it is a dominant strategy for agents to bid their true value. To see this, suppose that you bid lower than your reservation price. If you still won the auction, your payment would be one bid increment larger than the second-highest bid. Thus, you would pay the same as if you had bid your true reservation price.

Suppose instead you lost, and the winning bid was below your true reservation price. Then, you get zero payoff, while if you had bid your true value you might have won the auction purchased the good for less than your reservation price. This shows that you can never gain, and may lose, if you bid below your reservation price. The logic showing that bidding above your reservation prices is dominated by sincere bidding is similar. The most famous example of this are eBay auctions. Formally the game is as follows for the two agent case:

Players or Agents: A, B

Strategies: $S_i \equiv \mathbb{N}$

Payoff Functions: $F_i(s_A, s_B) = \{ 0 \text{ if } s_i \leq s_j \text{ and } P_i - s_j - b \text{ if } s_i > s_j \}$

You might ask yourself why we do not see second price open ascending bid auctions. What are the equilibrium properties we might expect to see here.

In both first and second price auctions, a **reserve price** is sometimes used. Typically, the winning bid must at least equal the reserve price or the auction is called off and the item left unsold. This price is usually kept secret from the bidders.

To how this works, suppose that agent A valued a painting at \$1000, while agent B valued it at \$2000. In a second price auction, agent B would win the painting and pay \$1001. Suppose, however, that the seller set a reserve price of \$1500. Since this minimum acts like another bid, agent

B would have to pay \$1501 when he bid his true value of \$2000. (You can verify that it is still a dominant strategy for all agents to bid their true values.)

The challenge for the seller is to get the reserve price right. The seller's target is to find a price between the first and second highest bid as close to the highest bid as possible. If the seller sets the reserve price below the second bid, it has no effect. If the reserve price is above the first bid, the good goes unsold and the seller does not profit. If the seller gets it right, however, he can increase his revenue above the second-highest bid.

In fact, it is a theorem that there exists a second price auction with some reserve price that gives the seller more revenue than a second price auction without any reserve price. Google, in fact, spends a great deal of effort using machine learning approaches to choose the best possible reserve price for its Google AdWords.

Dutch Auctions

In a **Dutch auction**, the auctioneer begins with a high asking price which is lowered until some participant is either willing to accept the auctioneer's price, or the seller's reserve price is reached. The winning participant pays the last announced price. This is also known as a clock auction or an open-outcry descending-price auction.

A variation on this is most useful when there are a fixed number of units to be sold and no single buyer is likely to want to take the entire lot. For example, fresh caught tuna and other seafood delivered for sale to the Tokyo fish market, and treasury bills (bonds) issued by the US government.

In this kind of auction, the bids consist of two numbers Q_i and p_i . Each bidder chooses a quantity he wishes to purchase, and a price per unit he is willing to pay and puts them in a sealed envelope handed to the auctioneer. The auctioneer orders the bids from highest to lowest price, and then sorts through the bids in this order until the sum of quantities is at least as big as Q , the amount to be sold. The auctioneer then gives each winning bidder the quantity he asked for at the lowest price in the winning pool. (Alternatively, each bidder might pay the price he bid.) If the lowest winning bidder asked for more than is left over, he just gets what is left over. If several bidders are lowest, then the remaining good is divided between them.

For example, suppose the bid submitted are as follows:

$$Q_1 = 10, \quad p_1 = 300$$

$$Q_2 = 40, \quad p_2 = 110$$

$$Q_3 = 50, \quad p_3 = 100$$

$$Q_4 = 25, \quad p_4 = 80$$

$$Q_5 = 10, \quad p_5 = 50$$

$$Q_6 = 20, \quad p_6 = 50$$

$$Q_7 = 60, \quad p_7 = 30$$

- If $Q = 100$, then $p = 100$ and first three bidders get all they asked for.
- If $Q = 120$, then $p = 80$, bidders 1, 2, and 3 get their orders filled and bidder 4 gets $Q_4 = 20$.
- If $Q = 140$, then $p = 50$, bidders 1, 2, 3, and 4 get their orders filled and bidders 5 and 6 get $Q_5 = Q_6 = 7.5$.

Other considerations:

- Does the number of bidders make a difference?
- Does it matter if this is a **private value** (each person values a painting differently) or common value (Drilling rights have the same value to all firms, but they are uncertain) auction.
- What happens if there are **multiple objects**? Bidding of the early objects may reveal something about the private value you place on future objects. This may work to the advantage or disadvantage of the bidders or the seller.
- What if there are multiple objects that become more valuable in groups. For example, owning the spectrum rights to have a cell phone network over the entire country is more than twice as valuable as having the rights to half the country.
- Reserve prices.
- Winner's curse.

Subsection 9.7.4. Lock-in

There are many situations in which agents have to make decisions that are difficult to reverse. In the most general sense, **lock-in** exists to the extent that it is costly to alter a choice. For example, if you marry someone, it is costly to undo your action. If you have kids, you are locked in even more strongly. Learning a skill set, especially acquiring what is called **specific human capital**, skills, or knowledge that are only of value to your current employer, produces employment lock-in. It is important, however, to distinguish switching costs from sunk costs.

For example, you may have spent 20 years with your current spouse or have invested a decade of weekends to keeping your boss's cobbled together IT system running. These are sunk costs and can never be recovered. In themselves, they do not create lock in. If your partner promises to take you for everything you have and turn your children against you, or you find that you are so specialized, that the next best job you can get is selling computers at Office Depot, then you will have to pay substantial costs if you try to switch, and so you are locked-in.

The ICT sector is rife with lock-in. Consider **ESM (enterprise systems management)** solutions, for example. ESM is top to bottom integration of an enterprise's information systems and management practices. It includes making IT systems interoperable, putting all sorts of data in compatible formats, implementing real time business intelligence, etc. This can be done using custom-built systems or putting together off-the-shelf elements from various vendors.

Aside from the obvious advantage that custom systems can be tailored to the exact needs of the company there are several reasons for companies to choose this path. Most SaaS and PaaS providers store data in their own proprietary way. If a company wanted to change vendors, extracting such data and building a new software system around it is an expensive and difficult task. In addition, employees get used to the work flow and interfaces of these proprietary systems. This also makes it costly to switch systems.

Similarly, the greater the degree of abstraction in the PaaS platform – for example special APIs that facilitate interactions between components like databases and email, proprietary libraries of code that the users can build applications with – the more difficult it is to move to a new provider. Thus, building ESM on more abstracted layers of XaaS makes users of cloud services more vulnerable to price increases. In addition, companies who build ESM from vendor's components are not in a good position to enforce high service quality.

More specifically, as technology and markets change, cloud providers may choose not to continue to support certain functions or features of their services that may be highly valued by a subset of customers. Updates to cloud systems may affect the way that they interact with the rest of a company's ESM solution and so crashes may result that require time to fix. In short, this approach leaves a company open to lock-in on many fronts.

Of course, the possibility of lock-in makes such cobbled together systems less valuable to users and so less profitable to providers. Thus, we might consider a game between cloud providers in which they choose how easy to make it for customers to cleanly move to another provider. While this would decrease their market power, it would increase their value and thus the price they could charge for services. Especially if service providers do not plan to take advantage of lock-in by suddenly switching from a low price/customer-base building phase to a high price/rent extraction phase, it would seem to be a dominant strategy to make it easy for customers to leave.

These same considerations give an advantage to rapidly growing companies and very large companies like Google, Oracle, and Amazon. Such companies are less likely to switch to a rent extraction phase since this would deter future customers. Growing companies get more benefit from main-

taining a good customer reputation, while large companies suffer more damage if they try to extract rent from a subset of customers.

Examples of lock-in

SIM locking: may be considered a vendor lock-in tactic, since phones purchased from the vendor will work with SIM cards only from the same network.

Gift certificates: are textbook examples of vendor lock-in as they can be used solely in the vendor's shops. Gift certificates are typically only worth their face price (no bonus credit is added), so generally, they do not represent any financial advantage over money.

Printer Manufacturers: using non-official cartridges will void and warranty Lexmark makes ink cartridges that contain an authentication system, the purpose of which is to make it illegal in the United States (under the DMCA) for a competitor to make an ink cartridge compatible with Lexmark printers.

Vacuum Cleaners: are only compatible with specific filter bags.

IBM: had significant lock-in of the punched card industry from its earliest days: before computers as we recognize them today even existed. From dominance of the card punches, readers, tabulators, and printers, IBM extended to dominance of the mainframe computer market, and then to the operating systems and application programs for computers. Third party products existed for some areas, but customers then faced the prospect of having to prove which vendor was at fault if, say, a third party printer didn't work correctly with an IBM computer, and IBM's warranties and service agreements often stipulated that they would not support systems with non-IBM components attached. This put customers into an all-or-nothing situation.

Microsoft: software carries a high level of vendor lock-in, based on its extensive set of proprietary APIs. The Windows API is so broad, so deep, and so foundational, that most ISVs would be crazy not to use it. It is so deeply embedded in the source code of many Windows apps that there is a huge switching cost to using a different operating system instead. Microsoft's application software also exhibits lock-in through the use of proprietary file formats. Microsoft Outlook uses a proprietary file format which are impossible to read without being parsed.

Apple Inc: digital music files with digital rights management were available for purchase from the iTunes Store, encoded in a proprietary derivative of the AAC format that used Apple's FairPlay DRM system. These files are compatible only with Apple's iTunes media player software on Macs and Windows, their iPod portable digital music players, iPhone smartphones, iPad tablet computers, and the Motorola ROKR E1 and SLVR mobile phones. As a result, that music was locked into this ecosystem and available for portable use only through the purchase of one of the above devices.

Sony: has used lock-in as a business tool in many other applications, and has a long history of engineering proprietary solutions to enforce lock-in. For many cases, Sony licenses its technology to a limited number of other vendors, which creates a situation in which it controls a cartel that collectively has lock-in on the product. Sony is frequently at the heart of format wars, in which two or more such cartels battle to capture a market and win the lock.

Examples of Sony's formats include:

- Audio Elcaset.
- Audio or computer data Minidisc and the related ATRAC3 encoding system.
- Super Audio CD.
- Betamax, Video-8, Hi8, Digital8, and MicroMV videotape formats.
- PlayStation Portable Universal Media Disc.
- Memory Sticks, used for a wide variety of applications in Sony products.

Proprietary Connectors: The reasons for such designs vary; some are intended to force customers quietly into a vendor lock-in situation, or force upgrading customers to replace more components than would otherwise be necessary.

Razors and Blades: business model involves products which regularly consume some material, part, or supply. In this system, a reusable or durable product is inexpensive, and the company draws its profits from the sale of consumable parts that the product uses. To ensure the original company alone receives the profits from the sales of consumable, they use a proprietary approach to exclude other companies. Ink-jet computer printers are a common example of this model.

Loyalty Programs: One way to create artificial lock-in for items without it is to create loyalty schemes. Examples include frequent flier miles or points systems associated with credit card offers that can be used only with the original company, creating a perceived loss or cost when switching to a competitor.

Subsection 9.7.5. Hold-up

Hold up arises when part of the return on an agent's relationship-specific investment is ex post expropriable by his trading partner. The hold-up problem has played an important role as a foundation of modern contract and organization theory, as the associated inefficiencies have justified many prominent organizational and contractual practices.

The hold-up problem is a situation where two parties may be able to work most efficiently by cooperating, but refrain from doing so due to concerns that they may give the other party increased bargaining power, and thereby reduce their own profits. When party A has made a prior commitment to a relationship with party B, the latter can 'hold up' the former for the value of that commitment. The hold-up problem leads to severe economic cost and might also lead to under-investment.

Subsection 9.7.6. Bank Runs and Bandwagons

A final economic issue of concern is that bankruptcy, reorganizations, or even changes in business focus on the part of cloud service providers pose significant risks for users. In the worst case, users may lose data stored in the cloud. At a less catastrophic level, features, and functionality may be dropped or no longer supported, service levels of companies distracted by internal problems may decline. Seeing this, users who can, may withdraw their data and find new providers of services. This further weakens the company, and induces more customers to leave. The dynamics are very much like a bank-run. There may be a role for regulations and standards to give users some assurance of stability. This also may provide a strategic advantage for established firms with a long run reputation for reliability.

Examples

- Web Storage, Document Sharing
- Social Platforms
- S-Curve, Tipping, Saturation
- Band Wagon Effect
- Innovation and Investment (First mover)
- Cryptocurrency

Subsection 9.7.7. Experimental Approaches

Experimental economists put the predictions of noncooperative game theory to the test. Typically, experimentalists get together a group of subjects (such as college students) and ask them to play a game in which they get real rewards depending on what they choose to do. One classic experiment is to ask pairs of randomly and anonymously matched subjects to play the ultimatum game.

In the **ultimatum game**, one player offers a division of a fixed payoff to the other. The second player has only two strategies: accept the proposed division and enjoy the payoff, or refuse it and receive nothing. In this sequential game, it is a dominant strategy to accept any positive offer. Accepting gives the second player something positive which is strictly better than refusing and receiving nothing. Knowing this, the first player should make the lowest positive offer possible, since by backward induction, the second player will accept it and this maximizes the first player's payoff. This is the only SPE.

Surprisingly, the experimental evidence suggests that people do not play the SPE. Instead, the average offer is about 30 % and this is usually accepted. Offers lower than this are often rejected.

What is going on? This evidence seems to say that agents playing a very simple game do not behave rationally, at least as defined by economists.

One explanation is that the experimenters have not captured the payoff function correctly. The monetary payoffs are only a fraction of the utility that agents get from playing the game. Agents also care about being treated “fairly” and not being taken advantage of. Thus, when the first player offers the second five cents and keeps 95 cents for himself, the second player really sees the payoffs as: say yes, get five cents but accept being treated unfairly, or, say no, get no money, but take 95 cents away from the person who tried to take advantage of me. The receiving agent might well prefer the second option and so this would not be an irrational choice.

If the first agent realizes that the second agent can credibly refuse any “unfair” division, it becomes an SPE to make the smallest offer that passes the fairness test. This begs the question of what agents think is a fair division. In the next chapter, we begin to think about this problem this formally.

Glossary

Battle of the Sexes: A game with two players in which there are two Nash equilibria which are not equally preferred by the agents. Thus, there are two stable outcomes, one of which favors the first agent and the other of which favors the second agent.

Bimatrix Form: A representation of a normal form game with two players. The strategies for one player are listed on the top row and those for the other player on the left column. This forms a matrix in which the payoffs for any specific choice of strategies for the row and column player are listed as an ordered pair with the row player's payoff listed first (or sometimes in the upper left-hand corner).

Continuation Strategy: Given $s \in S$ and a subgame starting at $d \in \mathcal{D}$, we call $s^d \in S^d$ the **continuation strategy from node d** if $\forall i \in \mathcal{I}$, and $\forall \bar{d} \in \mathcal{D}_i$, $s_i^d(\bar{d}) = s_i(\bar{d})$. That is, a continuation strategy describes how agents propose to play if a certain subgame happens to be reached.

Coordination Game: A game in which all the Nash equilibria require that agents choose the same strategy. These different Nash equilibrium outcomes may or may not favor different agents. It may even be the case that all agents agree on the ranking of these equilibrium outcomes. The battle of the sexes is an example of a coordination game.

Credible Threat: A strategy that agent promises to play in a subgame, should it be reached, which is a best-response in this subgame.

Disoordination Game: A game in which all the Nash equilibria require that agents choose the different strategies. These different Nash equilibrium outcomes may or may not favor different agents.

Dominant Strategy Equilibrium (DSE): We say $s \in S$ is a dominant strategy equilibrium if $\forall i \in \mathcal{I}$, $s_i \in S_i$ is a dominant strategy.

Dominant Strategy: We say $s_i \in S_i$ is a dominant strategy if $\forall \bar{s}_i \in S_i$ and $\forall s_{-i} \in S_{-i}$, $F_i(s_i, s_{-i}) \geq F_i(\bar{s}_i, s_{-i})$. Informally, a strategy is dominant for a player if it gives him the highest possible payoff in *every* possible situation.

Equilibrium Refinement: The set of Nash equilibria of a game may be large and some of these outcomes may be unrealistic or otherwise undesirable. To reduce the set of equilibria, restrictions may be placed on how agents play (for example, subgame perfect equilibrium) or how they form or use exceptions about the nature of uncertain parts of the game, how agents might play off the equilibrium path and other aspects of the game and players (for example, sequential equilibrium). These restrictions are justified by arguing that they are more refined encodings of what it means to be rational. Thus, an equilibrium refinement is a tightening of the definition of rational play that reduces the number of equilibrium outcomes.

Experimental Economics: Laboratory or field interactions between economists and subjects to see how agents makes decisions is the real world.

Extensive Form Game: An extensive form game is defined by the following list of items: $(\mathcal{I}, \mathcal{A}, \mathcal{D}, \mathcal{T}, \text{Player}, \text{Action}, \text{Next}, S, F)$, that is, a set of agents, actions, decision nodes, and terminal nodes, and a list of permissible actions at each node, linkages between nodes, strategies, and payoffs for each agent at each terminal node. This can be represented as a kind of tree formed of decision nodes and actions.

Folk Theorem: Consider an infinitely repeated game. Let p be a mixed strategy profile for the one-shot, stage game such that the expected payoff for each agent is strictly larger than his minimax payoff. Then there exists a grim trigger strategy for the repeated game for which playing p in every period is a subgame perfect equilibrium provided the discount rate is small enough.

Grim Trigger Strategy: In a grim trigger strategy, some stage game strategy profile (pure or mixed) is identified. An agent using the grim trigger plays his part of this strategy profile as long as all other agents do likewise. However, the first time any other agent defects from this profile, the rest of the agents play a punishment strategy from the next period on until the end of time that enforces the minimax payoff on the defector.

Imperfect Information: A situation where agents are not able to observe all the actions of agents higher in the game tree. Thus, agents must choose actions in subgames without knowing the history of play that lead them to the current stage of the game.

Incomplete Information Game: A situation where agents aspects of the game other than the history of play. For example, agents may not know type of agent they are playing with.

Information Set: In a game with imperfect information, agents do not know the history of play with certainty. Thus, when called upon to make a decision, agents may not know for certain the exact decision node at which they have arrived. The set of decision nodes an agent thinks he might be possible at any given decision point is his information set. This forms an **information partition**: $\text{Info} : \mathcal{D} \Rightarrow \{ \text{subsets of } \mathcal{D} \}$ that must satisfy a set of requirements that assure that it is logically consistent.

Minimax Payoff: The minimax payoff of an agent is found by first considering the maximum payoff a player could get for any given set of strategy choices of his opponent(s). Second, the minimum over all of these maximal payoffs is found and used as the punishment payoff. The idea is that the minimax is the worst payoff that the rest of the players are certain that they can impose on any defecting agent even when the agent chooses a best-response to minimize the damage.

Mixed Strategy: A strategy in which an agent randomizes over set of pure strategies in a normal form game, or over the set of permissible actions at any given decision node in an extensive form game. In other words, an agent chooses a lottery over all of his available actions whenever he is called upon to make a decision.

Nash Equilibrium (NE): A strategy profile $s \in S$ is a Nash equilibrium if $\forall i \in \mathcal{I}$ and $\forall \bar{s}_i \in S_i$, $F_i(s_i, s_{-i}) \geq F_i(\bar{s}_i, s_{-i})$. Nash equilibrium requires that all agents choose a strategy that is a best-response to the strategies chosen by the other agents in the game. That is, taking the actions of all the other agents as fixed, there is no choice that improves an agent's payoff. Formally,

Node: A decision or outcome points of an extensive form game that mark the starting point of a **subgame** or the ending of then entire game. The **initial node** is the first decision node of the game and typically leads to generic **decision nodes** at which a specific player is called upon to choose a strategy from some set. In turn this takes the game down a **decision branch** to other decision nodes controlled by other players and eventually to some **terminal node** which describes a **payoff** that each agent in the game receives.

Noncooperative Game Theory: A game in which agent cannot make binding commitments. For example, contracts may not be enforceable.

Normal Form Game: A simultaneous move, one-shot game with any number of players.

Payoff: $F \equiv (F_1, \dots, F_I)$ where $F_i: S \Rightarrow P$ and P is some payoff space (possibly a Euclidean space, but not necessarily). In the study group game, payoffs are letter grades, for example. We call F_i and F the **payoff function for agent i**, and the **payoff function for the game**, respectively.

Perfect Information: A situation in which all agents know precisely where they are in game tree whenever they are called upon to make a decision, Put another way, agents know the exact history of play and the identity of the decision node at which the game has arrived.

Player: An agent who makes strategic choices and receives payoffs. Usually, they are denoted by an index set: $i \in \{1, \dots, I\} \equiv \mathcal{I}$.

Prisoners' Dilemma Game: A game in which the dominant strategy equilibrium (called the "non-cooperative equilibrium") gives all agents in the game a lower payoff than if they choose another strategy profile (called the "cooperative strategy" in this context). In such games, agents who play rationally get lower payoffs than agents who somehow can agree to pay strategies that are contrary to their own rational self-interests.

Pure Strategy: A strategy in which an agent chooses a single strategy with certainty in a normal form game, or a single permissible action at any given decision node in an extensive form game.

Repeated Game: A special case of a game in extensive form a single normal form game is played several times or even infinity in succession. Thus, in each round of the repeated game, players are faced with the same normal form game which is called the **stage game**.

Strategies: In a normal form game, a strategy is an action available to an agent. In an extensive form game, a strategy is a plan that indicates what action an agent will choose in every decision node that is assigned to him.

Strategy Profile: A list that gives a strategy choice for each agent in a game.

Subgame Perfect Equilibrium (SPE): A strategy profile $s \in S$ is a *Subgame Perfect Equilibrium* if $\forall d \in \mathcal{D}$ (including $d = 1$), the continuation strategy s^d for the subgames beginning at node d , $(\mathcal{I}^d, \mathcal{A}^d, \mathcal{D}^d, \mathcal{T}^d, S^d, Player^d, Action^d, Next^d, F^d)$ is a Nash equilibrium and so $\forall i \in \mathcal{I}^d$, and $\forall \bar{s}_i^d \in S_i^d$, $F_i^d(s_i, s_{-i}) \geq F_i^d(\bar{s}_i, s_{-i})$. Informally, this requires that all agents choose strategies that are best-responses for the whole game, as well as in every subgame, even if those subgames are never seen in equilibrium play.

Subgame: A subgame starting from some node $d \in \mathcal{D}$ in an extensive form game (called the **super game**, in this context) is formed including only the nodes that can be reached starting from node d along with the players, actions, payoffs, etc. associated with these downstream nodes. All decision nodes above d , or that can only be reached by starting higher in the game tree are discarded.

Tit-for-Tat Strategy: A strategy in a repeated game in which each agent plays whatever his opponent played in the previous period.

Ultimatum Game: A game in which one player offers a division of a fixed payoff to the other. The second player has only two strategies: accept the proposed division and enjoy the payoff, or refuse it and receive nothing. In this sequential game, it is a dominant strategy to accept any positive offer. Accepting gives the second player something positive which is strictly better than refusing and receiving nothing. Knowing this, the first player should make the lowest positive offer possible, since by backward induction, the second player will accept it and this maximizes the first player's payoff. This is the only SPE.

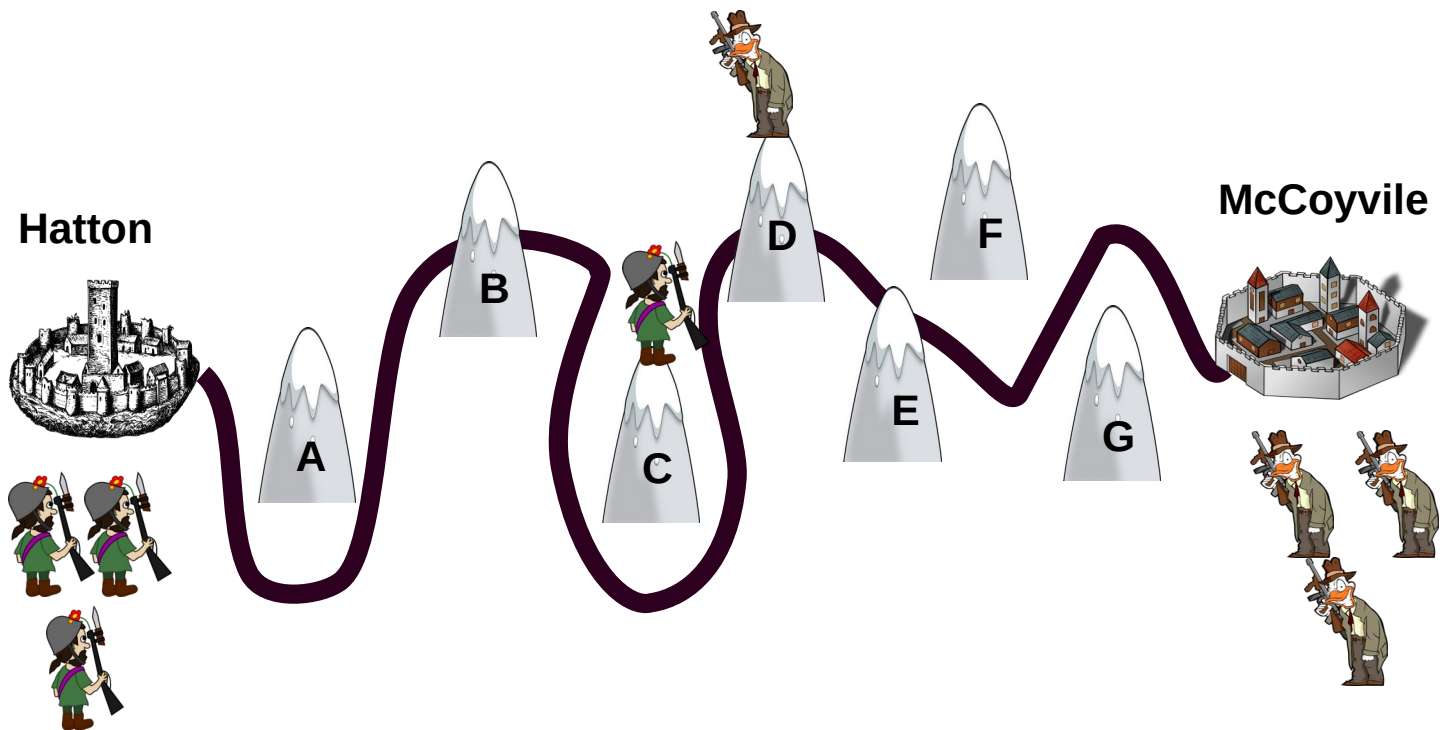
Zero Sum Game: A game in which the sum of the payoffs over all the agents is the same regardless of the strategy profile. Thus, in every outcome of the game, anything gained by one agent is exactly offset by losses spread over all the other agents.

Problems

1. There are three passengers on a ship sailing across the Atlantic who are carrying large sums of money. Suppose passenger A has \$1000, passenger B has \$1500 and passenger C has \$2000. They each put their money in separate envelopes and give them to captain who puts them in a chest for safe keeping. One night the ship is caught in a storm and the chest is broken open. The captain picks up all the money the next day, but does not know how much each of the three passengers gave him originally. He proposes the following solution. He asked each passenger to write down how much money they originally put into their envelope. They are to do this at the same time, without speaking to one another. If the numbers the agents write down add up to the total amount of money the captain picked up (assume he managed to find all \$4500) then the captain gives each passenger what he claims he put in. Otherwise, the captain throws the money into ocean and the passengers get nothing.
 - a. Write this out as a formal one-shot game. Note that there are three players (the passengers) and that the captain is just a mediator who defines the strategy spaces and payoff functions of the agents.
 - b. Is truth telling a Nash equilibrium. Why or why not?
 - c. Are there other Nash equilibria. If so, what are they?
2. Bill Clinton and George Bush Sr. hate each other. For security reasons they must live in one of only three places: New York, Chicago, and San Francisco. New York is about 1000 miles from Chicago, and 3000 from San Francisco. Chicago and San Francisco are about 2000 miles apart. For each thousand miles of separation from Bush, Clinton gets 5 units of utility. He gets an additional bonus of 10 units of utility if he lives in New York. Bush gets 5 units of utility for each thousand miles of separation from Clinton, plus a bonus of 5 units if he lives in his favorite city: Chicago.
 - a. Write this game in normal (matrix) form. Identify the Nash equilibria, if any.
 - b. Suppose Bush got 15 units of extra utility from living in Chicago. What are the Nash equilibria in this case (if any)?
3. In the town of Smallville, Ohio has one police officer, and ten dishonest citizens who are potential criminals. Each criminal has a choice of committing one crime each year or staying honest. If a citizen commits a crime, the probability of getting caught is equal to one divided by the total number of crimes committed $\frac{1}{\#crimes}$. If a criminal is caught, he goes to jail and suffers a loss equal to \$10,000 of income. If he gets away with the crime, he gets extra income \$2,000. Suppose potential criminals are risk neutral and seek only to maximize expected income.
 - a. Write this out as a normal form game. Be sure to include all three elements. (Note: it is best to write this as a payoff function rather than a matrix.)
 - b. What are the Nash equilibria of this game, if any?

c. How would this change if the criminal suffered a loss of \$100,000 if they were caught?

4. The Hatfields and McCoys are two families that have been fighting a feud in the West Virginia mountains for decades. They mostly live in Hatton and McCoyville. There are seven strategic mountain tops at ten mile intervals on the route between the two towns that have a commanding view of the countryside. Any family that controls one of these mountains can keep the other bottled up and enjoy the use of all the land between it and their hometown. These mountains are named A, B, C, D, E, F , and G . Both families have a brilliant strategist on their side and they each notice the following:



- Controlling a mountaintop N miles from their hometown gives them extra territory to trap and hunt and is worth $\$1000 \times N$ in extra income to the family. Thus, if the Hatfields control Mount C and the McCoys control Mount D , the Hatfields make \$30,000 since Mount C is 30 miles from Hatton, while the McCoys make \$40,000 since Mount D is 40 Miles from McCoyville.
- It is dangerous to get outflanked. If a family takes over a mountaintop only to find that they are trapped between the rival family's town and a mountain their rivals have taken over, the rival town will attack the outflanked mountaintop outpost while and their own family will be delayed getting help to them by the rival mountaintop outpost. As a result, outflanked family will be run-off and lose the outpost they have constructed. This costs the family \$20,000 in lost guns and timber. Note that if one family is outflanked then the other must be as well. Thus, overreaching in this way results in both mountaintop outposts being destroyed.
- If both families try to take over the same mountaintop, a war breaks out which costs each family \$100,000.

- a. Write down the set of players, strategies and the payoff function. Note that one of the strategies is to stay in town and do nothing.
 - b. Write this down in bimatrix form.
 - c. What are the NE if any? What are the DSE if any?
 - d. Is this a coordination game, a discoordination game, both, or neither? Support your position.
5. Poor farmer Jones has had a very bad year. In fact, so bad that he went bankrupt and the bank is about to auction off his farm. Suppose that there are two bidders Abby Hoffman and Bo Derick (A, and B). They each place the following values on the farm which are known only to themselves and not by the other bidders:

$$V_A = 1100, V_B = 1500.$$

The rules of the auction are these: 1) each bidder places a bid in a sealed envelope, 2) the highest bid wins the farm, 3) the price the winner pays for the farm is the bid of the next highest bidder, not his own bid. This is called a second price auction. Assume that if the bids are the same, Bo Derick wins the auction.

- a. Write the payoff function for each agent. Note that the payoff is zero if an agent loses and is the difference between his value and what he pays for the farm if he wins. Thus, you should specify a function for each agent of the form $F_A(Bid_A, Bid_B) = \dots$ which gives the payoff for every possible combination of bids.
 - b. Is bidding honestly a Nash equilibrium?
 - c. Can an agent ever do better than by bidding honestly? Why or why not?
6. Barges in ancient China used to be pulled through canals by human power. The bargemen received a bonus if they managed to get their cargo to its destination quickly. The problem was that each bargeman was tempted to shirk (by only pretending to pull at the tow-ropes). One man shirking did not result much difference in how long the journey took, but would make life much easier for the shirker. Of course, if everybody else were shirking, one man pulling hard would not make much difference in speeding the journey either. What is likely to happen in equilibrium? What is the name for this type of situation? Suppose that the laziest bargeman makes the following suggestion: Give me my full share of the bonus, and instead of pulling the barge I will follow the barge with a whip and will lash anyone I catch shirking. Should the bargemen accept this offer? Why or why not?
7. In the game below, identify the following:
- a. Nash equilibria, if any:
 - b. Dominant strategy equilibria, if any:
 - c. Dominated strategies, if any:

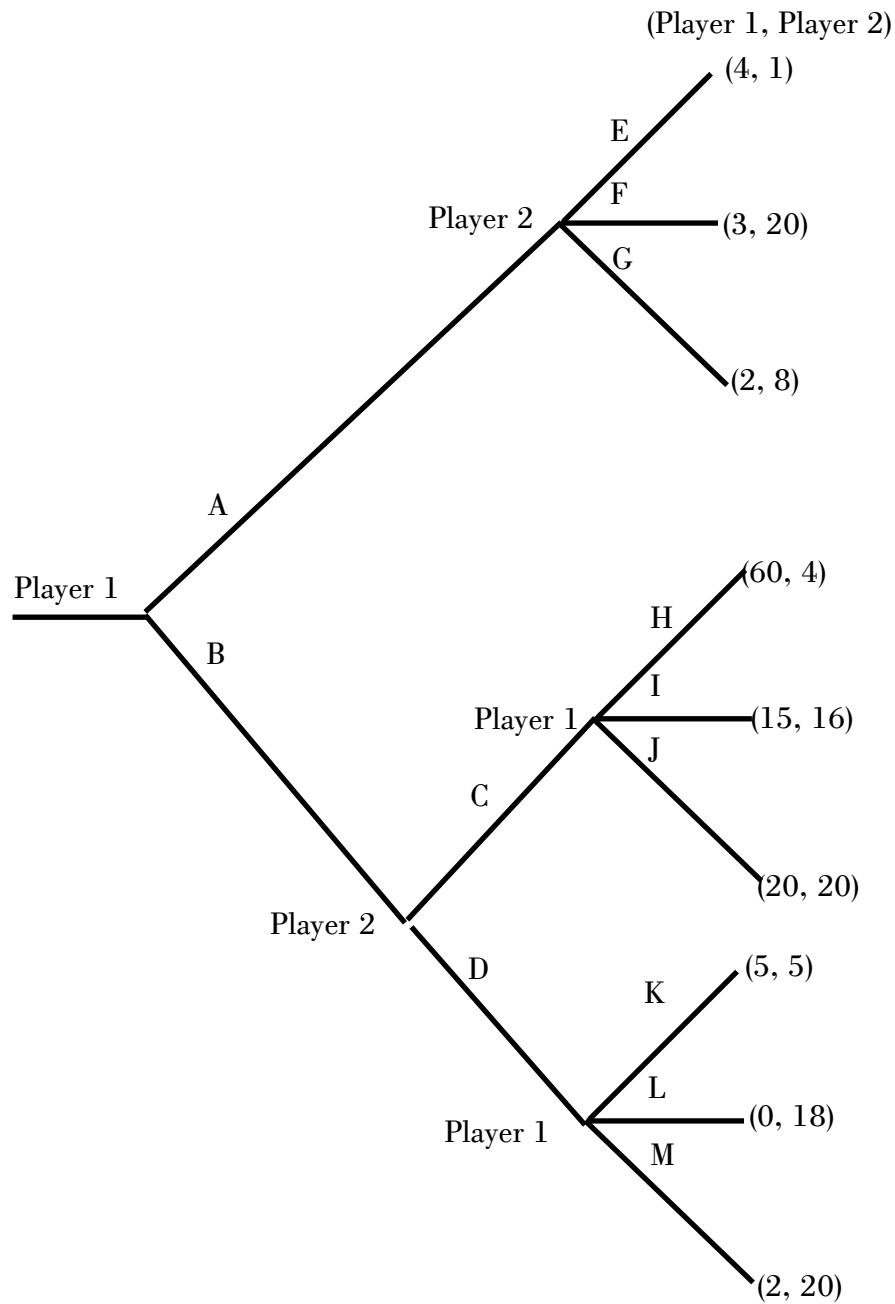
	Player 2			
		Left	Middle	Right
Player 1	Strategy A	25	4	34
		34	-3	5
	Strategy B	34	18	34
		122	7	87
	Strategy C	20	24	30
		1	3	49
	Strategy D	3	78	89
		33	22	12
	Strategy E	89	1	90
		89	58	28

8. The board of directors of Consolidated Fraternity Paddles consists of three members, Adam, Barbara, and the chairman of the board, Carl. Currently, there are 100 shares of the company held in common, and the board has decided to distribute them to the individual board members. They will be distributed according to the following rules:

- Adam, the most junior member of the board, will go first and propose a division of the shares (for example, 10 for him, 30 for Barbara, and 60 for Carl). The board members vote to either accept or reject. If the proposal is rejected, Adam is tossed off the board, and we go to the next round of voting.
- Barbara, the next most junior member of the board, will go next and propose a division between her and Carl. If her proposal is rejected, Barbara is tossed off the board, and all the shares go to Carl.
- Voting rules: Voting is by majority rule. **In the case of a tie vote, the proposal also passes.** Board members always vote no unless they know they are strictly better off over the whole game by voting yes. Shares cannot be divided, so giving an agent half a share is not allowed.

Adam wants to find the division that is most favorable to him that will also pass the majority vote in the first round. Using backwards induction, figure out what division Adam should propose. Hint: what is the outcome of the game if only one board member (Carl) remains? What does this imply about the outcome when both Carl and Barbara remain (remember that in the case of a tie vote, the proposal also passes and that board members vote no unless they know they are strictly better off by voting yes.) Using this, figure out what proposal Adam should make. Give your reasoning.

9. Consider the sequential game below. Using the rule of “look forward and reason backward”, also called backwards induction, find the subgame perfect equilibrium. Briefly explain your reasoning.



Chapter 10. Cooperative Game Theory

Section 10.1. What is Cooperative Game Theory?

Cooperative game theory addresses the question of how agents might come to an agreement over the sharing of some surplus. In this setting, agents may have utility or payoff functions, but they do not choose strategies. Instead, they agree on a rule or solution concept that suggests an allocation on which to settle. To be more precise, a **solution concept** is a mapping from a space of economies or games into the feasible set. Such mappings might be single valued or set valued. Of course, there are many different solution concepts available. We therefore make lists of properties called **axioms** that encode fair or desirable properties that solution concepts might have. If we can show that a given list of axioms is satisfied by one and only one solution concept, we say that the solution is **characterized** by this list. We will be more formal about this below.

Section 10.2. The Core

Probably the most important cooperative solution is called the “core”. Roughly speaking, we say an allocation over agents is a **core allocation** if it is stable against all possible coalitional deviations.

To illustrate this more formally, consider a **game in characteristic function form**:

Players: $i \in \{1, \dots, I\} \equiv \mathcal{I}$

Coalitions: $c \subseteq \mathcal{I}$

Characteristic Function: $V: \mathcal{C} \Rightarrow \mathbb{R}$

Allocations: $x^c \in \mathbb{R}^{|c|}$

where $|c|$ is defined as the **cardinality**, or the number of elements in a finite set c .

Coalitions are collections of players. The coalition consisting of all agents in the population, $\mathcal{I}$, is called the **grand coalition**. The set of all subsets of the grand coalition is denoted $\mathcal{C}$, and so $c \in \mathcal{C}$ for any specific coalition.⁷

The **characteristic function** simply gives \$120. the total transferable wealth that any specific coalition can produce. For example, suppose that the coalition of Bob, Bill, and Sue can start a lemonade stand and make \$120. In this case $V(\{Bob, Bill, Sue\}) = 120$.

We call x^c an **allocation for the coalition c** . This is a list of payoffs $(\dots x_i, x_j, x_k \dots)$ where $\{\dots, i, j, k, \dots\} \in c \in \mathcal{C}$.

Blocking Coalition: We say that an allocation for the grand coalition $x^{\mathcal{I}}$ is **blocked** by an allocation x^c for coalition c if:

- (1) $\sum_{i \in c} x_i^c < V(c)$ (x^c is feasible for c)
- (2) $\forall i \in c, x_i^c \geq x_i^{\mathcal{I}}$ (all agents in the blocking coalition are at least as well off)
- (3) $\exists j \in c, x_j^c > x_j^{\mathcal{I}}$ (some agents in the blocking coalition are strictly better off)

In this case, we would say that c is a **blocking coalition**, and x^c is a **blocking allocation**.

⁷ More formally $\mathcal{C} \equiv \mathcal{P}(\mathcal{I})$, that is, the **power set** of the set of agents, $\mathcal{I}$.

Core Allocation: $x^{\mathcal{I}}$ is a **core allocation** if it cannot be blocked by any $c \in \mathcal{C}$. We also say $x^{\mathcal{I}}$ is **in the core of a game**.

The intuition is that allocations in the core satisfy a minimal notion of social stability. We would never expect that any agent or group of agents would consent to be exploited if they had a more attractive opportunity. This would require that agents irrationally ignore their own best interests. Thus, the test the core imposes is that an allocation should be proof against group deviations. In short, we do not believe that an allocation would ever be seen in the real world if some group could do better by forming their own club and going off on their own. Any potential for group-wise improvement undermines the social order.

It is important to note that although the motivation for the core involves agents and groups of agents choosing to leave the grand coalition, the property “is proof against group deviation” or “unblockable” is defined without references to strategies or any noncooperative game. It is a property that can be checked by looking only at the feasible sets for groups and subgroups in the economy. In short, it is a “cooperative” notion in the sense that it considers what is feasible, and from this, defines what is desirable.

Subsection 10.2.1. Pareto Optimality and Individual Rationality

One very important property of the core is that it satisfies an axiom called “Pareto optimality”. At the most abstract level, suppose the feasible set for any coalition is given by the correspondence:

$$FS: \mathcal{C} \Rightarrow \text{allocations over coalitions}.$$

The allocation space could be almost anything, a consumption bundle in $\mathbb{R}^N$ for each agent a coalition, a job, a partner, or an office, a payment, and so on, for each agent. In the special case of characteristic functions form games, the feasible set for any coalition $c \in \mathcal{C}$ is:

$$FS(c) \equiv \{X^c \in \mathbb{R}^{|c|} \mid \sum_{i \in c} x_i^c \leq V(c)\}.$$

The **weak** and **strong Pareto sets** are defined as follows:

$$WPO(\mathcal{I}) \equiv \{\bar{x} \in FS(\mathcal{I}) \mid \nexists \hat{x} \in FS(\mathcal{I}) \text{ such that } \forall i \in \mathcal{I}, \hat{x}_i >_i \bar{x}_i\}$$

$$SPO(\mathcal{I}) \equiv \{\bar{x} \in FS(\mathcal{I}) \mid \nexists \hat{x} \in FS(\mathcal{I}) \text{ such that } \forall i \in \mathcal{I}, \hat{x}_i \geq_i \bar{x}_i \text{ and } \exists j \in \mathcal{I} \hat{x}_j >_j \bar{x}_j\}.$$

A second important property of the core is that agents must do better at a core allocation than they can do as individuals by themselves. Formally, we say that an allocation $\hat{x} \in FS(\mathcal{I})$ is **indi-**

vidually rational (IR) if $\nexists i \in \mathcal{I}$ and $\bar{x}_i \in FS(i)$ such that $\bar{x}_i \succ_i \hat{x}_i$. We define the **individually rational set** as:

$$IR(\mathcal{I}) \equiv \{ \hat{x} \in FS(\mathcal{I}) \mid \nexists i \in \mathcal{I} \text{ and } \bar{x}_i \in FS(i) \text{ such that } \bar{x}_i \succ_i \hat{x}_i \}.$$

Section 10.3. Examples of the Core

Subsection 10.3.1. Example 1

The figure below is a three-dimensional representation of the core allocations implied by the following characteristic function:

$$V(1) = V(2) = V(3) = 0$$

$$V(1,2) = V(2,3) = V(1,3) = 1.5$$

$$V(1,2,3) = 3$$

Each axis measures the payoff of one of the agents.

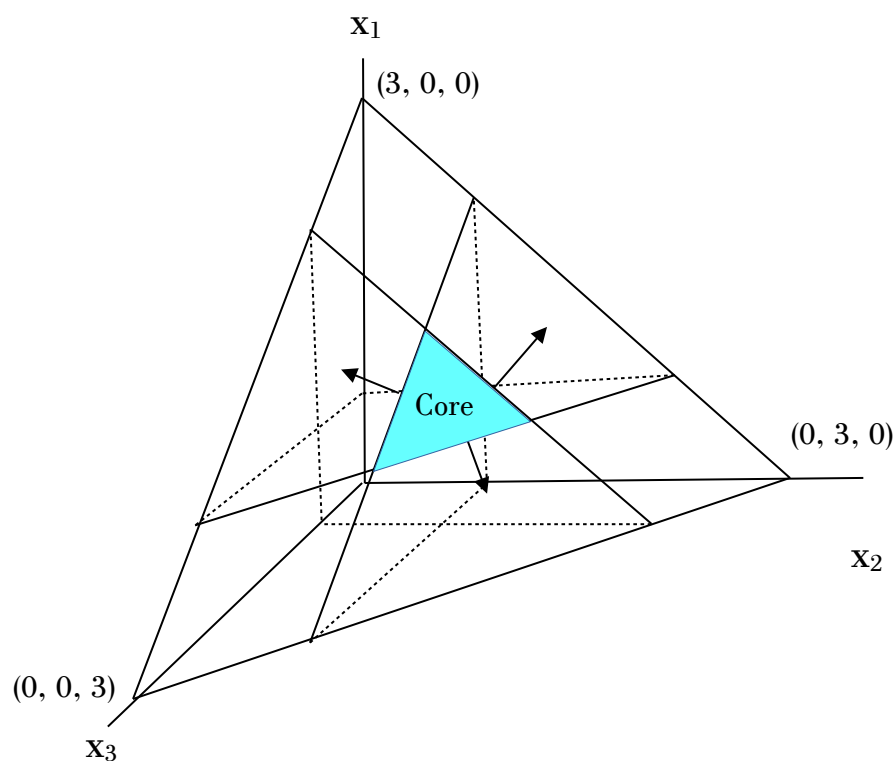


Figure 11: Core allocations of a three player game

To see this, first note that by IR, all agents must get nonnegative payoffs. Thus, all potential core allocations must be in the positive orthant. Otherwise, the allocation could be blocked by a one-person coalition.

By SPO, the sum of the allocations over agents must equal three. If the sum is less, then it is feasible to give *all* agents more, and so the grand coalition could block. This implies that the allocation must lie in the *2-dimensional simplex* shown (the triangle formed by points $(3,0,0)$, $(0,3,0)$, and $(0,0,3)$).

Finally, $x_1^I + x_2^I \geq 1.5$ or else the coalition of agents 1 and 2 could block. Since payoffs over all agents must add up to 3 by SPO, this means that $x_3^I \leq 1.5$. Graphically, this is shown by the triangle with dotted lines that intersects agent 3's payoff axis at 1.5. Since agent 3 must be closer to his origin than this, any allocation on the SPO simplex farther away from $(0,0,3)$ than this could be blocked by the coalition of agents 1 and 2. The other triangles show similar areas that can be blocked by the other two-person coalitions.

From all this we conclude that anything in the smaller triangle in the middle of the PO simplex is a core allocation since no coalition would be able to block any of these allocations. For example:

$(x_1, x_2, x_3) =$	$(1, 1, 1)$	core allocation
	$(0, 1.5, 1.5)$	core allocation
	$(.5, 1.2, 1.3)$	core allocation
	$(-1, 2, 2)$	not a core allocation
	$(.25, .25, 2.5)$	not a core allocation
	$(0, .5, 1)$	not a core allocation

Subsection 10.3.2. Example 2

We do a similar construction below with the following characteristic function:

$$V(1) = V(2) = V(3) = 0$$

$$V(1,2) = V(2,3) = V(1,3) = 2.5$$

$$V(1,2,3) = 3$$

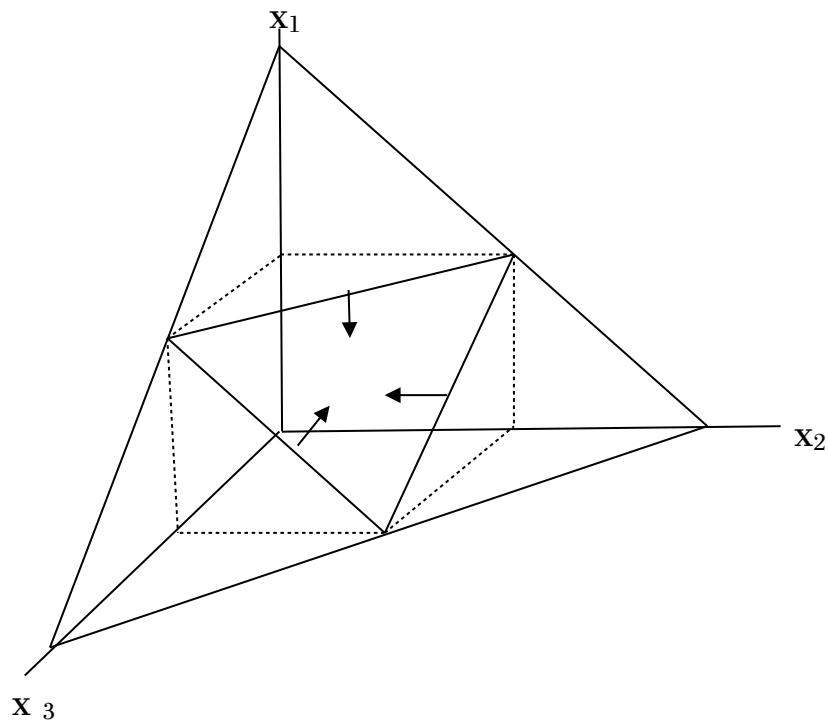


Figure 119: A three player game with an empty core

In this case, the core is empty. You can verify the following:

$$(x_1, x_2, x_3) = \begin{array}{ll} (1, 1, 1) & \text{not a core allocation} \\ (1, 1.5, .5) & \text{not a core allocation} \\ (0, 2.5, 2.5) & \text{not a core allocation} \end{array}$$

Subsection 10.3.3. Superadditivity

What is going on here? When do there exist stable allocations, and when is the core empty? This has a technical answer, but to be less formal, the core exists if the game is *super-additive* in the sense that there are increasing returns to coalitional size. More formally:

A game $(V, \mathcal{I})$ is **super-additive** if $\hat{c}, \bar{c} \subset \mathcal{C}$ where $\hat{c} \cap \bar{c} = \emptyset$, $V(\hat{c} \cup \bar{c}) \geq V(\hat{c}) + V(\bar{c})$.

It is immediate that if a game is super-additive, then:

$$\forall \hat{c}, \bar{c} \subset \mathcal{C} \text{ where } \hat{c} \cap \bar{c} = \emptyset, \frac{V(\hat{c} \cup \bar{c})}{|\hat{c} \cup \bar{c}|} \geq \frac{V(\bar{c})}{|\bar{c}|}.$$

In general, if a game is super-additive, then it has a nonempty core (superadditivity is a sufficient condition). This is because superadditivity implies that there is an increasing per capita return to coalition size. This implies that large coalitions are powerful since it is impossible for smaller coalitions to give as large an allocation per capita to its members. On the other hand, if a game is sub-additive, then small coalitions do better on average. In fact, the single person coalitions do the best of all. There is nothing the grand coalition can do to keep agents from defecting since it does not have enough to pay each agent as much as he could get on his own. The core is therefore empty.

Finally, note the following:

- Characteristic functions could be anonymous instead of nonanonymous. In the example above, V maps coalitions into payoff $V: \mathcal{C} \Rightarrow \mathbb{R}$. Thus, we need to know the name of each member of a coalition before we can tell the value of the coalition. We could instead have V map the size of the coalition into payoffs $V: \mathbb{N} \Rightarrow \mathbb{R}$.
- We could also go the other way and define a characteristic function that was not only nonanonymous, but also nontransferable. In this case $V: \mathcal{C} \Rightarrow \mathbb{R}^{|c|}$. For example, if $c = \{i, j, k\}$ then when coalition c gets together, agent i gets a nontransferable payoff of 3, agent j gets 2, and agent k gets 9.
- Finally, we do not need a characteristic function at all to define the core. All we need to know is what is feasible for each coalition and subcoalition. This implies how well subcoalitions can do when they defect and this in turn is enough to define the core.

Section 10.4. The Shapley Value

From the discussion above, we see that the core is a selection from the set of all feasible allocations that satisfy the “unblocked” condition. More formally, the core is a set-valued correspondence from some space of problems (perhaps economies or games) to the set of feasible allocations which satisfy some set of conditions or axioms.

We might desire a more precise recommendation for how to choose over the set of feasible allocations. Of course, we would have to consider a different set of axioms to do so. In this section, we discuss the **Shapley value** (or simply the **value**) which is probably the second most commonly used cooperative solution rule after the core. Of course there are many other alternatives such as the **kernel** or the **nucleolus**, but this is beyond our scope.

The objective of Shapley was to figure out how to value each agent’s contribution to a game or economy. To illustrate, let us return to the simple case of games in characteristic function form with transferable utility. In the abstract, a solution to a class of TU games is a mapping:

$$F : \text{Coalitions} \times \text{Games} \Rightarrow \text{Allocations}.$$

Thus, given a set of all the agents in the grand coalition $\mathcal{I}$ and a characteristic function of a game $V : \mathcal{C} \Rightarrow \mathbb{R}$, $F(\mathcal{I}, V) = x \in \mathbb{R}^I$ suggests an allocation, a payoff, or a value, for each agent $i \in \mathcal{I}$. The specific solution that Shapley proposed is the following for each agent $i \in \mathcal{I}$:

$$F_i^{SV}(\mathcal{I}, V) = \sum_{c \in \mathcal{C}} \frac{(|c|-1)!(I-|c|)!}{I!} (V(c) - V(c/i)).$$

where $c!$ denotes the **factorial** of c , is the number of distinct **permutations** of a set consisting of $|c|$ distinct elements.

The motivating idea is that the value of a player equals the average of his marginal contribution to each coalition he might potentially join. Thus,

- $V(c) - V(c/i)$ is what coalition c would lose if agent i were to disappear from the earth. Note that if coalition C does not contain agent i in the first place, it does not lose anything since in this case: $V(c) - V(c/i) = V(c) - V(c) = 0$.
- $I!$ is the total number of different orders in which agents could enter as the grand coalition forms. That is, the total number of permutations of the list of all agents.
- $(|c|-1)!(I-|c|)!$ is total number of these permutations in which any given agent would be exactly the $|c|^{th}$ agent to join any given coalition of size $|c|-1$ of which he is not already member.

For example, suppose that $\mathcal{I} \equiv \{1, 2, 3\}$ and $\hat{V}(c) = |c|^2$. Then there are $I! = 3! = 6$ possible orders of arrival as the grand coalition forms. These are given in the list below. Consider agent 1 for a moment. Notice that he arrives first, second or third exactly $\frac{1}{3}$ of the time.

$$\begin{aligned} &\{1, 2, 3\} \\ &\{1, 3, 2\} \\ &\{2, 1, 3\} \\ &\{3, 1, 2\} \\ &\{2, 3, 1\} \\ &\{3, 2, 1\} \end{aligned}$$

- Of the two cases in which agent 1 arrives first, there are two ways to complete the grand coalition. This does not matter for calculating the Shapley value of agent 1 in our example since in both cases the coalition goes from the empty set to 1 and thus: $\hat{V}(\{1\}) - \hat{V}(\emptyset) = 1^2 - 0^2 = 1$.
- Of the two cases in which agent 1 arrives second, there are two ways that the grand coalition could form. In general, this might matter since in one case, agent 1 joins agent 2 and in the other he joins agent 3. In our example, however, it turns out not to matter, since $\hat{V}(\{2, 1\}) - \hat{V}(\{3\}) = \hat{V}(\{3, 1\}) - \hat{V}(\{2\}) = 2^2 - 1^2 = 3$.
- Of the two cases in which agent 1 arrives last, there are also two ways that the grand coalition could have formed. However, in both cases, agent 1 joins the coalition of agents 2 and 3 to complete the grand coalition. How much value agent 1 adds to an existing coalition does not depend on the order of arrival of the members of existing coalition, only on which agents are present when agent 1 arrives. In our example: $\hat{V}(\{3, 2, 1\}) - \hat{V}(\{2, 3\}) = 3^2 - 2^2 = 5$.

Putting this together, there are a total of seven possible coalitions in $\mathcal{C}$. Agent 1 is in only four of these. In the remaining three, “removing” agent 1 from the coalition has no effect since he was not there in the first place. Thus, his marginal contribution is zero in these cases. The Shapley value is therefore the sum of the weighted marginal contributions given in the table below.

$$\begin{aligned}
c = \{1\}: & \quad \frac{(1-1)!(3-1)!}{3!}(\hat{V}(\{1\}) - \hat{V}(\{\emptyset\})) = \frac{(0!)(2!)}{3!}(1^2 - 0^2) = \frac{2}{6}(1) = \frac{2}{6} \\
c = \{2\}: & \quad \frac{(1-1)!(3-1)!}{3!}(\hat{V}(\{2\}) - \hat{V}(\{\emptyset\})) = \frac{(0!)(2!)}{3!}(1^2 - 1^2) = \frac{0}{6}(1) = 0 \\
c = \{3\}: & \quad \frac{(1-1)!(3-1)!}{3!}(\hat{V}(\{3\}) - \hat{V}(\{\emptyset\})) = \frac{(0!)(2!)}{3!}(1^2 - 1^2) = \frac{0}{6}(1) = 0 \\
c = \{1, 2\}: & \quad \frac{(2-1)!(3-2)!}{3!}(\hat{V}(\{1, 2\}) - \hat{V}(\{2\})) = \frac{(1!)(1!)}{3!}(2^2 - 1^2) = \frac{1}{6}(3) = \frac{3}{6} \\
c = \{1, 3\}: & \quad \frac{(2-1)!(3-2)!}{3!}(\hat{V}(\{1, 3\}) - \hat{V}(\{3\})) = \frac{(1!)(1!)}{3!}(2^2 - 1^2) = \frac{1}{6}(3) = \frac{3}{6} \\
c = \{2, 3\}: & \quad \frac{(2-1)!(3-2)!}{3!}(\hat{V}(\{2, 3\}) - \hat{V}(\{2, 3\})) = \frac{(1!)(1!)}{3!}(2^2 - 2^2) = \frac{0}{6}(3) = 0 \\
c = \{1, 2, 3\}: & \quad \frac{(3-1)!(3-1)!}{3!}(\hat{V}(\{1, 2, 3\}) - \hat{V}(\{2, 3\})) = \frac{(2!)(1!)}{3!}(3^2 - 2^2) = \frac{2}{6}(5) = \frac{10}{6} \\
& \quad F_1^{SV}(\{1, 2, 3\}, \hat{V}) = 3
\end{aligned}$$

In our example, $\hat{V}$ happens to be an anonymous game and so it is easy to see that:

$$F^{SV}(\{1, 2, 3\}, \hat{V}) = (3, 3, 3).$$

Notice that this allocation has two interesting features. First, it treats symmetric agents symmetrically. Second, the allocation is both feasible, and when summed over agents, adds up to value of the grand coalition. Thus, the allocation is Pareto optimal. In fact, these two, and two additional axioms, fully characterize the Shapley value. Formally:

Efficiency: $\sum_{i \in \mathcal{I}} F_i(\mathcal{I}, V) = V(\mathcal{I})$.

Symmetry: For any pair of agents $i, j \in \mathcal{I}$ if for all $c \subseteq \mathcal{I} / \{i, j\}$ it holds that $V(c \cup i) = V(c \cup j)$, then $F_i(\mathcal{I}, V) = F_j(\mathcal{I}, V)$.

Dummy: For any agent $i \in \mathcal{I}$ if for all $c \subseteq \mathcal{I} / \{i, j\}$ it holds that $V(c \cup i) - V(c) = 0$, then $F_i(\mathcal{I}, V) = 0$.

Additivity: Consider any two games V and W defined over $\mathcal{I}$. Define the composition of V and W as: $\forall c \in \mathcal{C}, (V+W)(c) = V(c) + W(c)$. Then, $F(\mathcal{I}, V+W) = F(\mathcal{I}, V) + F(\mathcal{I}, W)$.

Theorem (Shapley 1953): The Shapley Value is the unique solution satisfying efficiency, symmetry, dummy, and additivity.

Example of uses:

- Allocating costs in cloud computing or other resource splitting projects.
- Sharing the cost or benefits in a public venture or joint project.

Section 10.5. Cooperative Bargaining Theory

We mention above that experimental economics calls into question that agents behave in a way that could be called rational, at least by a narrow definition. Agents seem to dislike being treated unfairly and will sometimes act in apparently self-defeating ways to prevent it. **Cooperative bargaining theory** tries to address the question of what “fair” means by encoding it into list of axioms.

An **N-person bargaining problem** consists of a pair (S, d) where S is a nonempty subset of $\mathbb{R}^N$ and $d \in S$. The set S is interpreted as the set of utility allocations that are attainable through joint action on the part of all n agents and is called the **feasible set**. If the agents fail to reach an agreement, then the problem is settled at the point d , which is called the **disagreement point**. More formally:

Bargaining Problem: $(S, d) \in \Sigma$ where $d \in S \subset \mathbb{R}^N$ and Σ is a **domain bargaining problems** satisfying:

- (1) S is compact.
- (2) There exists $\exists x \in S$ and $x \gg d$.

Recall that a set $S \in \mathbb{R}^N$ is **compact** if it is closed and bounded. also recall that a set $S \in \mathbb{R}^N$ is **convex** if $\forall x, \bar{x} \in S \subseteq \mathbb{R}^N$, and $\forall \lambda \in [0, 1]$, it holds that $\lambda x + (1 - \lambda)\bar{x} \in S$.

Another property it seems reasonable to ascribe to feasible sets is called **free disposal**. This says that if a allocation is feasible, then all allocations that are smaller should also be feasible. For example, if it is feasible to give each of the two agents in an economy \$10, it should also be feasible to give each of them \$7, or go give the first agent \$10, and the second agent \$8. This process might be bounded below by zero since negative allocations would require taking money from agents that they might not have.

We express this formally by saying that a set S **comprehensive** with respect to a point d if it contains all points that are smaller than any given point $x \in S \subseteq \mathbb{R}^N$, but larger than d :

d-Comprehensive: $S \subset \mathbb{R}^N$ is d-comprehensive if $\forall x \in S, \{y \in \mathbb{R}^N \mid d \geq y \geq x\} \subseteq S$.

We will also be considering two subdomains of Σ that satisfy convexity or comprehensiveness:

Convex Bargaining Problem: $(S, d) \in \Sigma^{con} \subset \Sigma$ is an element of the domain of convex bargain problems if S is convex.

Comprehensive Bargaining Problem: $(S, d) \in \Sigma^{comp} \subset \Sigma$ is an element of the domain of comprehensive bargain problems if S is d-comprehensive.

A **bargaining solution** suggests a unique feasible allocation for any bargaining problem as a fair way to share the surplus over agents. Formally:

Bargaining Solution : $F: \tilde{\Sigma} \Rightarrow S$, such that $\forall (S, d) \in \tilde{\Sigma}$, $F(S, d) \in S$, where $\tilde{\Sigma}$ is a domain of bargaining problems, and F is single valued.

In the axiomatic approach to bargaining problems, we start by specifying a list of properties, called axioms, that we would like a solution to have. If it can be shown that there is one and only one solution that satisfies a given list of axioms, then the solution is said to be **characterized** by this list. We begin by defining three mapping that will be need to specify these axioms.

Permutation Operator: A permutation operator is a one-to-one mapping from a set of index numbers to itself: $\pi: \mathbb{N}^N \Rightarrow \mathbb{N}^N$ where $\mathbb{N}^N \equiv \{1, \dots, N\}$, and $\pi \in \Pi^N$ denotes the class of all such operators.

A permutation operation simply reorders index numbers. We will abuse notion and write:

$$\pi(x) \equiv (x_{\pi^{-1}(1)}, x_{\pi^{-1}(2)}, \dots, x_{\pi^{-1}(N)})$$

to denote the permutation of a vector. For example if $x = (5, 7)$ and $\pi(1) = 2$, $\pi(2) = 1$, then $\pi(x) = (7, 5)$. That is, the components of the vector x have been reordered (in this case, reversed) by the permutation operator. Similarly, we will write:

$$\pi(S) \equiv \{y \in \mathbb{R}^N \mid y = \pi(x) \text{ for some } x \in S\}$$

to denote the permutation of $S \in \mathbb{R}^N$ under some $\pi \in \Pi^N$.

Affine Transformation: $\psi: \mathbb{R}^N \Rightarrow \mathbb{R}^N$ such that $\psi(x) = a + bx$ for some $a \in \mathbb{R}^N$, $b \in \mathbb{R}_{++}^N$. Let Ψ^N denote the class of all such transformations.

As above, we will write:

$$\psi(S) \equiv \{y \in \mathbb{R}^N \mid y = \psi(x) \text{ for some } x \in S\}$$

to denote the affine transformation of $S \in \mathbb{R}^N$ under some $\psi \in \Psi^N$.

The **ideal point** of a bargaining problem is a kind of maximally aspirational, but not necessarily feasible, allocation. In words, it is the point composed of the largest payoff possible for each agent at any feasible allocation. Since the ideal point is composed by taking the largest value obtained by any element of S separately in each dimension, the ideal point is not feasible in general. Formally:

Ideal Point $a: \tilde{\Sigma} \Rightarrow \mathbb{R}^N$:

$$a(S, d) \equiv (\max_{\substack{x \in S \\ x \geq d}} x_1, \dots, \max_{\substack{x \in S \\ x \geq d}} x_N).$$

We now state several standard axioms which are meant to capture fair, or otherwise desirable, properties that bargaining solutions might satisfy.

Pareto Optimality (WPO): $F(S, d) \in WPO(S)$.

Pareto optimality is familiar and says that good solutions should not propose allocations that are wasteful in the sense that it remains feasible to make all agents better off. On the other hand, it does allow solutions to forego weak Pareto improvements, perhaps on the grounds that it is not fair to improve the welfare on one agent when you cannot improve the welfare of all.

Independence of Irrelevant Alternatives (IIA): If $\bar{S} \subseteq S$, $\bar{d} = d$, and $F(S, d) \in \bar{S}$, then $F(\bar{S}, \bar{d}) = F(S, d)$.

Suppose agents agreed that specific allocation was a fair way to solve a given bargaining problem. If some of the alternative allocations that the agent already rejected were suddenly to disappear from the feasible set, then IIA says that it is only fair that they stay with the solution they originally agreed to. After all, only things that have already been rejected have been removed from the feasible set.

Symmetry (SYM): If $\forall \pi \in \Pi^N$, $\pi(S) = S$ and $\pi(d) = d$, then $\forall i, j \in N$, $F^i(S, d) = F^j(S, d)$.

Suppose agents have identical positions in the game in terms of potential payoffs and the disagreement point. Then symmetry says that agents should also get identical payoffs.

Note that we use the permutation operator to determine symmetry as follows: Consider every possible way of relabeling the axis of the allocation space (for example switch the x-axis for the z-axis). If the shape of S does not change under any such permutation, then S is symmetric. Symmetry is sometimes called *anonymity* because it says that the names of the agents should not matter, only the shape of the feasible set.

Scale Invariance (S.INV): $\forall \lambda \in \Lambda^N$, $F(\lambda(S), \lambda(d)) = \lambda(F(S, d))$.

Utility is ordinal not cardinal. Thus, simply rescaling agents' payoffs should not alter the actual solution to a bargaining problem.

Translation Invariance (T.INV): $\forall x \in \mathbb{R}^N$, $F(S+x, d+x) = F(S, d)+x$ (where $S+x$ means set addition).

If the payoffs in the feasible set are cardinal, for example, if they are dollar, rather than utility, payoffs, that scale invariance would not be an appropriate or fair axiom to impose. Affine transformations would then have meaning. However, suppose that two feasible sets differ simply because of the addition or subtraction of a vector to both the feasible sets and the disagreement points. Since this would not change the *potential gains relative to the disagreement point*, translation invariance suggests that it is only fair that the solution should change by exactly the same vector.

Strong Monotonicity (S.MON): If $S \subset \bar{S}$ and $d = \bar{d}$, then $F(\bar{S}, \bar{d}) \geq F(S, d)$.

Suppose the feasible set gets strictly larger, but the disagreement point did not change. Strong monotonicity suggests that it is only fair that no agent should be harmed by this increase in possibilities.

Restricted Monotonicity (R.MON): If $S \subset \bar{S}$, $d = \bar{d}$, and $a(S, d) = a(\bar{S}, \bar{d})$, then $F(\bar{S}, \bar{d}) \geq F(S, d)$.

Suppose the feasible set gets strictly larger, but both the disagreement point and ideal point stay the same. Restricted monotonicity suggests that it is only fair that no agent should be harmed by the improved opportunities. However, it might be the case that most of the feasible points that are added improve the potential payoffs for some agents more than others. Thus, the aspirations of some agents are justifiable higher in the new situation. In this case, it might be fair that some agents benefit and some are harmed by in a fair allocation given the new feasible set.

Subsection 10.5.1. The Nash Solution

The **Nash bargaining solution** is defined as:

$$N(S, d) \equiv \underset{\substack{x \in S \\ x \geq d}}{\operatorname{argmax}} (x_1 - d_1)(x_2 - d_2) \dots (x_N - d_N).$$

The figure gives an illustration:

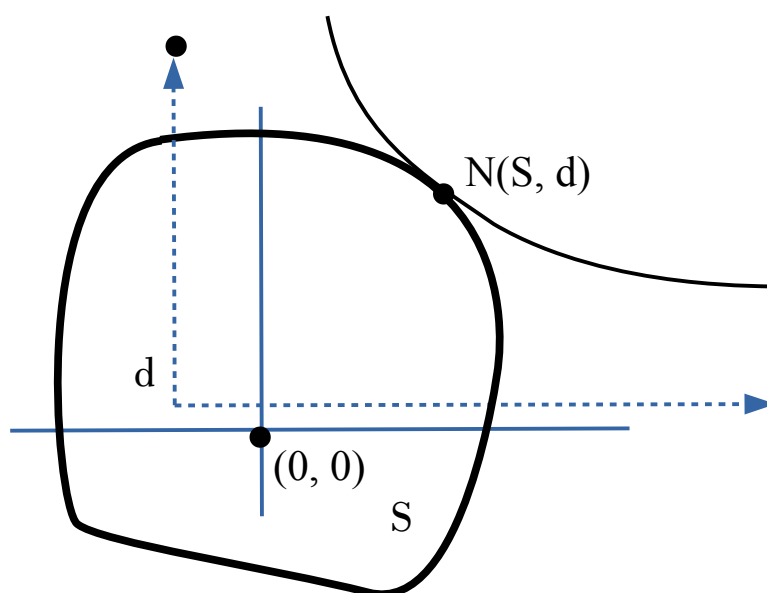


Figure 120: The Nash equilibrium bargaining solution

The notion here is that gains over the disagreement point show balance equity and efficiency. By taking the product of the relative gain as an objective function, the Nash solution chooses a point on the efficient boundary of the feasible set that compares the cost of giving additional payoff to one agent against the losses to the other.

The axioms that characterize the Nash solution are WPO, SYM, IIA and S.INV. The argument for this being a good list of axioms is the following:

- WPO: it is in no one's interest to waste resources.

- **SYM:** if agents have identical positions in the game in terms of potential payoffs and the disagreement point, they should get identical payoffs.
- **IIA:** The deletions of feasible allocations that have been rejected as solutions should not change the solution so a problem.
- **S.INV:** Since utility is not cardinal, simply rescaling agents' payoffs should not alter the actual outcome.

Theorem: A solution F defined on Σ^{con} satisfies WPO, SYM, IIA and S.INV if and only if it is the Nash bargaining solution.

Proof: First, we show that the solution satisfies the axioms. That is, the solution concept used will satisfy the four axioms for all problems in the domain IF it is the Nash solution.

WPO: Suppose there were a point $y \in S$ such that $y > N(S, d) \equiv x$. Then it must be the case that $(y^1 - d^1)(y^2 - d^2) \dots (y^N - d^N) \gg (x^1 - d^1)(x^2 - d^2) \dots (x^N - d^N)$. It is immediate that y could not be the Nash solution.

SYM: Suppose that (S, d) was symmetric but $x \equiv N(S, d)$ was not. Observe that since S is convex and $N(S, d)$ is the argmax of a strictly quasi-concave function, $N(S, d)$ is unique. But then $\forall \pi \in \Pi^N$:

$$(x^1 - d^1)(x^2 - d^2) \dots (x^N - d^N) = (\pi(x^1) - \pi(d^1))(\pi(x^2) - \pi(d^2)) \dots (\pi(x^N) - \pi(d^N)).$$

But $\exists \pi \in \Pi^N$ such that $N(S, d) \neq \pi(N(S, d))$, a contradiction of uniqueness.

IIA: Suppose $\bar{S} \subseteq S$, $\bar{d} = d$, and $N(S, d) \in \bar{S}$, but $N(\bar{S}, \bar{d}) \equiv y \neq N(S, d) \equiv x$. Then since $N(\bar{S}, \bar{d}) \equiv y$, $x \in N(\bar{S}, \bar{d})$ and $d = \bar{d}$, by definition:

$$(y^1 - d^1)(y^2 - d^2) \dots (y^N - d^N) \geq (x^1 - d^1)(x^2 - d^2) \dots (x^N - d^N).$$

However, $\bar{S} \subseteq S$, $\bar{d} = d$, and $y \in S$. Thus:

$$(x^1 - d^1)(x^2 - d^2) \dots (x^N - d^N) \geq (y^1 - d^1)(y^2 - d^2) \dots (y^N - d^N).$$

This contradicts the uniqueness of the Nash bargaining solution.

S.INV: Left to the reader.

Second, we should the other direction. That is, the solution concept used can satisfy the four axioms for all problems in the domain ONLY IF it is the Nash solution.

Consider any problem (S, d) . First rescale this problem with an affine transformation $\lambda \in \Lambda^N$ so that:

$$d = 0$$

$$(1, \dots, 1) = \lambda(N(S, d)) = N(\lambda(S), 0).$$

Since the transformed feasible set, $\lambda(S)$, is compact, there exists $B \in \mathbb{R}_{++}$ such that $\forall y \in \lambda(S)$:

$$(-B, \dots, -B) \geq (y_1, \dots, y_N) \geq (B, \dots, B).$$

Now, let Z be the symmetric, closed hypercube defined by:

$$Z \equiv \{y \in \mathbb{R}^N \mid \forall n \in \mathcal{N}, |y_n| \leq B\}.$$

This is shown in dark red in the figure. Now consider the symmetric problem, $(T, 0)$, where:

$$T = \{y \in Z \mid \sum_{n \in \mathcal{N}} y_n \leq N(\lambda(S), 0)\}.$$

In words, the set T is constructed by taking $H_{(1, \dots, 1), N(\lambda(S), 0)}$, the the "45 degree" hyperplane through point $(1, \dots, 1)$, and intersecting its lower half space with the symmetric hypercube Z . This is shown in with a thick, dashed blue in the figure. Then by SYM, and PO,

$$F(T, 0) = (1, \dots, 1)$$

since $(T, 0)$ is a symmetric bargaining problem, and $(1, \dots, 1)$ is the only symmetric and Pareto optimal element of T .

By construction, no element of T is above the symmetric hyperplane, $H_{(1, \dots, 1), N(\lambda(S), 0)}$. It is easy to see that Nash objective function is maximized at the symmetric point $(1, \dots, 1)$ on this hyperplane, and thus, over the entire set T . We conclude that:

$$N(T, 0) = (1, \dots, 1) = F(T, 0).$$

Also by construction, $\lambda(S) \subseteq T$ and $F(T, 0) = (1, \dots, 1) \in \lambda(S)$. Thus, by IIA,

$$N(\lambda(S), 0) = (1, \dots, 1).$$

Finally, by S.INV, we conclude that:

$$F(S, d) = N(S, d).$$

QED

Figure 121: Proving the Nash Equilibrium Characterization Theorem

Subsection 10.5.2. The Kalai-Smorodinsky Solution

The **Kalai-Smorodinsky bargaining solution** is defined as:

$$K(S, d) \equiv \max \{ x \in S \mid x \in \text{con}(a(S, d), d) \}.$$

The figure below illustrates the solution:

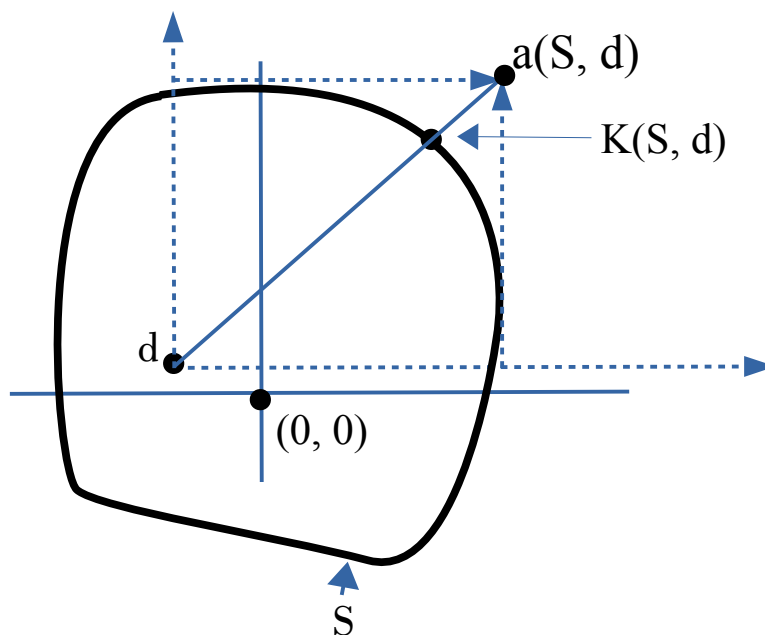


Figure 122: The Kalai-Smorodinsky bargaining solution

The notion here is that losses from the ideal or aspirational point should be shared proportionately in terms of the potential for each agent to gain over the disagreement point.

The axioms that characterize the Kalai-Smorodinsky solution are similar to those that characterize the Nash solution. The only difference is that R.MON replaces IIA. This says that if the feasible set gets larger but the both the disagreement point and ideal point stay the same, no agent should be harmed by the improved opportunities. However, it might be the case that if some feasible points were deleted (perhaps points that had high payoffs for some player i), and that this altered the ideal point, then even if the old outcome were still feasible, we might choose a new outcome (perhaps was less in favor of player i since his claim now seems weaker).

Theorem: A solution F on Σ^{comp} satisfies SYM, S.INV, WPO, and R.MON if and only if it is the Kalai-Smorodinsky solution.

Subsection 10.5.3. The Egalitarian Solution

The **Egalitarian bargaining solution** is defined as:

$$E(S, d) \equiv \max \{ x \in S \mid \forall i, j \in \mathcal{I}, x_i = x_j \}.$$

The figure below illustrates this solution:

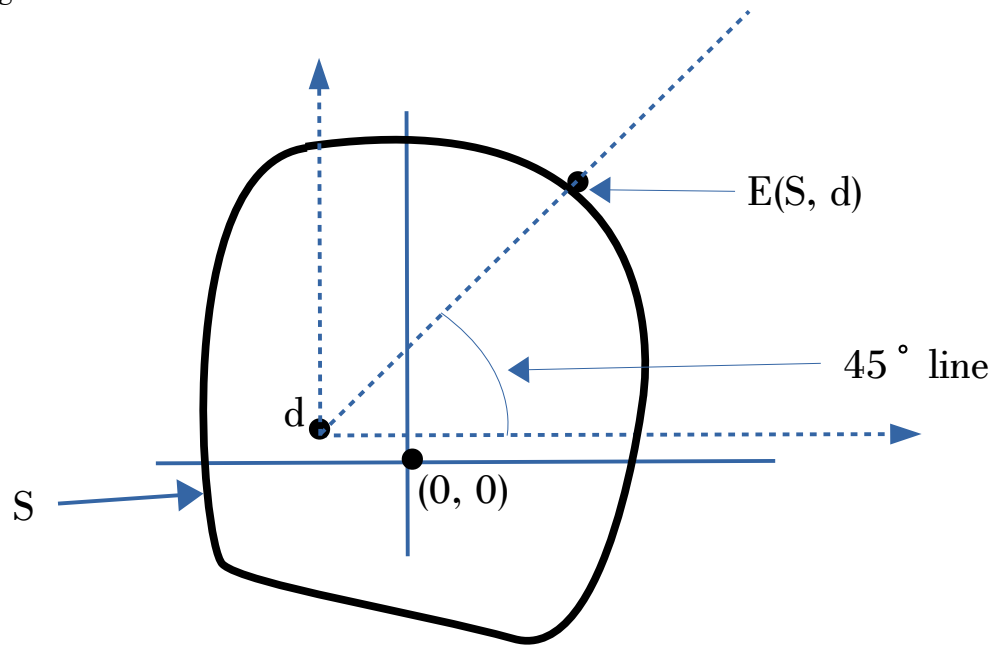


Figure 123: The egalitarian bargaining solution

The notion here is that agents should share any potential gains over the disagreement point equally.

The solution is characterized using the axioms for the Kalai-Smorodinsky, but weakening S.INV. to T.INV. This suggests that the payoffs have strong cardinal meaning, but that the payoffs should move in lockstep if we just added or subtracted a constant to or from all payoffs as well as disagreement point.

Theorem: A solution F on Σ^{comp} satisfies SYM, T.INV, WPO, and R.MON if and only if it is the Egalitarian solution.

Glossary

Affine Transformation: $\psi: \mathbb{R}^N \Rightarrow \mathbb{R}^N$ such that $\psi(x) = a + bx$ for some $a \in \mathbb{R}^N$, $b \in \mathbb{R}_{++}^N$. Let Ψ^N denote the class of all such transformations.

Axiom: An assumption or starting point from which logical implication may be derived. Axioms are foundational premises that are justified because they are self-evidently true, can be supported as factual by empirical evidence, or are in some way desirable points from which to proceed. In the case of cooperative game theory, axioms are meant to capture ethical notions of fair division, or fair procedure. Since these are not derived from primitives, they are assumptions without foundation that each individual may embrace or reject accounting to preference.

Bargaining Solution: A bargaining solution suggests a unique feasible allocation for any bargaining problem as a fair way to share the surplus over agents. Formally:

$$F: \tilde{\Sigma} \Rightarrow S, \text{ such that } \forall (S, d) \in \tilde{\Sigma}, F(S, d) \in S,$$

Blocking Coalition: Given a feasible allocation for the grand coalition, a set of agents for whom there exists an allocation that is feasible on their own, without the complementary coalition, and which makes all blocking agents better off than they were in the proposed allocation for the grand coalition.

Characteristic Function Form Game: A characteristic function form game is defined by a set of players and a characteristic function that maps coalitions of these players into payoffs.

Characteristic Function: A characteristic function is a mapping from coalitions of agents into payoffs. If the feasible set of payoffs is freely transferable over agents the mapping is from coalitions into $\mathbb{R}$, while if payoffs are not transferable, the mapping is from coalitions into $\mathbb{R}^{|C|}$. In effect, a characteristic function is a measure of the value or productivity of various coalitions of agents.

Characterization: If it can be shown that there is one and only one solution that satisfies a given list of axioms, then the solution is said to be **characterized** by this list.

Coalition: A set of agents drawn from some specified population of agents (called the grand coalition).

Comprehensive Hull: The comprehensive hull of a set $S \subset \mathbb{R}^N$ with respect to a point $d \in \mathbb{R}^N$ is the smallest d -comprehensive set containing S :

$$\text{comp}(S; d) \equiv \{x \in \mathbb{R}^N \mid x \in S \text{ or } \exists y \in S \text{ such that } d \geq x \geq y\}.$$

Convex Hull: The convex hull of a set $S \subset \mathbb{R}^N$ is the smallest convex set containing the set S :

$$\text{con}(S) \equiv \{z \in \mathbb{R}^N \mid \exists x_1, \dots, x_N \in S, \text{ and } \lambda \in \Delta^{N-1} \text{ such that } z = \sum_{n \in N} \lambda_n x_n\}.$$

Cooperative Bargaining Theory: A branch of cooperative game theory that focuses on finding single-valued solution concepts that can be characterized by attractive lists of axioms to specific classes bargaining problems.

Cooperative Game: A game in which agents can sign binding agreements to share and divide pay-offs.

Core Allocation: A feasible allocation for the grand coalition that is stable against all possible coalitional deviations in the sense that no coalition could do better for all of its members using only its own resources if it defected.

Egalitarian Bargaining Solution: A bargaining solution that gives all agents the maximum feasible payoff that gives all agents the same amount:

$$E(S, d) \equiv \max \{ x \in S \mid \forall i, j \in \mathcal{I}, x_i = x_j \}.$$

Grand Coalition: The coalition consisting of all agents in the population, $\{1, \dots, I\} \equiv \mathcal{I}$.

Ideal Point: The **ideal point** of a bargaining problem is a kind of maximally aspirational, but not necessarily feasible, allocation. In words, it is the point composed of the largest payoff possible for each agent at any feasible allocation. Since the ideal point is composed by taking the largest value obtained by any element of S separately in each dimension, the ideal point is not feasible in general. Formally:

$$a(S, d) \equiv (\max_{\substack{x \in S \\ x \geq d}} x_1, \dots, \max_{\substack{x \in S \\ x \geq d}} x_n).$$

Individually Rational Set: The set of all feasible allocations that gives all agents in the grand coalition at least as much as they could get on their own. Formally:

$$IR(\mathcal{I}) \equiv \{ \hat{x} \in FS(\mathcal{I}) \mid \nexists i \in \mathcal{I} \text{ and } \bar{x}_i \in FS(i) \text{ such that } \bar{x}_i \succ_i \hat{x}_i \}.$$

Kalai-Smorodinsky Bargaining Solution: A bargaining solution that shares losses from the ideal or aspirational point relative to the disagreement point proportionally over agents. Formally:

$$K(S, d) \equiv \max \{ x \in S \mid x \in \text{con}(a(S, d), d) \}.$$

Nash Bargaining Solution: A bargaining solution that balances equity and efficiency and chooses a point on the efficient boundary of the feasible set that compares the cost of giving additional payoff to one agent against the losses to the other. Formally:

$$N(S, d) \equiv \underset{\substack{x \in S \\ x \geq d}}{\text{argmax}} (x_1 - d_1)(x_2 - d_2) \dots (x_N - d_N).$$

N-Person Bargaining Problem: This is defined by an ordered pair (S, d) where S is a nonempty subset of $\mathbb{R}^N$ and $d \in S$. The set S is interpreted as the set of utility allocations that are attainable through joint action on the part of all N agents and is called the **feasible set**. If

the agents fail to reach an agreement, then the problem is settled at the point d , which is called the **disagreement point**.

Permutation Operator: A permutation operator is a one-to-one mapping from a set of index numbers to itself: $\pi: \mathbb{N}^N \Rightarrow \mathbb{N}^N$ where $\mathbb{N}^N \equiv \{1, \dots, N\}$, and $\pi \in \Pi^N$ denotes the class of all such operators.

Shapley Value (or Value): Given a set of agents, $\mathcal{I}$, and a characteristic function of a game $V: \mathcal{C} \Rightarrow \mathbb{R}$, the Shapley value proposes the following payoff for each agent $i \in \mathcal{I}$:

$$F_i^{SV}(\mathcal{I}, V) = \sum_{c \in \mathcal{C}} \frac{(|c|-1)!(I-|c|)!}{I!} (V(c) - V(c/i)).$$

Solution Concept: is a mapping from a space of economies or games into the feasible set. Such mappings might be single valued or set valued.

Strong Pareto Set: The set of feasible allocations for which no weak Pareto improvement exists:

$$SPO(\mathcal{I}) \equiv \{ \bar{x} \in FS(\mathcal{I}) \mid \nexists \hat{x} \in FS(\mathcal{I}) \text{ such that } \forall i \in \mathcal{I}, \hat{x}_i \geq_i \bar{x}_i \text{ and } \exists j \in \mathcal{I} \hat{x}_j >_j \bar{x}_j \}.$$

Super-additive Characteristic Function Form Game: A game $(V, \mathcal{I})$ is **super-additive** if $\forall \hat{c}, \bar{c} \subset \mathcal{C}$ where $\hat{c} \cap \bar{c} = \emptyset$, $V(\hat{c} \cup \bar{c}) \geq V(\hat{c}) + V(\bar{c})$.

Weak Pareto Set: The set of feasible allocations for which no strong Pareto improvement exists:

$$WPO(\mathcal{I}) \equiv \{ \bar{x} \in FS(\mathcal{I}) \mid \nexists \hat{x} \in FS(\mathcal{I}) \text{ such that } \forall i \in \mathcal{I}, \hat{x}_i >_i \bar{x}_i \}$$

Problems

1. Consider the following:

- a. In words define the notion of “blocking” and use this to give a definition of the “core”. Now consider the following problem:
- b. Britney and KFed get married. At first, they are deeply in love. Britney gets 20 units of utility if she is with KFed while she gets only 10 if she is with Justin. KFed gets 20 when he is with Britney and only 2 if they break up (he has fewer outside options). They also have a dog (Muffin) who needs to be walked every day. Muffin is particular and insists that the same person do it every day. Both Britney and KFed hate this and whoever takes on this job loses 8 units of utility. In a core allocation, will Britney or KFed walk the dog, or can you tell?
- c. Now suppose that things are getting rocky. Britney now only gets 12 units of utility from being with KFed and still gets 10 if she chooses Justin. KFed also gets 12 from being with Britney, and 2 from being with that cute waitress at IHOP (his best outside option). If they break up, they give away Muffin and no one has to walk him. In a core allocation will Britney or KFed walk the dog, is it splitsville (with Britney and KFed taking their best outside options), or can you tell?

2. Consider the following game in characteristic function form with three agents:

$$3. \{ \text{Fred, Joe, Sue} \} \equiv I.$$

The payoffs the various coalitions can get are the following:

$$4. F(\text{Fred}) = F(\text{Joe}) = F(\text{Sue}) = 20$$

$$5. F(\text{Fred}) = F(\text{Joe}) = 30, F(\text{Sue}) = 35$$

$$6. F(\text{Fred, Joe, Sue}) = 60$$

- a. Write down a core allocation. Can you give more than one?
- b. Write down an allocation that is not in the core. Write down the coalition and allocation that blocks it.
- c. Is this game super-additive?

Appendix A. Conventions

Section A.1. A List of Mathematical Symbols and Notation

A summary list of these mathematical symbols used in this book, with English interpretations.

$\forall$: For all, For every

$\exists$: There exists, For some

$\nexists$: There does not exist, For no

$\in$: Element of, In

$\notin$: Not an element of, Not in

$\subset$: Strict Subset of, Contained in but not equal to

$\subseteq$: Subset of, Contained in or equal to

$\cup$: Union of sets

$\cap$: Intersection of sets

$\wedge$: Logical “and”

$\vee$: Logical “or”

$\Rightarrow$: Implies, If

$\Leftarrow$: Implied by, Only if

$\Leftrightarrow$: If and only if

$\emptyset$: Empty set

$\mathbb{R}$: Real numbers

$\mathbb{N}$: Natural numbers

Section A.2. Index and Vector Notation

Indices keep track of attributes of vectors or other objects. For example, which good is being consumed or produced, by which agent or firm. We will use the following conventions:

Index Set: An ordered set natural numbers, $\mathbb{N}$, which may be finite or infinite. If an index set is finite, it runs from 1 to some upper bound. We will use a lower-case letter to denote an **element of an index set**, an upper-case letter to denote the **upper bound of an index set**, and a script letter to denote the **entire index set**.

$$(1, \dots, n, \dots, N) \equiv \mathcal{N}.$$

N-Dimensional Vector: An ordered set of N real numbers. We will denote vectors with lower-case letters, the n^{th} **component of a vector** with a subscripted index, and the **set that an element is a member of** with an upper-case letter:

$$(x_1, \dots, x_n, \dots, x_N) \equiv x \in X \subseteq \mathbb{R}^N.$$

When we consider multiple agents producing or consuming multiple goods, we need two index sets to distinguish the elements of a vector along these two separate sets of attributes. Thus, agent i 's consumption of good n is the scalar:

$$x_{i,n} \in \mathbb{R}.$$

Here, the first subscript indicates the agent while the second indicates the good.

Dropping a subscript will indicate a vector composed whole list over the entire index set of the attribute. Thus, we indicate the consumption vector of agent i , and the vector of consumption amounts of good n for each agent, respectively:

$$x_i = (x_{i,1}, \dots, x_{i,N}) \in \mathbb{R}^N, \text{ and } x_n = (x_{1,n}, \dots, x_{I,n}) \in \mathbb{R}^I,$$

You may notice that this creates a potential ambiguity since x_2 could be either a consumption vector for agent 2, or the consumption level of good 2 for all agents. In practice, we will refer to generic agents or goods, as in x_i and x_n , and so the choice of index will provide context. We will not use more complex notational forms that would prevent such ambiguities in the interest of readability.

To build up the space from which an economy-wide consumption or productions plan is drawn, we need the following:

N-fold Cartesian Product: The set of ordered N-tuples where the n^{th} component is an element of the n^{th} set in a concatenated list:

$$(x_1, \dots, x_I) \equiv x \in X \equiv X_1 \times \dots \times X_I \equiv \prod_{i \in \mathcal{I}} X_i.$$

As an example the consumption plan is an allocation of each good for each agent and is denoted x . Since the consumption vector for each agent must be in that agent's consumption set, the plan as a

whole must be in Cartesian product of these consumption sets over the index set for agents: $X_1 \times \dots \times X_I \equiv X$. while the combined consumption vector over all agents and goods is written:

$$(x_{1,1}, \dots, x_{i,n}, \dots, x_{I,N}) \equiv (x_1, \dots, x_I) \equiv x \in X_1 \times \dots \times X_I \equiv \prod_{i \in \mathcal{I}} X_i \equiv X \in \mathbb{R}^{I \times N}.$$

Note that we are making use of another piece of notation: $\prod$. This is very much like the summation operator, $\sum$, you are already familiar with, but instead of summing over an index set, $\prod$ is the product operator and so tells you to multiply the elements in the index set together. Thus, we indicate a production plan for the economy as:

$$(y_1, \dots, y_F) \equiv y \in Y \equiv Y_1 \times \dots \times Y_F \equiv \prod_{f \in \mathcal{F}} Y_f.$$

Finally, we note that it is conventional to distinguish ordered and unordered sets with round and curly brackets, respectively.

Ordered Set: A set for which the order of elements is part of its definition, and are which denoted by round brackets:

$$(1, \dots, N) \subset \mathbb{N}^N, (x_1, \dots, x_N) \in \mathbb{R}^N, (x_1, \dots, x_I) \in \mathbb{R}^{I \times N}, \text{ and } (x^s)_{s \in \mathbb{N}}$$

Examples of ordered sets are index sets, the set of scalar components of vectors ordered by their competent index, sets of vectors ordered by agent or other index, and infinite sequences ordered in a defined, but otherwise arbitrary, way.

Unordered Sets: A set containing a collection of elements which is defined only by the elements in the collection, not by the order in which they are listed, and which are denoted by curly brackets:

$$\{a, b, c\} = \{b, c, a\} \in S, \text{ Spref}(x) \equiv \{z \in X \mid z \succ x\}.$$

Examples of unordered sets are arbitrary collections elements, which may or may not be vectors, and abstractly defined sets. In the first case, sets that contain the same elements are equivalent, regardless of the order of the elements In the latter case, there may be no meaningful way, or need, to order the elements.

Section A.3. Index Sets

The following is a list of the most common index sets used in the text

$i, I, \mathcal{I}$	agents (individuals)
$f, F, \mathcal{F}$	firms
$n, N, \mathcal{N}$	commodities

$t, T, \mathcal{T}$ time, period

$c, C, \mathcal{C}$ coalitions

Subsection A.3.1. Sets, Vectors, and Functions

The following is a list of the most commonly used letters to express certain sets, vectors, and functions.

x, y, z generic vectors or elements in a set

S, T generic subsets $\mathbb{R}^N$

k a scalar constant $k \in \mathbb{R}$

f, g, h generic functions or correspondences in a Euclidean space

Subsection A.3.2. Markets

Subsubsection A.3.2.1. Partial Equilibrium Analysis

The following is a list of the most commonly used letters to express certain economic variables in a partial equilibrium environment.

P price

Q quantity

W wealth or income

Subsubsection A.3.2.2. General Equilibrium Economies

The following is a list of the most commonly used letters to express certain economic variables in a general equilibrium environment.

p commodity prices

x commodity quantities for consumers

w wealth or income $w \in \mathbb{R}_+$

ω commodity endowment, $\omega \in \mathbb{R}_+^N$

X	consumption set
u	utility function
y	net production vector for firms
Y	production set
f	production function

Subsection A.3.3. Games

Subsubsection A.3.3.1. Noncooperative Games

The following is a list of the most commonly used letters to express certain elements of noncooperative games

$s \in S$	strategies
$a \in \mathcal{A}$	actions in a sequential game
$d \in \mathcal{D}$	decision nodes in a sequential game
$t \in \mathcal{T}$	terminal nodes in a sequential game
F	payoff functions: a mapping from strategies into a payoff space over agents

Subsubsection A.3.3.2. Cooperative Games and Coalitions

The following is a list of the most commonly used letters to express certain elements of cooperative games.

$i \in c \in \mathcal{C}$	agent i is in coalition c , an element of the set of all possible coalitions $\mathcal{C}$
V	value functions, characteristic functions
F	solution concept
S	set of feasible allocations in a cooperative bargaining game
$d \in S$	disagreement point in a cooperative bargaining game

Appendix B. Mathematical Review

Section B.1. Vectors, Sets, and Order

Subsection B.1.1. Euclidean Spaces and Vector Order

Most of what we will do in this book takes place in $\mathbb{R}^N$:

Real Numbers: $\mathbb{R}$ denotes the set of real numbers. You can think of this as the **real line**.

$\mathbb{R}^N \equiv \mathbb{R} \times \dots \times \mathbb{R}$ is the **N-fold Cartesian product** of the set of real numbers.

We distinguish three types of vector inequality or vector ordering.

Consider two vectors, $x, y \in \mathbb{R}^N$:

Strictly Greater than: $x \gg y$ if $\forall n \in \mathcal{N}, x_n > y_n$

for example: $(2,3,7) \gg (1,0,6)$.

Greater than: $x > y$ if $\forall n \in \mathcal{N}, x_n \geq y_n$ and $\exists m \in \mathcal{N}$ such that $x_m > y_m$

for example: $(2,3,7) > (1,0,6)$, and $(2,3,7) > (2,3,6)$.

Greater than or Equal to: $x \geq y$ if $\forall n \in \mathcal{N}, x_n \geq y_n$

for example: $(2,3,7) \geq (1,0,6)$, $(2,3,7) \geq (2,3,6)$, and $(2,3,7) \geq (2,3,7)$.

Of course, it may be that none of these apply. For example: $(2,3,7)$ and $(1,8,6)$ are not vector ordered by any of the above inequalities.

More formally, $\mathbb{R}^N$ this is an **N-dimensional real vector space** endowed with the Euclidean distance metric, a linear structure, and an inner product operation.

The **Euclidean metric** is the distance measure for Euclidean spaces and is defined as:

$$d(x, y) \equiv \|x - y\| \equiv \sqrt{\sum_{n \in \mathcal{N}} (x_n - y_n)^2}.$$

This is sometimes called the **Euclidean norm**. More generally, given an arbitrary space, S , a **metric** is a measure of distance, $d: S \times S \Rightarrow \mathbb{R}$, that satisfies four conditions $\forall x, y, z \in S$:

Identity Condition: $d(x, y) = 0 \Leftrightarrow x = y$

Symmetry: $d(x, y) = d(y, x)$

Nonnegativity: $d(x, y) \geq 0 \Leftarrow x \neq y$

Triangle Inequality: $d(x, y) + d(y, z) \geq d(x, z)$

The Euclidean metric measures the distance between two points in a real space, and you can easily see that walking from point A to point B on the straight line connecting them is the quickest path. If you take a detour to include a third point C, the trip will take longer unless C happens to be on the line between A and B, in which case, it will take the same amount of time (this is the triangle inequality).

A space becomes a **linear space** (equivalently, a **vector space**) if it satisfies certain familiar axioms you also learned in middle school. Let $\alpha, \beta \in \mathbb{R}$ be real numbers which are called scalars. Let $x, y, z \in \mathbb{R}^N$ be vectors. First, define three important operations using vectors and scalars:

Vector Addition: $x + y = (x_1 + y_1, x_2 + y_2, \dots, x_N + y_N)$

Scalar Multiplication: $\alpha x = (\alpha x_1, \alpha x_2, \dots, \alpha x_N)$

Inner Product or Dot Product: $x \cdot y \equiv \sum_{n \in \mathcal{N}} x_n y_n$

The **inner product** or **dot product operation** for vectors comes up a lot in economic contexts. Let $p \in \mathbb{R}^N$ be a price vector, that is, vector that gives the price of each of the i 's goods in the economy. Similarly, let $x \in \mathbb{R}^N$ be a consumption vector. Then the dot product of these two vectors is defined as follows:

$$p \cdot x = p_1 x_1 + p_2 x_2 + \dots + p_N x_N \equiv \sum_{n \in \mathcal{N}} p_n x_n$$

which is simply the total cost of bundle x under prices, $>, \geq, \sim$,

If a space is **linear**, the following axioms hold for these operations:

Commutativity of Vector Addition: $x + (y + z) = (x + y) + z$

Associativity of Vector Addition: $x + y = y + x$

Associativity of Scalar Multiplication: $\alpha(\beta x) = (\alpha\beta)x$

Distributivity of Scalar Addition: $(\alpha + \beta)x = \alpha x + \beta x$

Distributivity of Vector Sums: $\alpha(x + y) = \alpha x + \alpha y$

Scalar Multiplication Identity: $1(x) = x$

Existence of an Additive Inverse: $x + (-x) = 0 \equiv (0, \dots, 0)$

Additive Identity (existence of a zero): $x + 0 = 0 + x = x$

Subsection B.1.2. Sequences

Sequence: A sequence in $\mathbb{R}^N$ is an infinite, ordered, set of points, $(x^s)_{s \in \mathbb{N}}$, such that $\forall s \in \mathbb{N}, x^s \in \mathbb{R}^N$.

Convergence: A sequence in $\mathbb{R}^N$, $(x^s)_{s \in \mathbb{N}}$, **converges** to $\bar{x} \in \mathbb{R}^N$ if $\forall \varepsilon > 0, \exists \hat{s} \in \mathbb{N}$ such that $\forall t \in \mathbb{N}$ where $t > \hat{s}$, $\|x^t - \bar{x}\| < \varepsilon$.

Limit: $\bar{x}$ is the **limit** of a sequence $(x^s)_{s \in \mathbb{N}}$: if the sequence converges to $\bar{x}$, denoted:

$$\lim_{s \rightarrow \infty} x^s = \bar{x} \text{ or } x^s \rightarrow \bar{x}.$$

Subsection B.1.3. Open, Closed and Bounded Sets

Epsilon Ball: An open epsilon ball around a point, $x \in \mathbb{R}^N$

$$B_\varepsilon(x) \equiv \{z \in \mathbb{R}^N \mid \|x - z\| < \varepsilon\}.$$

We will also call this an **open neighborhood of x** or simply a **neighborhood of x**.

Interior: $x \in S \subseteq \mathbb{R}^N$ is an **interior point** of a set S if $\exists \varepsilon > 0$ such that $B_\varepsilon(x) \subseteq S$.

$$interior(S) \equiv \{x \in S \mid \exists \varepsilon > 0 \text{ and } B_\varepsilon(x) \subseteq S\}.$$

Boundary: $x \in S \subseteq \mathbb{R}^N$ is a **boundary point** of S if $\forall \varepsilon > 0, \exists y \in B_\varepsilon(x)$ such that $y \notin S$.

$$boundary(S) \equiv \{x \in S \mid \forall \varepsilon > 0, \exists y \in B_\varepsilon(x) \text{ and } y \notin S\}.$$

Open Set: $S \subseteq \mathbb{R}^N$ is an **open set** if $S = interior(S)$.

Closed Set: $S \subseteq \mathbb{R}^N$ is a **closed set** if $boundary(S) \subseteq S$.

Complement of a Set: The **complement** of a set S relative to $\mathbb{R}^N$ (or any $T \subseteq \mathbb{R}^N$) consists of all elements of $\mathbb{R}^N$ (or T) that are not also in S :

$$complement(S) \equiv \{x \in \mathbb{R}^N \mid x \notin S\}.$$

Bounded: A set $S \subseteq \mathbb{R}^N$ is **bounded** if $\exists B \in \mathbb{R}$ such that $\forall x \in S, \|x\| < B$.

Bounded Above: A set $S \in \mathbb{R}^N$ is **bounded from above** if $\exists B \in \mathbb{R}^N$ such that $\forall x \in S$, $x \leq B$.

Bounded Below: A set $S \in \mathbb{R}^N$ is **bounded from below** if $\exists B \in \mathbb{R}^N$ such that $\forall x \in S$, $x \geq B$.

Compactness: A set $S \in \mathbb{R}^N$ is **compact** if it is both *closed* and *bounded*.

Set Subtraction: The **set subtraction** of a $T \in \mathbb{R}^N$ from a $S \in \mathbb{R}^N$ removes any elements in T from S , and is denoted: “/”.

$$S/T \equiv \{x \in S \mid x \notin T\} \equiv \{x \in S \mid x \notin S \cap T\}.$$

If not elements in T are also in S then $S/T = S$.

Note the following:

- The union of an arbitrary number of open sets is also an open set.
- The intersection of a finite number of closed sets is also a closed set.
- The set of all open sets in $\mathbb{R}^N$ is defined as follows:

$$\mathcal{O} \equiv \{S \subseteq \mathbb{R}^N \mid \forall x \in S, \exists \varepsilon > 0 \text{ and } B_\varepsilon(x) \subseteq S\}.$$

- The interior of a set S is equal the union of all open sets contained in S :

$$\text{interior}(S) \equiv \bigcup_{T \in \mathcal{O}} T,$$

or, equivalently:

$$\text{interior}(S) \equiv \{x \in S \mid \exists T \subseteq S, \text{ such that } x \in T \text{ and } T \text{ is an open set}\}.$$

- The boundary of a set S is equal to the intersection of its closure and the closure of its complement:

$$\text{boundary}(S) \equiv \text{closure}(S) \cap \text{closure}(\text{complement}(S)).$$

- The closure of a set S is equal to the intersection of all closed sets containing S .
- A closed set is the complement of an open set.
- Both open and closed sets can be unbounded. For example, the positive orthant without the axes is an open set, while if we add the axes and the origin, it becomes a closed set.
- The empty set ($\emptyset$) and the entire space ($\mathbb{R}^N$) are technically both open and closed.

Subsequence: Given a sequence, $(x^s)_{s \in \mathbb{N}}$, a **subsequence** $(y^t)_{t \in \mathbb{N}} \subseteq (x^s)_{s \in \mathbb{N}}$ is a selection of an infinite number of elements from the original sequence.

For example, a subsequence might consist of every other element of the original sequence. Thus, a subsequence is a **countably infinite** subset of the original sequence, which also countably infinite.

Not every sequence converges. For example, it is possible for a sequence in an unbounded set such as $\mathbb{R}^N$ to go to infinity. Sequences in bounded sets also may not converge. For example, a sequence could cycle between two points forever. We have an important result, however, that tells us that if a sequence drawn from a bounded set but does not converge, we can always find a subsequence that will converge. For example, in the case of an oscillating sequence, we could just choose one of the points forever as our subsequence, which is trivially convergent.

Bolzano–Weierstrass Theorem: Let $(x^s)_{s \in \mathbb{N}}$ be a bounded sequence in such that $\mathbb{R}^N$. Then there exists a subsequence $(y^t)_{t \in \mathbb{N}} \subseteq (x^s)_{s \in \mathbb{N}}$ that converges to some $\bar{y} \in \mathbb{R}^N$.

Subsection B.1.4. Convex and Comprehensive Sets

Convex Set: $S \subseteq \mathbb{R}^N$ is **convex** if $\forall x, \bar{x} \in S \subseteq \mathbb{R}^N$, and $\forall \lambda \in [0,1]$, $\lambda x + (1-\lambda)\bar{x} \in S$.

That is, a set is (weakly) convex if the weighted average between any two points in the set is also in the set. Graphically, this defines is a line segment between x and y including the endpoints. In other words, a set is convex if the line connecting any two points in the set remains entirely within the set.

Extreme Point: A vector $x \in \mathbb{R}^N$ is an **extreme point** of a convex set $S \subset \mathbb{R}^N$ if it cannot be expressed as $x = \lambda y + (1-\lambda)z$ for any $y, z \in S$ and $\lambda \in (0,1)$.

Carathéodory's Theorem: Let $S \subset \mathbb{R}^N$ be a convex and compact set. Then every $x \in S$ can be expressed as a convex combination of at most $N+1$ extreme points.

(N-1)-Dimensional Unit Simplex: The **simplex** in $\mathbb{R}^N$ is defined as follows:

$$\Delta^{N-1} \equiv \{p \in \mathbb{R}_+^N \mid \sum_{n \in \mathcal{N}} p_n = 1\}.$$

Convex Hull:

$$\text{con}(S) \equiv \{z \in \mathbb{R}^N \mid \exists x_1, \dots, x_N \in S, \text{ and } \lambda \in \Delta^{N-1} \text{ such that } z = \sum_{n \in \mathcal{N}} \lambda_n x_n\}.$$

It turns out that the convex hull of a set is the smallest convex set that contains original set.

Comprehensiveness: A set $S \in \mathbb{R}^N$ is **comprehensive** if $\forall x \in S$, $y \leq x \Rightarrow y \in S$.

In words, a set is comprehensive if it includes all the points that we weakly smaller than every point in the set. Put another way, if a point is in a comprehensive set, then the negative orthant that extends below that point is also in the set. This allows us to define the following related idea:

Comprehensive Hull:

$$\text{comp}(S) \equiv \{x \in \mathbb{R}^N \mid \exists y \in S \text{ such that } y \leq x\}.$$

Section B.2. Hyperplanes and Separation Theorems

Hyperplane: Let $p \in \mathbb{R}^N$ with $p \neq 0$ and $k \in \mathbb{R}$, then the **hyperplane generated by p and k** is defined as:

$$H_{p,k} \equiv \{z \in \mathbb{R}^N \mid pz = k\}.$$

Upper Half Space: The upper half space of $H_{p,k}$ is defined as: $\{z \in \mathbb{R}^N \mid pz \geq k\}$.

Upper Lower Space: The lower half space of $H_{p,k}$ is defined as: $\{z \in \mathbb{R}^N \mid pz \leq k\}$.

Separation: Consider two sets $S, T \subset \mathbb{R}^N$. A hyperplane $H_{p,k}$ is said to **separate S from T** if:

$$\forall x \in S, \text{ and } y \in T, \quad px \leq k \leq py.$$

Separating Hyperplane Theorem: Let $S \subset \mathbb{R}^N$ be convex and closed, and $x \notin S$ be a vector in $\mathbb{R}^N$. Then $\exists p \in \mathbb{R}^N$ with $p \neq 0$ and $k \in \mathbb{R}$ such that $px > k$ and $\forall y \in S, py < k$.

More generally:

Minkowski Separation Theorem: Let $S, T \subset \mathbb{R}^N$ be disjoint convex sets, $S \cap T = \emptyset$. Then $\exists p \in \mathbb{R}^N$ with $p \neq 0$ and $k \in \mathbb{R}$ such that $H_{p,k}$ separates S from T .

Supporting Hyperplane Theorem: Let $S \subset \mathbb{R}^N$ be a convex set, and $x \in \text{boundary}(S)$. Then $\exists p \in \mathbb{R}^N$ with $p \neq 0$ and $k \in \mathbb{R}$ such that $\forall y \in S, k = px \geq py$.

Section B.3. Rules for Calculus

Subsection B.3.1. Derivatives

Consider two functions of one variable, $f: \mathbb{R} \Rightarrow \mathbb{R}$ and $g: \mathbb{R} \Rightarrow \mathbb{R}$, and let $k \in \mathbb{R}$ be a constant:

Constant Rule: $\frac{\partial k}{\partial x} = 0$

Power Rule: $\frac{\partial x^s}{\partial x} = sx^{s-1}$

Derivatives of Exponents: $\frac{\partial k^x}{\partial x} = k^x \ln(k)$

Derivatives of Natural logs: $\frac{\partial \ln|x|}{\partial x} = \frac{1}{x}$

Derivatives of Exponential Functions: $\frac{\partial e^x}{\partial x} = e^x$

Product Rule: $\frac{\partial f(x)g(x)}{\partial x} = \frac{\partial f(x)}{\partial x}g(x) + \frac{\partial g(x)}{\partial x}f(x)$

Quotient Rule: $\frac{\partial \frac{f(x)}{g(x)}}{\partial x} = \frac{\frac{\partial f(x)}{\partial x}g(x) - \frac{\partial g(x)}{\partial x}f(x)}{g(x)^2}$

Chain Rule: $\frac{\partial f(g(x))}{\partial x} \equiv \frac{\partial f \circ g(x)}{\partial x} = \frac{\partial f(g)}{\partial g} \frac{\partial g(x)}{\partial x}$

Subsection B.3.2. Algebra of Exponents

Multiplication Rule: $x^s x^t = x^{s+t}$

Division Rule: $\frac{x^s}{x^t} = x^{s-t}$

Power Rule: $(x^s)^t = x^{s \times t}$

Inverse Rules:

$$\frac{1}{x^t} = x^{-t} \text{ and } \frac{1}{x^{-t}} = x^t$$

Subsection B.3.3. Algebra of Natural Logarithms

Euler's Number:

$$e \approx 2.71828183$$

Natural Logarithm:

$$\ln(x) \equiv \log_e(x)$$

Product Rule:

$$\ln(x \cdot y) = \ln(x) + \ln(y)$$

Quotient Rule:

$$\ln(x/y) = \ln(x) - \ln(y)$$

Power Rule:

$$\ln(x^y) = y \ln(x)$$

Derivative:

$$\frac{\partial \ln(x)}{\partial x} = \frac{1}{x}$$

Integral:

$$\int \ln(x) dx = x \cdot (\ln(x) - 1) + C$$

Euler's Identity:

$$\ln(-1) = i\pi$$

Inverses:

$$e^{\ln(x)} = x \text{ and } \ln(e^x) = x \text{ for } x > 0$$

Negative number:

$$\ln(x) = \text{undefined for } x \leq 0$$

Zero:

$$\ln(0) = \text{undefined}$$

One:

$$\ln(1) = 0$$

Infinity:

$$\lim_{x \rightarrow \infty} \ln(x) = \infty$$

Section B.4. Functions

Subsection B.4.1. Continuous Functions

A **function** $f: S \Rightarrow T$ is a single-valued mapping that associates one and only one element of the **domain** T with each element of the **range** S . Thus, $\forall x \in S, \exists y \in T$, such that $y \in f(x)$ and $\forall z \in f(x), z = y$.

A function $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ is **continuous at a point** $x \in \mathbb{R}^N$ if for any sequence such that $\lim_{s \rightarrow \infty} x^s = x$, it holds that: $\lim_{s \rightarrow \infty} f(x^s) = f(x)$.

A function $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ is **continuous on** $\mathbb{R}^N$ if it is continuous at every point, $x \in \mathbb{R}^N$.

Weierstrass Maximum Theorem: Let $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ be a continuous function and $S \subset \mathbb{R}^N$ be a compact set. Then $\exists x \in S$ such that $\forall y \in S, f(x) \geq f(y)$ (A continuous function has a maximum on any compact set.)

Subsection B.4.2. Monotonic Functions

Monotonic Function: $f: \mathbb{R} \Rightarrow \mathbb{R}$ is a **monotonic function** if

$$\forall x, y \in \mathbb{R}, x \geq y \Leftrightarrow f(x) \geq f(y).$$

Note that if f is differentiable, monotonicity implies $\forall x \in \mathbb{R}, \frac{\partial f}{\partial x} \geq 0$.

Strictly Monotonic Function: $f: \mathbb{R} \Rightarrow \mathbb{R}$ is a **strictly monotonic function** if

$$\forall x, y \in \mathbb{R}, x > y \Leftrightarrow f(x) > f(y).$$

Note that if f is differentiable, strict monotonicity implies $\forall x \in \mathbb{R}, \frac{\partial f}{\partial x} > 0$.

Strictly monotone functions have the property that they preserve order. Thus, if we compose a strictly monotone function f with another function:

$$g: \mathbb{R}^N \Rightarrow \mathbb{R}: f \circ g(x) \equiv f(g(x))$$

we say that the resulting composite function is a **monotonic transformation** of g since:

$$\forall x, y \in \mathbb{R}^N, g(x) > g(y) \Leftrightarrow f(g(x)) > f(g(y)).$$

Subsection B.4.3. Homogeneous, Linear and Affine Functions

Linear Function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}^N$ is a **linear function** if $\forall x \in \mathbb{R}^N, \forall \alpha \in \mathbb{R}_{++}^N$:

$$f(\alpha x) = \alpha f(x).$$

Affine Function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}^N$ is an **affine function** if $\forall x \in \mathbb{R}^N, g(x) \equiv f(x) - f(0)$ is a linear function.

Homogeneous Function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}^N$ is **homogeneous of degree k** , if $\forall x \in \mathbb{R}^N, \forall \alpha \in \mathbb{R}_{++}^N$, and some $k \in \mathbb{N}$, $f(\alpha x) = \alpha^k f(x)$.

For example, let $\alpha \in \mathbb{R}_{++}^N$ and $\beta \in \mathbb{R}^N$, then $f(x) = \alpha x$ is a linear function and $f(x) = \alpha x + \beta$ is an affine function. Of course all linear functions are affine, but not the reverse. You can also verify that all linear functions, but not all affine functions, are homogeneous of degree 1.

Subsection B.4.4. Concave and Convex Functions

Concave Function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ is a **concave function** if $\forall x, y \in \mathbb{R}^N$ and $\forall \lambda \in (0,1)$:

$$\lambda f(x) + (1-\lambda)f(y) \leq f(\lambda x + (1-\lambda)y).$$

Strictly Concave Function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ is a **strictly concave function** if

$$\forall x, y \in \mathbb{R}^N \text{ and } \forall \lambda \in (0,1):$$

$$\lambda f(x) + (1-\lambda)f(y) < f(\lambda x + (1-\lambda)y).$$

Convex Function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ is a **convex function** if $\forall x, y \in \mathbb{R}^N$ and $\forall \lambda \in [0,1]$:

$$\lambda f(x) + (1-\lambda)f(y) \geq f(\lambda x + (1-\lambda)y).$$

Strictly Convex Function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ is a **strictly convex function** if $\forall x, y \in \mathbb{R}^N$ and $\forall \lambda \in (0,1)$:

$$\lambda f(x) + (1-\lambda)f(y) > f(\lambda x + (1-\lambda)y).$$

Quasi-Concave Function: $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ is a **quasi-concave function** if $\forall x, y \in \mathbb{R}^N$ such that $f(x) \geq f(y)$, and $\forall \lambda \in [0, 1]$:

$$f(\lambda x + (1-\lambda)y) \geq f(y).$$

Note that the utility function: $u(x) = x_1 x_2$ is not concave, since, for example:

$$\frac{1}{2}u(2,2) + (1-\frac{1}{2})f(1,1) = \frac{1}{2}4 + \frac{1}{2}1 = 2.5 > u(\frac{1}{2}(2,2) + \frac{1}{2}(1,1)) = u(1.5, 1.5) = 2.25,$$

You can check that if a utility function is quasi-concave, then the weakly preferred sets are convex. That is, $\forall x, y, z \in \mathbb{R}^N$ such that $\forall U(x), U(y) \geq U(z)$, and $\forall \lambda \in [0,1]$:

$$U(\lambda x + (1-\lambda)y) > (z).$$

We sometimes assume what are called **Inada Conditions** when it is useful to have extremely well-behaved functions. Formally, there are five separate requirements needed for any $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ to satisfy the Inada conditions:

en it is useful to have extremely well-behaved functions. Formally, there are five separate requirements needed for any $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ to satisfy the Inada conditions:

- $f(0) = 0$.
- f is continuously differentiable.
- f is strictly increasing all arguments: $\forall n \in \mathcal{N}, \frac{\partial f}{\partial x^n} > 0$.
- f is concave: $\forall n \in \mathcal{N}, \frac{\partial^2 f}{\partial (x^n)^2} < 0$.
- f is asymptotic to the boundary of the positive orthant:

$$\forall n \in \mathcal{N}, \lim_{x^n \rightarrow \infty} \frac{\partial f}{\partial x^n} = 0 \text{ and } \lim_{x^n \rightarrow 0} \frac{\partial f}{\partial x^n} = \infty.$$

Section B.5. Optimization

Subsection B.5.1. Optimization without Constraints

Consider a function $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ and suppose that all the first derivatives exist. We will often be interested in finding maxima and minima of such functions.

Formally:

Local Maximum: $x^* \in \mathbb{R}^N$ is a **local maximum** of f if for all x in a small enough neighborhood of x^* , $f(x^*) \geq f(x)$.

Local Minimum $x^* \in \mathbb{R}^N$ is a **local minimum** of f if for all x in a small enough neighborhood of x^* , $f(x^*) \leq f(x)$.

Extreme Points: The union of all local maxima and minima of a function f .

Global Maximum: The largest local maximum of a function f , if it exists.

Global Minimum: The smallest local minimum of a function f , if it exists.

In the simple case where $f: \mathbb{R} \Rightarrow \mathbb{R}$ is a function of only one variable, $x^* \in \mathbb{R}$ is a local extreme point of f only if:

$$\frac{\partial f(x^*)}{\partial x} = 0.$$

This is called the **first order condition (FOC)** and is a **necessary condition** for x^* to be a local maximum or a local minimum.

Satisfying the zero derivative condition, however, is not enough on its own to conclude that a point must be a maximum or a minimum, either local or global. In fact, the condition is also satisfied at “inflection points” of the functions. An **inflection point** is a point in the domain of f where the curvature changes sign, for example, where a function goes from being convex to concave, or inversely. If a point in the domain satisfies the zero derivative condition, it is called a **critical point**.

Determining whether a critical point is local maximum or minimum requires looking at the sign of the second derivative. In contrast to the first order conditions above, the **second derivative test** provides conditions that are sufficient, but not necessary, for a point to be local extrema.

$$\frac{\partial^2 f(x^*)}{\partial x^2} < 0 \text{ implies that } x^* \text{ is a local maximum.}$$

$$\frac{\partial^2 f(x^*)}{\partial x^2} > 0 \text{ implies that } x^* \text{ is a local minimum.}$$

$$\frac{\partial^2 f(x^*)}{\partial x^2} = 0 \text{ implies nothing in particular about } x^* .$$

In many economic situations, however, we will make additional assumptions on the objectives or constraints (generally known as **convexity assumptions**) that guarantee that there is one and only one extreme point, and that this must be a maximum or a minimum. In such a case, we do not need to check second order conditions, and we know that the extreme point is also a global extreme.

Let $f: \mathbb{R}^N \Rightarrow \mathbb{R}$ be a function of many variables. Then $x^* \in \mathbb{R}^N$ is a local extreme point of f only if $\forall n \in \mathcal{N}$

$$\frac{\partial f(x^*)}{\partial x_n} = 0.$$

The second order sufficient conditions are much more complicated in this case. We include them here only for completeness.

First we must form what is called the **Hessian matrix**. This is an array of the second derivatives of the function. We know by **Young's Theorem** that this matrix is symmetric around the main diagonal:

$$\forall m, n \in \mathcal{N}, \frac{\partial^2 f(x)}{\partial x_m \partial x_n} = \frac{\partial^2 f(x)}{\partial x_n \partial x_m}.$$

The Hessian matrix is written formally:

$$H = \begin{bmatrix} \frac{\partial^2 f(x)}{\partial x_1 \partial x_1} & \frac{\partial^2 f(x)}{\partial x_1 \partial x_2} & \cdots & \frac{\partial^2 f(x)}{\partial x_1 \partial x_N} \\ \frac{\partial^2 f(x)}{\partial x_2 \partial x_1} & \frac{\partial^2 f(x)}{\partial x_2 \partial x_2} & \cdots & \frac{\partial^2 f(x)}{\partial x_2 \partial x_N} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 f(x)}{\partial x_N \partial x_1} & \frac{\partial^2 f(x)}{\partial x_N \partial x_2} & \cdots & \frac{\partial^2 f(x)}{\partial x_N \partial x_N} \end{bmatrix}.$$

Let $|H|$ denote the **determinate** of the matrix H .

Then $\forall n \in \mathcal{N}$, let H^n denote a **principle minor** of H .

Generically, H^n is defined as the upper left-hand submatrix of H out to the n^{th} row and the n^{th} column and deleting the remaining rows and columns.

Given this:

H is **positive definite** if and only if $|H^n| > 0 \ \forall n \in \mathcal{N}$.

H is **negative definite** if and only if $(-1)^n |H^n| < 0 \ \forall n \in \mathcal{N}$.

The second order conditions are as follows:

If H evaluated at x^* is **negative definite**, then x^* is a local maximum of f .

If H evaluated at x^* is **positive definite**, then x^* is a local minimum of f .

Subsection B.5.2. Optimization with Constraints

Very frequently in economics, we wish to find a maximum or minimum for a function within some set of constraints. For example, consumers maximize their utility functions but must not violate their budget constraints, while firms minimize their costs within the constraints of their production technology the cost of inputs. Fortunately, we can use the **method of Lagrange** to solve such problems.

Consider an objective function $f: \mathbb{R}^N \Rightarrow \mathbb{R}$, and suppose that all the first derivatives exist. Suppose we want to find a maximum, but we must also satisfy a set of M differentiable equality constraints of the form $g_m: \mathbb{R}^N \Rightarrow \mathbb{R} \ \forall m \in \mathcal{M}$. Formally, the problem is this:

$$\max f(x) \text{ subject to } g_m(x) = 0 \ \forall m \in \mathcal{M}.$$

We can rewrite this using **Lagrangian multipliers** as follows:

$$\max f(x) + \sum_{m \in \mathcal{M}} \lambda_m g_m(x).$$

Notice the following:

- We have converted a constrained maximization problem into an unconstrained one. Thus, the same methods described above can be used, including the application of the necessary and sufficient conditions defined for extrema.
- We have added M new variables. In addition to taking derivatives with respect to x_n for all $n \in \mathcal{N}$, we also must take derivatives with respect to $\lambda_m \ \forall m \in \mathcal{M}$.
- Sometimes constraints (budget constraints, for example) are of the form $px = w$, that is, they equal a constant different from zero. This is not a problem, of course, since we can just define:

$$g(x) = px - w = 0.$$

- Taking these derivatives gives the following first order conditions (FOCs):

$$\frac{\partial f(x)}{\partial x_n} + \sum_{m \in \mathcal{M}} \lambda_m \frac{\partial g_m(x)}{\partial x_n} = 0 \quad \forall n \in \mathcal{N}$$

$$g_m(x) = 0 \quad \forall m \in \mathcal{M}.$$

Now notice two more things:

- We have $N+M$ FOCs with $N+M$ unknowns (the new unknowns are the Lagrangian multipliers).
- From the second set of FOCs, you can see that the M constraints must be satisfied at the solution.

What do these Lagrangian multipliers mean? It turns out they have a clear interpretation. Each λ_m gives the marginal increase or decrease in the value of the objective function in the neighborhood of a local maximum or minimum that would be result from an infinitesimal relaxing of given constraint, n . Take the simple case of a utility function and a single budget constraint, for example. Then the λ in front of the budget constraint equals the marginal increase in utility resulting from a small increase in income, in other words, the marginal utility of income.

One final note: The method of Lagrange solves maximization problems that have equality constraints. That is, all constraints are equations that must be exactly satisfied. Many problems call for inequality constants instead. For example, agents don't need to spend all of their income, but they can't spend more, so the constraint is really: $px \leq w$. Similarly, agents can't consume negative amounts of any good, so we should add this constraint: $x \geq 0$. Incorporating inequality constraints requires using the **Kuhn-Tucker Theorem**, which is a generalization of the **Theorem of Lagrange**. Take this as a pointer for further reading, but we won't go into greater detail here.

Appendix C. Logic and Proofs

Section C.1. Logical Statements and Propositions

Statement: A declarative sentence or a proposition that has a **truth value** which must either be **true** or **false**.

zzz

Compound Statement: A composition of component **atomic statements**, joined by **logical connectives** such as “and”, “or”, or “not”, that collectively have a truth value.

Proposition: A compound statement built from component atomic statements, joined by **logical connectives** such as “if”, “only if”, or “if and only if” that collectively have a truth value.

Examples of statements:

- $(x \geq_i \bar{x})$ and (my dog is brown) are atomic statements.
- $((x \geq_i \bar{x}) \text{ or } (\bar{x} \not\geq_i x))$ and ((my dog is brown) and (my dog has four legs)) are compound statements.
- (if $(x > \bar{x})$ then $(x \geq_i \bar{x})$) and (if (my dog is brown) then (my dog is four legs)) are propositions.

Examples of sentences are not statements.

- (make me a sandwich) is neither true nor false. It is an imperative sentence, or a command.
- (John is very beautiful) is neither true nor false unless we have a precise, logical, way to quantify both “beauty” and the meaning of “very” as a comparative. Without these, it is just an opinion.

Propositional statements are written:

$$P : \text{If } A \text{ then } B$$

and often contain components that we identify as an **hypothesis** and **conclusion**. In this case, A is the hypothesis and B is the conclusion.

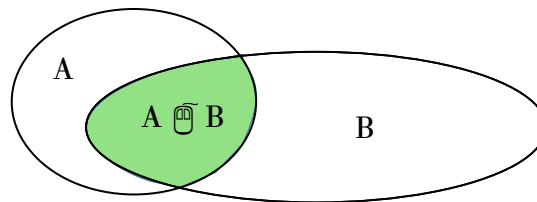
Section C.2. Basic Logical Connectives

The three essential logical connective used to construct compound statements are “and”, “or”, and “not”. Each of these connectives can be used in ordinary English or written in symbolic logic, and have corresponding interpretations in set theory.

Logical “And”: The logical “and” means that two or more statements connected by “and” must all be true for the compound statement as whole to be true.

We can write this in three ways:

- A and B
- $A \wedge B$
- $A \cap B$ (the intersection of the sets A and B)



For example:

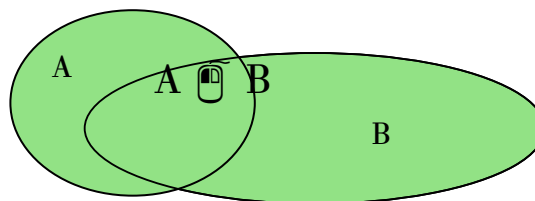
$(\text{water is wet}) \wedge (\text{I am mostly made of water}) \wedge (\text{the rain falling on my planet is mostly water})$

is a compound statement that has a value of true if I am a human on Earth. On the other hand, if I am a rock, the second statement is false, and so the whole compound statement is false. Similarly, if I am a human on Jupiter, the compound statement is false (scientists think it rains diamonds on Jupiter).

Logical “Or”: The logical “or” means that at least one or more statements connected by “or” must all be true for the proposition as whole to be true.

We can write this in three ways:

- A or B
- $A \vee B$
- $A \cup B$ (the union of the sets A and B)



For example:

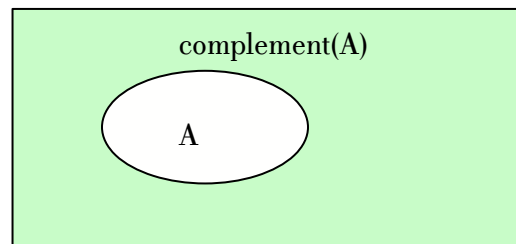
$(\text{water is wet}) \vee (\text{I am mostly made of water}) \vee (\text{the rain falling on my planet is mostly water})$

is a compound statement that has a value of true regardless of whether I am human or on Earth. As long as water is wet, the compound statement is true.

Logical “Not”: The logical “not” means the that a statement is false. Thus, if A is true, then the negation of A is false.

We can write this in four ways:

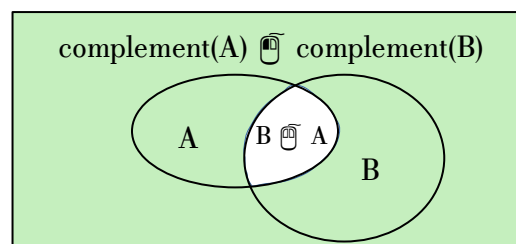
- A is false
- $\neg A$
- $\text{complement}(A)$ (complement of the set A)
- $x \notin A$



The logical $\wedge$ and $\vee$ are **dual** to each other under negation. For example consider the negations of compound statement $((A) \wedge (B))$.

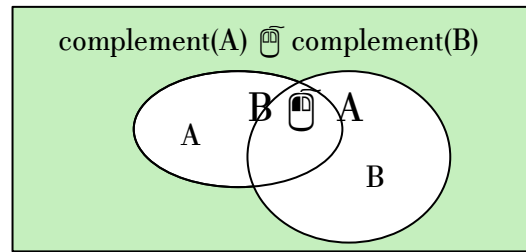
We can write this in three ways:

- $\neg((A) \wedge (B)) \Leftrightarrow ((\neg A) \vee (\neg B))$ (“ $\Leftrightarrow$ ” means that the two statements are logically equivalent.)
- $\text{complement}(A) \cup \text{complement}(B)$
- $x \notin A \cap B$



Similarly, the negation of $((A) \vee (B))$ can be written:

- $\neg((A) \vee (B)) \Leftrightarrow ((\neg A) \wedge (\neg B))$
- $\text{complement}(A) \cap \text{complement}(B)$
- $x \notin A \cup B$



Section C.3. Propositional Connectives

The three basic connectives used to construct propositions are “if”, “only if”, or “if and only if”. We define the meaning of these connectives via **truth tables**.

Truth Table: A truth table considers every possible combination of **truth values** of each atomic statement in proposition, and then determines the truth value of proposition given the definitions of the connectives. Truth tables are also used **define** the meaning of connectives as primitive elements of logic.

Logical “If”: The logical “if” creates a proposition defined as by the following truth table:

A	B	P
T	T	T
T	F	F
F	T	T
F	F	T

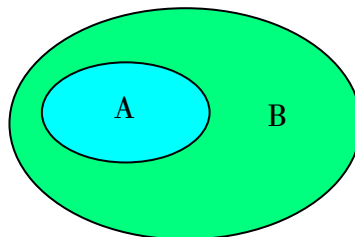
Truth Table for $P : A \Rightarrow B$

For example, suppose that A and B are both true. Then the proposition “ A implies B ” is defined to be true by this truth table. On the other hand, if A is true, but B is false, we have a counterexample to the proposition that “ A implies B ” and so the proposition is false, by definition.

The logical if can be written as follows:

- $A \Rightarrow B$
- A implies B
- if A then B
- A is a **sufficient condition** for B
- $A \subseteq B$

In the language of set theory, $A \Rightarrow B$ is equivalent to $A \subseteq B$. To understand this, note that if $A \subseteq B$, then $(x \in A)$ is a sufficient condition for $(x \in B)$. We can illustrate this with a Venn diagram.



$$A \subseteq B$$

Note that in the Venn diagram, $A = \emptyset$ is equivalent to A having a truth value of F . Notice that the empty set is always trivially and element of every other set. Similarly, if a statement A is false, then $A \Rightarrow B$ is trivially true. Consider the following proposition, for example:

Proposition: All 8' tall economists are married to Brittney Spears.

This is a true proposition! There are no 8' tall economists. Since we can't point to any 8' economists who are not married to Brittany Spears, it follows that all existing 8' economists are, in fact, married to Brittney Spears. It is also true that "all lions who speak English live in Boston". On the other hand, the converse: "all lions who live in Boston speak English" is probably false, assuming there is at least one lion who lives in Boston, and given that very few lions anywhere speak English.

Logical "Only If": The logical "only if" creates a proposition defined as by the following truth table:

A	B	P
T	T	T
F	T	F
T	F	T
F	F	T

Truth Table for $P : A \Leftarrow B$

As you might expect, this table is exactly like the first one except the truth values in the first two columns are swapped. This is because $A \Leftarrow B$ is the logical converse of $A \Rightarrow B$.

Converse: The converse of a proposition is formed by inverting the hypothesis and the conclusion. We denote the converse of a proposition P by P^c .

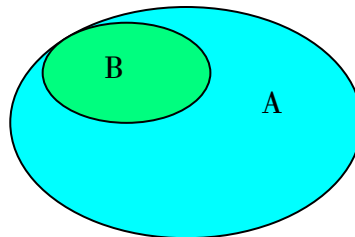
For example, the converse of $P : A \Rightarrow B$ is written $P^c : B \Rightarrow A$ or, equivalently, $P^c : A \Leftarrow B$

The logical only if can be written as follows:

- $A \Leftarrow B$
- A is implied by B
- A only if B
- A is a **necessary condition** for B
- $B \subseteq A$

or using the equivalent converse:

- $B \Rightarrow A$
- B implies A
- if B then A
- B is a **sufficient condition** for A



$$B \subseteq A$$

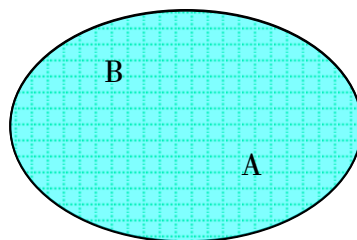
Logical “If and Only If”: The logical “if and only if” creates a proposition defined as by the following truth table:

A	B	P
T	T	T
T	F	F
F	T	F
F	F	T

Truth Table for $P : A \Leftrightarrow B$

The logical if and only if can be written as follows:

- $A \Leftrightarrow B$
- A is implied, and implied by B
- A if and only if B
- A is a **necessary and sufficient condition** for B .
- $A = B$



$$B = A$$

Negation: The negation of a statement A is denoted $\neg A$ and is read “not A ”

If the negation of a statement can be shown to be false, then the statement itself must be true. For example, (John is 6 feet tall) being true is logically identical to (John is not 6 feet tall) being false.

Proof by Contradiction: Proving that a proposition is true by proving that its negation is false.

The negation of the proposition $P: A \Rightarrow B$ is $\neg P: A \Rightarrow \neg B$. As above P is true, if and only if $\neg P$ is false. Thus, a proof by contradiction would start by saying suppose P is false. Then $\neg P$ must be true. However, I can show $\neg P$ is false. Therefore, P is true.

For example suppose A = (the sun is shining) and B = (it is daytime), and so proposition: $P: (the\ sun\ is\ shining) \Rightarrow (it\ is\ daytime)$. Note that the proposition always has truth value of true if the sun is not shining. The proposition only tells what must be true if the sun is shining.

The negation is $\neg P: (the\ sun\ is\ shining) \Rightarrow (it\ is\ not\ daytime)$. Suppose this is true. We observe that when the sun is shining it is always daytime. Thus, $(the\ sun\ is\ shining) \Rightarrow (it\ is\ not\ daytime)$, or $A \Rightarrow \neg B$, if proven to have a truth value of false, implying $P: A \Rightarrow B$ is true.

We will have more to say about this in the next Section.

Contrapositive: The contrapositive of a proposition is formed by negating both the hypothesis and the conclusion, and then *inverting* them. We denote the contrapositive of a proposition P by P' .

For example, the contrapositive of $P: A \Rightarrow B$ is $P': \neg B \Rightarrow \neg A$ which can be read as:

- if not B then not A
- if B is false, then A is false.
- $\text{complement}(B) \subseteq \text{complement}(A)$

Inverting the negated statements is key, and what makes the contrapositive different from a simple negation. As we say above, proposition is true if and only if its negation is false. On the other hand a proposition is true if and only if its contrapositive is true. We can construct a truth table for the contrapositive, and see that it is logically identical to the original proposition:

$\neg B$	$\neg A$	$\neg P$
T	T	T
T	F	F
F	T	T
F	F	T

Truth Table for P' : $\{\neg B \Rightarrow \neg A\}$

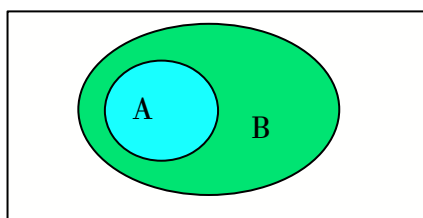
c

A	B	P
T	T	T
T	F	F
F	T	T
F	F	T

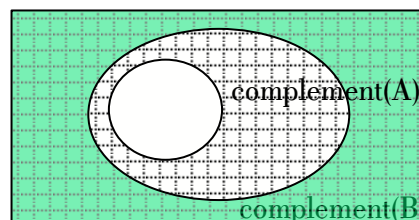
Truth Table for P : $A \Rightarrow B$

For example, suppose that both A and B are true, implying the P is true. Then $\neg B$ and $\neg A$ are false, also implying the P' is true. Similarly, if A is true, and B is false, then the P is false. In this case, $\neg A$ would be false, and $\neg B$ would be true, implying the P' is false.

This equivalence means that one strategy for proving theorems or propositions is to demonstrate the truth of its contrapositive. This is often much easier than the direct proof. Looking at the Venn diagram, we can see the same thing:



c



$$A \subseteq B \text{ c } \text{complement}(B) \subseteq \text{complement}(A)$$

The complement of A with respect to the rectangle is the crosshatched area. The complement of B with respect to the rectangle is the green area. You can see that the green area is a subset of the crosshatched area if and only if $A \subseteq B$.

Section C.4. Quantifiers

Consider the following statement using the quantifier “for all”

$$\forall x \in S, P(x) \text{ is true.}$$

This is a compound statement which says the following: for every element x in the set S , the statement $P(x)$ is true. For example, if $S = \{1, 2\}$ then this is equivalent to the statement:

$$P(1) \wedge P(2).$$

Example: “for all people in this classroom, their current home is Tennessee” is a true statement.

Consider instead a statement using the quantifier “there exists”

$\exists x \in S$, such that $P(x)$ is true.

This is a compound statement in which says the following: for at least one element x in the set S , the statement $P(x)$ is true. For example, if $S = \{1, 2, 3\}$ then this is equivalent to the statement:

$$P(1) \vee P(2) \vee P(3)$$

Example: “there exist a person in this classroom who is from Texas” is a true statement. I happen to be from Texas, and I am in this classroom. Are you from Texas too?

These two quantifiers have a dual relationship under negation:

$$\neg(\forall x \in S, P(x)) \Leftrightarrow (\exists x \in S, \neg P(x))$$

and

$$\neg(\exists x \in S, P(x)) \Leftrightarrow (\forall x \in S, \neg P(x))$$

Section C.5. Proofs

There are many ways to approach proving propositions. The following are the most common:

Direct Proof: A direct proof of $P: A \Rightarrow B$ starts with the hypothesis A , and use the rules of logic with other true statements to derive implications until finally, B is implied. Thus, the chain of logic proves $A \Rightarrow B$.

Example: $A \Rightarrow A_1, A_1 \Rightarrow A_2, A_2 \Rightarrow B$. Thus, you would show that statement A implies statement A_1 , and that statement A_1 implies statement A_2 , and finally that statement A_2 implies statement B . Thus, $A \Rightarrow B$.

Suppose the A = (class is over) and B = (Professor Conley will have a beer). How would we show $P: A \Rightarrow B$ in this case? Let's start by restating the B as (Professor Conley $\in$ {people having a beer}). The proposition is then:

Proposition: If (class is over), then (Professor Conley $\in$ {people having a beer}).

Proof:

If (class is over) then Professor Conley $\in$ {tired people}

If person $\in$ {tired people} $\cap$ {Irish people} then person $\in$ {people having a beer}

Professor Conley is a member of the Loyal Order of Hibernia.

Thus, Professor Conley $\in$ {Irish people} $\subset$ {people}

Therefore, if (class is over) then Professor Conley $\in$ {tired people} $\cap$ {Irish people}

We conclude that if (class is over) then Professor Conley $\in$ {people having a beer}

QED

QED: *quod erat demonstrandum*, which means literally, “that which was to be demonstrated” is a common way of indicating a proof is concluded.

A direct proof of $A \Leftarrow B$ would proceed in an analogous way:

$$A \Leftarrow A_1, A_1 \Leftarrow A_2, A_2 \Leftarrow A_3, A_3 \Leftarrow B.$$

A direct proof $A \Leftrightarrow B$ could proceed this way:

$$A \Leftrightarrow A_1, A_1 \Leftrightarrow A_2, A_2 \Leftrightarrow A_3, A_3 \Leftrightarrow B$$

but it is more common to break it into two propositions, $A \Rightarrow B$ and $B \Rightarrow A$, and prove them separately.

Proof by Contrapositive: Construct a direct proof that says that if B is false, A must be false. Formally: $\neg B \Rightarrow \neg A$.

For the example above, this would be rendered as:

Proposition: If (Professor Conley $\notin$ is not having a beer) then (class is not over),

Proof: Give it a try.

Note that proving this proposition is true, proves the original proposition is true.

Proof by Contradiction: If the proposition $A \Rightarrow B$ is true if and only if its negation is false. Thus, we can try to prove $\neg P: A \Rightarrow \neg B$ has a truth value of false.

Proposition: If (class is over), then (Professor Conley $\in$ {people having a beer}).

Proof:

Suppose the negation is true.

That is, suppose that (class is over) and Professor Conley $\notin$ {people having a beer} is true.

But we know:

$\forall$ person $\in$ {tired people} $\cap$ {Irish people}, it holds that person $\in$ {people having a beer}

We also know that (class is over) $\Rightarrow$ Professor Conley $\in$ {tired people} and that

Professor Conley $\in$ {Irish people} $\subset$ {people}

Thus, Professor Conley $\in$ {tired people} $\cap$ {Irish people} $\Rightarrow$

Professor Conley $\in$ {people having a beer}, contradicting the hypothesis that the negation is true.

QED

Proof by Induction: The goal here is to show that a proposition, $P(n)$, indexed by a set of integers, S , is true for all integers in the set. Rather than showing the proposition is true for every integer separately, an induction proof as follows is used:

Initial step: Show that a proposition $P(n)$ is true for the smallest integer in S (say $n = 1$).

Induction step: Show that if $P(n)$ is true for any arbitrary $n \in S$, it must be true for next largest integer in S (say if $P(n)$ is true, then $P(n+1)$ is true).

For example:

Proposition:: If you place K dominoes $\frac{1}{2}$ " apart in a ring and knock one over, they will all fall down.

Proof:

If I knock down domino #1, it will fall.

For all $n < K$, if domino # n falls, it will hit domino # $n+1$. Since any domino that is hit by another domino it will fall, domino # $n + 1$ will fall.

Therefore, if you knock down domino #1, all K dominoes will fall.

QED

Appendix D. Game Theory

Game theory is the study of rational agents who make strategic choices. This appendix will give the basic outline of three basic types.

Noncooperative Static Games: Games between agents who make simultaneous strategic choices.

These games take place in a single period and are sometimes called normal-form or one-shot games.

Noncooperative Sequential Games: Games between agents who make strategic choices at more than one point in time. Such games can take place over a finite, or infinite, number of periods. One special case is repeated play of one-shot games, where all players make simultaneous moves each period. In more general extensive-form games, one, or a subset of players, choose a strategy at a given time, after which, another player, or set of players, choose a strategy, until the game arrives at a terminal decision node, if they exist.

Cooperative Games: Cooperative games don't model strategic choice directly, but instead consider the properties of the solutions to games where the welfare of agents is in conflict. The idea behind this is that agents will participate and accept institutions that resolve disputes if they feel they are being treated fairly or appropriately in some way.

Section D.1. Noncooperative Static Games

Subsection D.1.1. Definitions

One-Shot, Simultaneous Move Game: $(\mathcal{I}, S, F)$ where:

Players or Agents: $i \in \{1, \dots, I\} \equiv \mathcal{I}$

Strategies: $s = \{s_1, \dots, s_I\} \in S_1 \times \dots \times S_I \equiv S$ where $s_i \in S_i$

Payoff Functions: $F \equiv (F_1, \dots, F_I)$ where $F_i: S \Rightarrow P$

We denote by P some **payoff space**. This could be any sort of finite or infinite set where the elements represent amounts of money or utility, market shares, discrete objects like houses or art, possible grades of a test, jobs, binaries such as winning or losing a game or a war, etc. To give some vocabulary:

Strategy Set for an Agent: S_i is called a strategy set for an agent i . This is also sometimes called the **Action Space**.

Strategy for an Agent: $s_i \in S_i$ is called a strategy for agent i .

Strategy Profile: $s = \{s_1, \dots, s_I\} \in S_1 \times \dots \times S_I \equiv S$ is called a strategy profile. In other words, a specific strategy choice for each agent.

Deviation Strategy Profile for Agent i : $(\bar{s}_i, s_{-i}) \equiv (s_1, \dots, s_{i-1}, \bar{s}_i, s_{i+1}, \dots, s_I)$ denotes the strategy profile $s \in S$ with the i^{th} element $s_i \in S_i$ deleted, and replaced with an alternative strategy $\bar{s}_i \in S_i$.

Subsection D.1.2. Equilibrium Concepts

The following are the three most important solution concepts for noncooperative one shot games.

Nash Equilibrium (NE): A strategy profile $s \in S$ is a Nash equilibrium if $\forall i \in \mathcal{I}$ and $\forall \bar{s}_i \in S_i$, $F_i(s_i, s_{-i}) \geq F_i(\bar{s}_i, s_{-i})$.

Dominant Strategy: We say $s_i \in S_i$ is a dominant strategy if $\forall \bar{s}_i \in S_i$ and $\forall s_{-i} \in S_{-i}$, $F_i(s_i, s_{-i}) \geq F_i(\bar{s}_i, s_{-i})$.

Dominant Strategy Equilibrium (DSE): We say $s \in S$ is a dominant strategy equilibrium if $\forall i \in \mathcal{I}$, $s_i \in S_i$ is a dominant strategy.

Subsection D.1.3. Mixed Strategies

So far, this outline has been quite general about the nature of strategies. It is conventional, however, to distinguish pure strategies from mixed strategies (although the strategies discussed above could, in fact, have been mixed strategies).

Suppose that for each agent $i \in \mathcal{I}$ in a one-shot game, there are finite number $A_i \in \mathbb{N}$ of pure strategies available in his strategy set: $\{s_i^1, \dots, s_i^{A_i}\} \equiv S_i$. Note that:

$$a_i \in \{1, \dots, A_i\} \equiv \mathcal{A}_i$$

is the index set for agent i 's strategy set.

Mixed Strategy: $p_i = (p_i^1, \dots, p_i^{a_i}, \dots, p_i^{A_i}) \in \Delta^{A_i-1}$ where $p_i^{a_i}$ is interpreted as the probability that agent i chooses strategy $s_i^{a_i} \in S_i$.

Pure Strategy: A pure strategy is a degenerate mixed strategy where $p_i^{a_i} = 1$, for some $s_i^{a_i} \in S_i$. and so, $p_i^{\hat{a}_i} = 0$, $\forall \hat{a}_i \neq a_i$.

Mixed Strategy Profile: $p = (p_1, \dots, p_i, \dots, p_I) \in \Delta^{A_1-1} \times \dots \times \Delta^{A_i-1} \times \dots \times \Delta^{A_I-1}$.

Note that given a mixed strategy profile p , the probability that agents will end up jointly playing any given pure strategy profile $(s_1^{a_1}, \dots, s_I^{a_I}) \in S_1 \times \dots \times S_I$ is:

$$\prod_{i \in \mathcal{I}} p_i^{a_i}.$$

Expected Value: The expected payoff or expected value to agent i of participating in any given a mixed strategy profile p is:

$$EV_i(p) = \sum_{(s_1^{a_1}, \dots, s_I^{a_I}) \in \mathcal{S}} \left(\prod_{i \in \mathcal{I}} p_i^{a_i} \right) F_i(s_1^{a_1}, \dots, s_I^{a_I}).$$

In words, for any given pure strategy profile: $(s_1^{a_1}, \dots, s_I^{a_I}) \in \mathcal{S}$, multiply the payoff that agent i receives, $F_i(s_1^{a_1}, \dots, s_I^{a_I})$, by the probability that $(s_1^{a_1}, \dots, s_I^{a_I}) \in \mathcal{S}$, ends up being played given the mixed strategy profile p , then add this expected payoff up over all possible pure strategy profiles in the game. Note that there are at total of $\prod_{i \in \mathcal{I}} A_i$ such pure strategy profiles.

Section D.2. Extensive Form Games

Subsection D.2.1. Definitions

Extensive Form Game: $(\mathcal{I}, \mathcal{A}, \mathcal{D}, \mathcal{T}, Player, Action, Next, S, F)$ where:

Players: $i \in \{1, \dots, I\} \equiv \mathcal{I}$

In the example: $\mathcal{I} \equiv \{Student, Professor\}$

Actions: $a \in \{1, \dots, A\} \equiv \mathcal{A}$. This is the set of every action available to any agent at any node.

In the example: $\mathcal{A} \equiv \{drink, cram, test, no test\}$

Decision Nodes: $d \in \{1, \dots, D\} \equiv \mathcal{D}$

In the example: there are $D = 3$ three decision nodes: $\mathcal{D} \equiv \{1, 2, 3\}$

Initial Node: Decision node 1, by definition.

In the example: the initial node belongs to the student who decides to either drink, or study.

Terminal Nodes: $t \in \{1+D, \dots, T+D\} \equiv \mathcal{T}$

In the example: there are $T = 4$ terminal nodes at the end of the game tree:
 $\mathcal{T} \equiv \{4, 5, 6, 7\}$. Terminal nodes have no successor nodes.

Player-Node Mapping: $Player: \mathcal{D} \Rightarrow \mathcal{I}$

In the example:

$$Player(1) = Student$$

$$Player(2) = Player(3) = Professor$$

Action-Node Mapping: $Action: \mathcal{D} \Rightarrow \mathcal{A}$

In the example:

$$Action(1) = \{drinking, cram\}$$

$$Action(2) = Action(3) = \{test, no\ test\}.$$

Next Node Mapping: $Next: \mathcal{D} \times \mathcal{A} \Rightarrow \mathcal{D}$

In the example:

$$Next(1, cram) = 2, \quad Next(1, drink) = 3$$

$$Next(2, test) = 4, \quad Next(2, no\ test) = 5$$

$$Next(3, test) = 6, \quad Next(3, no\ test) = 7$$

(Note that each of these four strategies must specify a feasible action for the professor for each of his decision nodes, even if those nodes are not ever seen in equilibrium.)

Strategies: $s = \{s_1, \dots, s_I\} \in S_1 \times \dots \times S_I \equiv S$, such that:

$$(a) \mathcal{D}_i \equiv \{d \in \mathcal{D} \mid Player(d) = i\}$$

$$(b) \mathcal{A}_i \equiv \{a \in \mathcal{A} \mid \exists d \in \mathcal{D}_i \text{ with } a \in Action(d)\}$$

$$(c) s_i: \mathcal{D}_i \Rightarrow \mathcal{A}_i \quad \forall d \in \mathcal{D}_i, s_i(d) \in Action(d)$$

In our example: $S_{Student}$ contains two strategies:

$$(i) s_{Student}(1) = drink$$

$$(ii) \bar{s}_{Student}(1) = cram$$

On the other hand, $S_{Professor}$ contains four strategies:

$$(i) s_{Professor}(2) = test, \quad s_{Professor}(3) = test$$

$$(ii) \hat{s}_{Professor}(2) = test, \quad \hat{s}_{Professor}(3) = no\ test$$

$$(iii) \bar{s}_{Professor}(2) = notest, \quad \bar{s}_{Professor}(3) = test$$

$$(iv) \tilde{s}_{Professor}(2) = no\ test, \tilde{s}_{Professor}(3) = no\ test$$

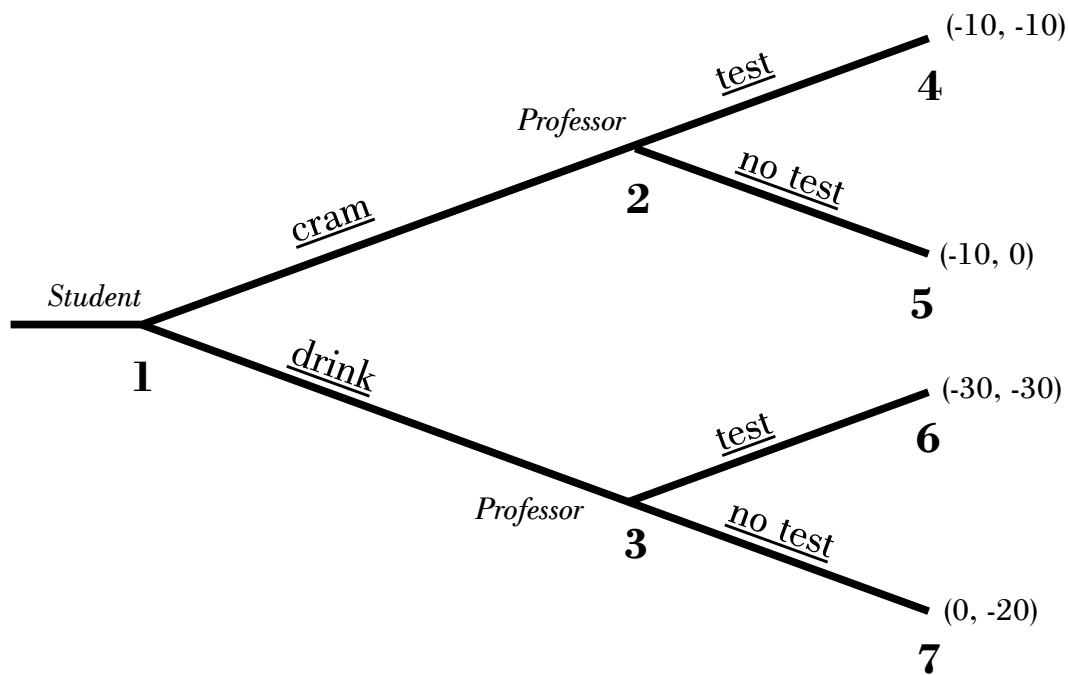
(Note that each of these four strategies must specify a feasible action for the professor for each of his decision nodes, even if those nodes are not ever seen in equilibrium.)

Payoff Functions: $F \equiv (F_1, \dots, F_I)$ where $\forall i \in \mathcal{I}, F_i: S \Rightarrow \mathbb{R}$

In our example:

$$F_{student}(\bar{s}_{student}, \tilde{s}_{professor}) = -10$$

$$F_{professor}(\bar{s}_{student}, \tilde{s}_{professor}) = 0.$$



Subsection D.2.2. Subgames and Subgame Perfect Equilibrium

Subgame: Let $(\mathcal{I}, \mathcal{A}, \mathcal{D}, \mathcal{T}, Player, Action, Next, S, F)$ be an extensive form game with perfect information. For any $d \in \mathcal{D}$, $(\mathcal{I}^d, \mathcal{A}^d, \mathcal{D}^d, \mathcal{T}^d, Player^d, Action^d, Next^d, S^d, F^d)$ denotes the subgame beginning at node d where:

Players: $\mathcal{I}^d \equiv \{i \in \mathcal{I} \mid \exists d \in \mathcal{D}^d \text{ such that } i = Player(d)\}$

Actions: $\mathcal{A}^d \equiv \{a \in \mathcal{A} \mid \exists \bar{d} \in \mathcal{D}^d \text{ such that } a \in Action(\bar{d})\}$

Decision Nodes: $\mathcal{D}^d \equiv \{\bar{d} \in \mathcal{D} \mid \bar{d} \notin \mathcal{T} \text{ and } \exists a_1, \dots, a_K \in \mathcal{A}, \text{ for some } K \in \mathbb{N}, \text{ such that } \bar{d} = \text{Next}(a_K \dots \text{Next}(a_2, \text{Next}(a_1, d)) \dots) \cup d\}$

Terminal Nodes: $\mathcal{T}^d \equiv \{t \in \mathcal{T} \mid \exists \bar{d} \in \mathcal{D}^d, \bar{a} \in \text{Action}(\bar{d}) \text{ such that } t = \text{Next}(\bar{a}, \bar{d})\}$

Player-Node Mapping: $\text{Player}^d(\bar{d}) = \text{Player}(\bar{d}) \forall \bar{d} \in \mathcal{D}^d$

Action-Node Mapping: $\text{Action}^d(\bar{d}) = \text{Action}(\bar{d}) \forall \bar{d} \in \mathcal{D}^d$

Next Node Mapping: $\text{Next}^d(\bar{a}, \bar{d}) = \text{Next}(\bar{a}, \bar{d}) \forall \bar{a} \in \mathcal{A}^d, \bar{d} \in \mathcal{D}^d$

Strategies: $s^d \in \prod_{i \in \mathcal{I}^d} S_i^d \equiv S$ such that $\forall i \in \mathcal{I}^d$:

- (a) $\mathcal{D}_i^d \equiv \{d \in \mathcal{D}^d \mid \text{Player}^d(d) = i\}$
- (b) $\mathcal{A}_i \equiv \{a \in \mathcal{A} \mid \exists d \in \mathcal{D} \text{ with } a \in \text{Action}(d)\}$
- (c) $s_i^d: \mathcal{D}_i^d \Rightarrow \mathcal{A}_i^d$ where $\forall \bar{d} \in \mathcal{D}_i^d, s_i^d(\bar{d}) \in \text{Action}^d(\bar{d})$

Payoff Function: $F^d(\bar{s}) = F(\bar{s}) \forall \bar{s} \in \mathcal{S}^d$

We can now state the definition of a continuation game for a subgame, and a refinement of Nash equilibrium called subgame perfection:

Continuation Strategy: Given $s \in S$ and a subgame starting at $d \in \mathcal{D}$, we call $s^d \in S^d$ the **continuation strategy from node d** if $\forall i \in \mathcal{I}^d$, and $\forall \bar{d} \in \mathcal{D}_i^d, s_i^d(\bar{d}) = s_i(\bar{d})$.

Subgame Perfect Equilibrium (SPE): A strategy profile $s \in S$ is a **subgame perfect equilibrium** if $\forall d \in \mathcal{D}$ (including $d = 1$), the continuation strategy s^d for the subgames beginning at node d , $(\mathcal{I}^d, \mathcal{A}^d, \mathcal{D}^d, \mathcal{T}^d, S^d, \text{Player}^d, \text{Action}^d, \text{Next}^d, F^d)$ is a Nash equilibrium and so $\forall i \in \mathcal{I}^d$, and $\forall \bar{s}_i^d \in S_i^d, F_i^d(s_i, s_{-i}) \geq F_i^d(\bar{s}_i, s_{-i})$.

Subsection D.2.3. Information Partitions

Everything above describes a game with complete information. However, it may be the case that players don't know exactly where they are in the game tree when they have to make a decision. This might be because, they can't observe the strategic choice the other player made at the previous node. This gives rise to the notion of an information partition:

Information Partition: $\text{Info}: \mathcal{D} \Rightarrow \{\text{subsets of } \mathcal{D}\}$ such that:

- (a) $\forall i \in \mathcal{I}, \forall d \in \mathcal{D}_i$ it holds that $d \subseteq \text{Info}(d) \subseteq \mathcal{D}_i$

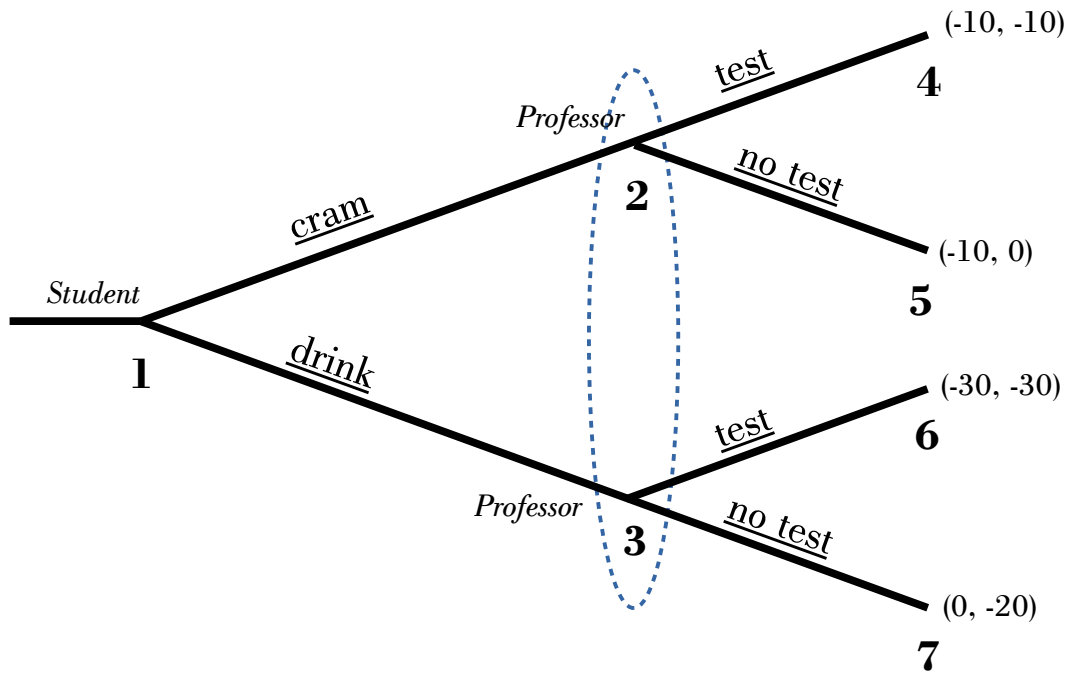
(b) $\forall i \in \mathcal{I}, \forall d, \bar{d} \in \mathcal{D}_i$:

$$Info(d) \cap Info(\bar{d}) = \emptyset \text{ or } Info(d) = Info(d) \cap Info(\bar{d}) = Info(\bar{d})$$

(c) $\forall i \in \mathcal{I}, \cup_{d \in \mathcal{D}_i} Info(d) = \mathcal{D}_i$

(d) $\forall i \in \mathcal{I}, \forall d, \bar{d} \in \mathcal{D}_i$, if $\bar{d} \in Info(d)$, then $A(d) = A(\bar{d})$

In our previous example of the student-professor game, the information partition is trivial since we have perfect information: $Info(1) = 1$, $Info(2) = 2$, $Info(3) = 3$.



$$Info(1) = \{1\}, \quad Info(2) = \{2, 3\}, \quad Info(3) = \{2, 3\}$$

We now need to revise our notion of strategies.

Strategies: $s = \{s_1, \dots, s_I\} \in S_1 \times \dots \times S_I \equiv S$ where $\forall i \in \mathcal{I}, \mathcal{D}_i \equiv \{d \in \mathcal{D} \mid Player(d) = i\}$:

(a) $\mathcal{A}_i \equiv \{a \in \mathcal{A} \mid \exists d \in \mathcal{D}_i \text{ such that } a \in Action(d)\}$

(b) $s_i: \mathcal{D}_i \Rightarrow \mathcal{A}_i \quad \forall d \in \mathcal{D}_i, s_i(d) \in Action(d)$

(c) $\forall d, \bar{d} \in \mathcal{D}_i$, if $Info(d) = Info(\bar{d})$ then $s_i(d) = s_i(\bar{d})$

In our new example with an information partition, $S_{Student}$ contains two strategies:

(i) $s_{Student}(1) = drink$

(ii) $\bar{s}_{Student}(1) = cram$

On the other hand, $S_{Professor}$ now contains only two strategies as well:

- (i) $s_{Professor}(2) = test$, $s_{Professor}(3) = test$
(ii) $\tilde{s}_{Professor}(2) = no\ test$, $\tilde{s}_{Professor}(3) = no\ test$

Section D.3. Cooperative Games

Subsection D.3.1. Games in Characteristic Function Form

Characteristic Function Form Game:

Players: $i \in \{1, \dots, I\} \equiv \mathcal{I}$

Coalitions: $c \subseteq \mathcal{I}$

Characteristic Function: $V: \mathcal{C} \Rightarrow \mathbb{R}$

Allocations: $x^c \in \mathbb{R}^{|c|}$

Blocking Coalition: We say that an allocation for the grand coalition $x^{\mathcal{I}}$ is **blocked** by an allocation x^c for coalition c if:

- (1) $\sum_{i \in c} x_i^c < V(c)$ (x^c is feasible for c)
- (2) $\forall i \in c, x_i^c \geq x_i^{\mathcal{I}}$ (all agents in the blocking coalition are at least as well off)
- (3) $\exists j \in c, x_j^c > x_j^{\mathcal{I}}$ (some agents in the blocking coalition are strictly better off)

In this case, we would say that c is a **blocking coalition**, and x^c is a **blocking allocation**.

Core Allocation: $x^{\mathcal{I}}$ is a **core allocation** if it cannot be blocked by any $c \in \mathcal{C}$. We also say $x^{\mathcal{I}}$ is **in the core of a game**.

Subsection D.3.2. Bargaining Theory

Bargaining Problem: $(S, d) \in \Sigma$ where $d \in S \subset \mathbb{R}^N$ and Σ is a **domain bargaining problems** satisfying:

- (1) S is compact.
- (2) There exists $\exists x \in S$ and $x \gg d$.

Two subdomains are the following:

Convex Bargaining Problem: $(S, d) \in \Sigma^{con} \subset \Sigma$ is an element of the domain of convex bargain problems if S is convex.

Comprehensive Bargaining Problem: $(S, d) \in \Sigma^{comp} \subset \Sigma$ is an element of the domain of comprehensive bargain problems if S is d-comprehensive.

Bargain solutions are defined for particular subdomains of problems. If it can be shown that there is one and only one solution that satisfies a given list of **axioms**, then the solution is said to be **characterized** by this list. We begin by defining three mapping that will be need to specify these axioms.

Bargaining Solution : $F: \tilde{\Sigma} \Rightarrow S$, such that $\forall (S, d) \in \tilde{\Sigma}$, $F(S, d) \in S$, where $\tilde{\Sigma}$ is a domain of bargaining problems, and F is single valued.

Axiom: Formal statements that offer various notions of fairness with regard to division or procedure.

Pareto Optimality (WPO): $F(S, d) \in WPO(S)$.

Independence of Irrelevant Alternatives (IIA): If $\bar{S} \subseteq S$, $\bar{d} = d$, and $F(S, d) \in \bar{S}$, then $F(\bar{S}, \bar{d}) = F(S, d)$.

Symmetry (SYM): If $\forall \pi \in \Pi^N$, $\pi(S) = S$ and $\pi(d) = d$, then $\forall i, j \in N$, $F^i(S, d) = F^j(S, d)$.

Scale Invariance (S.INV): $\forall \lambda \in \Lambda^N$, $F(\lambda(S), \lambda(d)) = \lambda(F(S, d))$.

Translation Invariance (T.INV): $\forall x \in \mathbb{R}^N$, $F(S+x, d+x) = F(S, d)+x$ (where $S+x$ means set addition).

Strong Monotonicity (S.MON): If $S \subset \bar{S}$ and $d = \bar{d}$, then $F(\bar{S}, \bar{d}) \geq F(S, d)$.

Restricted Monotonicity (R.MON): If $S \subset \bar{S}$, $d = \bar{d}$, and $a(S, d) = a(\bar{S}, \bar{d})$, then $F(\bar{S}, \bar{d}) \geq F(S, d)$.

Finally, we define the three most important solutions concepts for bargaining theory.

Nash Bargaining Solution: $N(S, d) \equiv \underset{\substack{x \in S \\ x \geq d}}{\operatorname{argmax}} (x_1 - d_1)(x_2 - d_2) \dots (x_N - d_N)$.

Theorem: A solution F defined on Σ^{con} satisfies WPO, SYM, IIA and S.INV if and only if it is the Nash bargaining solution.

Kalai-Smorodinsky Bargaining Solution is defined as:

$$K(S, d) \equiv \max \{ x \in S \mid x \in \text{con}(a(S, d), d) \}.$$

Theorem: A solution F on Σ^{comp} satisfies SYM, S.INV, WPO, and R.MON if and only if it is the Kalai-Smorodinsky solution.

Egalitarian Bargaining Solution is defined as:

$$E(S, d) \equiv \max \{ x \in S \mid \forall i, j \in \mathcal{I}, x_i = x_j \}.$$

Theorem: A solution F on Σ^{comp} satisfies SYM, T.INV, WPO, and R.MON if and only if it is the Egalitarian solution.