

1 – Introduction to microeconomics

There are two major parts of economics: microeconomics and macroeconomics. The names are somewhat misleading, since the former suggests small objects, while the latter – big research objects. Despite the terminology, both microeconomics and macroeconomics can address problems of the entire economy. The difference between the two is in a perspective adopted. Microeconomics is about economy as seen from the point of view of individual decision-makers. Macroeconomics is about economy as seen from the point of view of relationships between such aggregates as: GDP, exchange rate, unemployment, growth rate etc. Microeconomics tries to explain economic phenomena by asking questions on motivations decision-makers are likely to have when they do what they do. In contrast, macroeconomics is not necessarily interested in individual motivations. It tries to find relationships between various phenomena as they are observed.

Microeconomics is often considered more credible, because it offers more accurate predictions. Yet, despite frequently erroneous predictions, macroeconomics addresses very important questions. For instance, what will be the exchange rate of euro to dollar next month? Macroeconomics looks at what happens in European and American economies, and tries to predict future rates. Mistakes it makes are spectacular sometimes. Nevertheless the question is extremely interesting, and it is quite obvious that many people would like to see a specific answer rather than an analysis of hypothetical motives of various decision-makers. The distinction between microeconomics and macroeconomics is in a perspective adopted – not in an area studied.

Abstraction is a useful methodological approach to economics in general and microeconomics in particular. Abstraction means that we try to disregard details that are of secondary importance, and concentrate on what is of primary importance. In this process we make simplifying assumptions in order to get rid of unnecessary details. The selection of what is essential and what is not is very difficult, and it can be even ridiculed, but there is no better approach.

The power of abstraction can be envisaged when one looks at the development of physics over the last several thousand years. Aristotle claimed that movement of objects cannot be analysed, because there are too many contradictions involved. For instance if a force pushes an object in one direction then friction acts in the opposite direction, and so on. He concluded that movement is subject to a number of dialectical tendencies and it is beyond comprehension. His intellectual authority was so strong that he stopped the development of physics for two thousand years. Isaac Newton made simplifying assumptions like (1) the mass of an object is concentrated in one point, and (2) if no force acts, the object moves along a straight line with a constant speed. Of course, these assumptions can be ridiculed. For example, everybody knows that the mass of an object is not concentrated in one point. Nevertheless – by getting rid of unnecessary details, and by disregarding Aristotle's concerns about complexity – Newton managed to discover principles which allowed physics to start unprecedented development it has enjoyed over the last centuries.

The subject matter of economics is difficult, since it addresses human choices in routine voluntary transactions. Many people think of economics as something which gives answers to questions: "how to become rich?", "why is the price of a given stock sky-rocketing?", "what to do in order to eradicate poverty?" and so on. These are important questions, but they do not

delineate economics. For at least 100 years, most economists have understood their discipline as the science of making choices. Of particular interest are choices when people cannot have everything they would like to have, and if they wish to have something, they need to give up something else.

It is important to emphasise that economics is about choices made by ordinary people – not necessarily by the clever and the virtuous ones. We often say: "she must be an idiot if she spends so much on something" or "if he were ethical, he would not have done this". As citizens, neighbours, teachers, parents we can say so, but economists are not expected to judge choices made by other people. They are supposed to analyse people's decisions as they are, but they should not call them right or wrong.

Transactions establish typical proportions certain goods are exchanged for each other. For instance, two apples go for one orange. We say that one orange is worth two apples, or one apple is worth a half of an orange. Often one good becomes so common that people refer to it as a *numeraire*. They say: one orange is worth 1 €. The good referred to as a *numeraire* is called money. We are used to consider some currencies – like euro – as money. Yet many other goods served in this capacity for a long time: cattle, gold, vodka, man-hour, or what have you.

It is important to emphasise that economics is about routine transactions. Economists have little to say about transactions that are unique. If a criminal kills somebody "in exchange" for 100 € it does not mean that the victim's life was worth 100 €. Or if somebody spends 1 million € to cure herself it does not mean that the sickness was worth 1 million €. Only routine transactions are analysed in economics. Unique choices can be studied in psychology (or in other disciplines), but not in economics.

The last adjective in the description of economics reads "voluntary". Economics does not analyse transactions which are involuntary (that are forced). For instance, the behaviour of slaves in ancient Egypt, or prices charged in Soviet stores cannot be explained by economic methods. Voluntariness makes an important assumption, since many economic analyses are based on it. For example, economists argue that a purchase makes the buyer better off. Unless it is voluntary, the argument cannot be applied.

Thus economics is about how people make choices. They make choices every day, whenever they take decisions about what to do, what to buy, what to sell, whom to support, etc. *Homo oeconomicus*, 'economic man' (an often criticised concept) makes a convenient abstraction. Economists assume that people make choices that let them achieve well-being; people contemplate alternative decisions and choose what is most profitable given the circumstances. This does not preclude spontaneity, altruism, etc.

Optimization is a convenient behavioural assumption (not necessarily an empirical fact). Economists assume that economic agents maximise something (profit, utility, market share, satisfaction, 'power' etc.) or minimise something (costs, disutility, risk, etc.). We know that sometimes people make decisions that are difficult to explain – their decisions seem to make them worse off rather than better off. Nevertheless, when a large number of decisions is scrutinised, we usually discover, that there was a good reason why people did what they did.

Very often decisions are constrained by external circumstances. A typical economic problem indicates external constraints, such as production capacity, availability of credit, budget, etc. In the absence of such constraints, production or consumption could have been higher.

What was said above can be summarised by the following mathematical representation:

$$\text{Maximise } \{f(\mathbf{x}): \text{subject to } \mathbf{g}(\mathbf{x}) \leq \mathbf{0}\}$$

Standard assumptions made in microeconomics read:

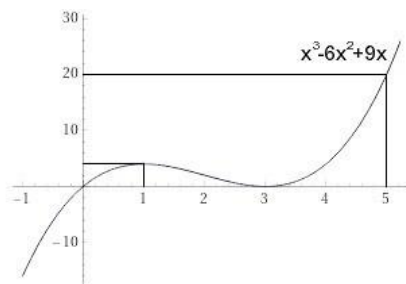
$$\begin{aligned} \mathbf{x} &\in \mathfrak{R}^n \\ f: \mathfrak{R}^n &\rightarrow \mathfrak{R} - \text{concave,} \\ \mathbf{g}: \mathfrak{R}^n &\rightarrow \mathfrak{R}^k - \text{convex.} \end{aligned}$$

If these assumptions are satisfied, then $\{\mathbf{x} \in \mathfrak{R}^n: \mathbf{g}(\mathbf{x}) \leq \mathbf{0}\}$ is a convex set, and any local maximum of f in this set is a global one.

Let us look at the following high-school example:

$$\text{Maximize } \{x^3 - 6x^2 + 9x: \text{subject to } x - 5 \leq 0\}.$$

Perhaps some students remember that in order to find a maximum one has to differentiate the function and to check where the derivative "vanishes" (equals to zero). Let us do it. The derivative is $3x^2 - 12x + 9$. This is a quadratic term; its roots are $x=1$ and $x=3$. Is any of these points the solution to our problem? It is not. If you make a graph you will easily see that for $x=3$ the function achieves its local minimum of 0 (in order to check that this is a minimum rather than a maximum, the second derivative can be calculated). For $x=1$ it achieves a local maximum of 4 which is less than 20 (when $x=5$). Thus the solution of the maximisation problem is 20 when $x=5$ (see the graph below). Calculating the derivatives was useless.



The high-school method works for so-called unconstrained maximisation, and the problem above was a constrained maximisation problem. Jean Luis Lagrange (1736-1813) – one of the greatest mathematicians – developed a powerful method to solve such problems. One needs to calculate derivatives of a somewhat different function called Lagrange function or 'Lagrangian'. For $\lambda \geq 0$, it is defined in the following way:

$$L(\mathbf{x}, \lambda) = f(\mathbf{x}) - \lambda \mathbf{g}(\mathbf{x}) = f(\mathbf{x}) - (\lambda_1 g_1(\mathbf{x}) + \dots + \lambda_k g_k(\mathbf{x})).$$

Auxiliary variables $\lambda_1, \dots, \lambda_k$ are called Lagrange multipliers; each corresponds to a constraint. Economists calculate so-called 'First Order Conditions' (FOC), i.e. conditions for the

derivatives of the Lagrange function to vanish. Under the standard (concavity/convexity) assumptions, FOC are sufficient to identify the solution. The FOC read:

$$\partial L / \partial \mathbf{x} = \mathbf{0},$$

that is:

$$\partial f / \partial x_1 - (\lambda_1 \partial g_1(\mathbf{x}) / \partial x_1 + \dots + \lambda_k \partial g_k(\mathbf{x}) / \partial x_1) = 0,$$

...

$$\partial f / \partial x_n - (\lambda_1 \partial g_1(\mathbf{x}) / \partial x_n + \dots + \lambda_k \partial g_k(\mathbf{x}) / \partial x_n) = 0.$$

Partial derivatives of L with respect to λ return respective coordinates of the original 'constraint' function $\mathbf{g}$.

Let us solve the numerical example we looked at before: Maximize $\{x^3 - 6x^2 + 9x\}$: subject to $x - 5 \leq 0$. Its Lagrange function reads: $L(x, \lambda) = f(x) - \lambda g(x) = x^3 - 6x^2 + 9x - \lambda(x - 5)$. Please note that in this case $x \in \mathbb{R}$ and $\lambda \in \mathbb{R}_+$ (they are numbers, not vectors). The partial derivatives of the Lagrange function are: $\partial L / \partial x = 3x^2 - 12x + 9 - \lambda$, and $\partial L / \partial \lambda = x - 5$. When equated to zero, the first one makes an equation ($3x^2 - 12x + 9 - \lambda = 0$) with two unknowns (x and λ). It seems that we substituted the original problem (with one unknown only) with a more complicated one. Indeed, it is difficult to solve one equation with two unknowns.

Lagrange's method offers a way out from this predicament. First we need to check what happens when $\lambda = 0$. Having substituted zero for λ , we get the equation $3x^2 - 12x + 9 = 0$. We already checked that it has two roots: $x = 1$ and $x = 3$. We need to reject the second one, since the second derivative reveals that this was a minimum rather than maximum. The first one is a local maximum, but f is not a concave function, so there is no guarantee that a local maximum found is the global one.

Thus we have to check what happens when $\lambda > 0$. Lagrange's method states that if a Lagrange multiplier is positive then the corresponding constraint is binding (it is satisfied as an equation rather than inequality). Hence we get an additional equation: $x - 5 = 0$. Finally we have a system of two equations with two unknowns. Its solution is $x^* = 5$ and $\lambda^* = 24$.

Everybody knows how to interpret $x^* = 5$. The interpretation of $\lambda^* = 24$ relates to the "opportunity cost" of the constraint $x - 5 \leq 0$. It informs about the slope of the function $x^3 - 6x^2 + 9x$ for $x = 5$. The slope coefficient of the straight line tangent to its graph in $(5, 20)$ is 24. It informs how the maximum value of the original function could change if the constraint was relaxed (say, $x \leq 5 + \varepsilon$ rather than the original $x \leq 5$). It measures sensitivity of the maximisation problem to relaxing constraints (a "marginal measure", as one of the students admitted correctly).

The Lagrange method is used in economics widely. The result of optimization yields an equilibrium, i.e. an outcome that cannot be improved by the optimizing agent. In typical economic problems, there is also a constraint $\mathbf{x} \geq \mathbf{0}$. This – in the form of $-\mathbf{x} \leq \mathbf{0}$ – can be included in $\mathbf{g}(\mathbf{x}) \leq \mathbf{0}$ (adding an additional constraint to the original ones). Despite this, economists prefer to deal with the $\mathbf{x} \geq \mathbf{0}$ constraint by applying the so-called Kuhn-Tucker theorem which is proved in the last lecture (QF-15).

Equilibrium is merely a reference point. Microeconomists do not claim that economies are in equilibrium; they may be in dis-equilibrium. Nevertheless microeconomists identify forces which motivate economic agents to undertake certain actions.

The equilibrium concept is sometimes ridiculed by observing that in the real world things are in dis-equilibrium which is a more typical state of affairs. An example of the water cycle is quoted. The river water flows whereas in "equilibrium" it should be motionless, stored in the ocean. Yet the fact that water molecules are forced to move down (by the gravity forces) explains why rivers flow from the mountains to the seas. Without the equilibrium concept in mind it would be difficult to understand why the nature does what it does.

Likewise in economics. It would be difficult to understand why consumers buy certain things, or why producers apply certain technologies unless their willingness to achieve some equilibrium is taken into account. In this course we will identify incentives economic agents have to optimise their predicament. Their activities are aimed at maximising some objectives subject to constraints established by the environment or by other agents' behaviour.

Questions and answers to lecture 1

1.1 Aristotle (385 BC - 323 BC) made a distinction between "chrematistics" (*Χρηματιστική*), and "economics" (*Οἰκονομικός* is a person in charge of managing a household). The former was about how to make profits, while the latter was about how to manage a household. Is this distinction relevant for contemporary economics?

It is questionable. Social critics complain that economists are preoccupied with money while neglecting problems of how to manage the world (the "household" of the humankind) in an appropriate way. They suggest that what economists do is "chrematistics" while "true economics" should deal with how to solve global problems. I disagree with this point of view. Managing the world in an appropriate way is an important problem indeed, but it goes much beyond what economics can analyse. It requires insights from physics, biology, psychology political science, theology, and perhaps a number of other scientific disciplines. At the same time, narrowing economics to making profits – a conviction revealed by non-economists very often – is not adequate either. What economists have analysed is how people make choices. This was explained by Lionel Robbins (1898-1984) most convincingly, but such an attitude was present already before. QF students are probably fond of "chrematistics". Yet I trust that this course of advanced microeconomics will demonstrate that there are a number of other fascinating economic problems beyond what can be observed in financial markets.

1.2 Should choices of ordinary people imply market prices?

Many people would prefer to face prices determined by the behaviour of clever consumers. Otherwise, if stupid or sinister people rush to buy harmful products (illegal drugs, cigarettes, weapons, etc.), their prices are driven up, and their suppliers have an incentive to continue and perhaps even to expand. If some products or services are harmful without any doubt, then their supply should be banned and they should not enter economic circulation at all. But there are thousands of products which seem to be questionable; people doubt whether they are useful. There is a temptation to override the market logic and to determine politically how to distribute goods. To some extent this is what we practice all over the world. For instance, we

insist that certain health services are available irrespective of whether patients pay for them or not. Nevertheless this method of economic organisation implies problems, and – when abused – works against people. After all, that is why centrally planned systems collapsed in Central and Eastern Europe. Despite what many of us feel, there is no better way to organise economy than allowing prices to be determined by ordinary people choices.

1.3 Is abstraction a good method of solving economic questions?

Yes, it is. If we tried to take into account every detail of everything, we could not make any useful statement, and we could not make any useful prediction. The conviction of ancient scholars that movement of objects was too complex to be studied provides an example of what may happen, if scientists wish to grasp everything at once. In order to find a useful explanation of phenomena observed, they need to make things easier by adopting simplifying assumptions. When Newton assumed that the mass of an object was concentrated in one point, he obviously knew that this was not true. Likewise, when we assume that firms maximise profits we know that this is not necessarily true. Firms may wish to accomplish something different, or within their management boards there are conflicting interests calling for different decisions. The world is complicated, and we cannot explain everything at once. Abstraction is inevitable if one wants to learn something useful about a process or a mechanism.

1.4 *Homo oeconomicus* (economic man) is a concept assuming that men prefer to be better off rather than worse off. It is criticised as an unrealistic assumption. Is the criticism justified?

Economists who rely on the *homo oeconomicus* concept do not necessarily believe that people can tell what makes them better off or worse off; the next lecture (QF-2) will address this problem in a greater detail. Neither do they believe that people always choose what makes them better off. *Homo oeconomicus* provides an abstraction. It helps to explain why the demand for certain goods is as it is. The criticism would be justified, if economists ignored the fact that people sometimes choose spontaneously and they regret this once they realise what they did. Spontaneous – perhaps irrational – decisions have to be studied by psychologists. Economists acknowledge such choices, but they can do little to explain them competently.

1.5 Why are FOC considered by economists an important characterisation of economic phenomena?

First Order Conditions identify solutions where partial derivatives of the Lagrange function are equal to zero. Calculating derivatives makes sense only for points that are not located at the border of the set where they can be found (called feasibility set). Points that are located at the boundary are called boundary ones (or corner solutions). The rest are called internal ones. In particular, in economic analyses we assume that $\mathbf{x} \geq \mathbf{0}$. Solutions $\mathbf{x}^*$ such that at least one coordinate is zero ($x_i^* = 0$) are corner solutions. For internal solutions First Order Conditions must be satisfied. If they are violated, then it would be possible to increase the value of the function to be maximised without leaving the feasibility set (which is impossible). Hence they are necessary. They are also sufficient if the function to be maximised is a concave one.

Solutions of optimisation problems are non-internal sometimes. First Order Conditions do not have to be satisfied then. Nevertheless in many applications solutions are internal ones, the

First Order Conditions must be satisfied, are therefore they provide useful information on what can be expected.

1.6 Prove that a strictly concave function cannot have two different local maxima over a convex set. Any local maximum of a concave function (not necessarily strictly concave) is a global one.

First we can prove that a concave function does not have two local maxima with different values. In other words, if x_1 and x_2 are local maxima of such a function then $f(x_1)=f(x_2)$. To see this, let us assume to the contrary, for instance, that $f(x_1)<f(x_2)$. Then the segment linking $f(x_1)$ and $f(x_2)$ cannot lie entirely below the graph of f , since it must lie above it in some small neighbourhood of x_1 . Thus if a concave function has two local maxima they need to be equal (this ends the proof of the second sentence), and – moreover – the entire segment linking $f(x_1)$ and $f(x_2)$ must coincide with the graph of f . Therefore what needs to be proved is that for a strictly concave function no portion of the graph can be flat. But this is precisely what the definition of strict concavity states ($f(\lambda y+(1-\lambda)z)>\lambda f(y)+(1-\lambda)f(z)$).

1.7 Is equilibrium a typical state of an economic variable?

No. It is not. Let us assume that according to analyses of demand and supply, p^* was found to be the equilibrium price, that is the price which makes the demand and the supply equal to each other. It may happen that indeed the market clears completely, and there are no incentives to change anything. A more typical state is that p^* turns out to motivate buyers to buy more than they did before (perhaps as a result of changing preferences) or to motivate sellers to sell more than they did before (perhaps as a result of technological progress). When the market senses these tendencies, the price is likely to rise in the first case, and to fall in the second case. Hence nothing is in equilibrium. Nevertheless, equilibrium is a useful reference for what happens. The demand goes up because the price faced was too low, or the supply goes up, because the price was too high. The assessments "too low" or "too high" refer to an equilibrium that may never materialise.

1.8 In the numerical example solved in the class, the original constraint is made more stringent. Originally it was $x \leq 5$; now it is $x \leq 4$. Was the value of the original Lagrange multiplier $\lambda^*=24$ useful in predicting the change of the maximum?

No. The original maximum was 20, and the Lagrange multiplier of 24 suggested that it would go down. It went down by 16, since the new maximum (for $x^*=4$) is equal to 4. If the slope of the tangent line was looked at, the anticipated decrease would be $(4-5) \times 24 = -24$. In fact it was smaller, because the shape of the function changed between 5 and 4 (the graph became less steep). Incidentally, the new maximum of 4 is attained for two points: for $x=1$, and for $x=4$. The first solution is an internal one, and the second – a corner one. First Order Conditions are satisfied in both cases: (when $x^*=4$, $\lambda^*=9$, and $3 \times 4^2 - 12 \times 4 + 9 - 9 = 0$; and when $x^*=1$, $\lambda^*=0$, and $3 \times 1^2 - 12 \times 1 + 9 + 0 = 0$).

2 – Consumer's choice

There are two important economic agents: consumers and producers. We will carry out microeconomic analyses of their behaviour. We will start with consumers. A consumption bundle is a key concept used in this analysis. It refers to what a consumer uses within a period of time. A bundle can consist of 3 kg of bread, a theatrical ticket, 10 litres of gasoline, 2 CDs with classical music, 0.5 kg of apples, and so on. A consumer is understood either as a physical person or a household. Interpretation depends on what decision making process is to be reflected. If one assumes that decisions are made by individual people then a consumer should be understood as a person. If one assumes that decisions are made collectively by people living in the same place, and perhaps controlling common money resources, then a household is a more appropriate interpretation. In our analyses we will not specify whether consumers are understood as physical persons or households.

Using mathematical notation, a consumption bundle is defined as a collection of goods $\mathbf{x}=(x_1,\dots,x_n)$ taken out of a consumption set $X: \mathbf{x}\in X\subset\Re^n$. Normally it is also assumed that $\mathbf{x}\geq\mathbf{0}$. It will be assumed that there is a preference relation defined on the consumption set. Interpretation of this relationship is consistent with common sense. If there are two bundles (they can be almost identical to each other), and, say, one includes three apples and one orange, and the other includes one apple and two oranges, a consumer is likely to prefer the first over the second, or *vice versa*. Symbol $\geq$ will be used to indicate that the consumer prefers one bundle over another one. In mathematical terms: preference relation (for a given consumer) is understood as a relation $\geq$ defined on the consumption set X .

One of the key concepts used in consumer theory is rationality of preferences. The formal definition reads:

Definition: Relation $\geq$ is rational if

1. $\forall \mathbf{x}, \mathbf{y} \in X [\mathbf{x} \geq \mathbf{y} \vee \mathbf{y} \geq \mathbf{x}]$ (completeness); and
2. $\forall \mathbf{x}, \mathbf{y}, \mathbf{z} \in X [(\mathbf{x} \geq \mathbf{y} \ \& \ \mathbf{y} \geq \mathbf{z}) \Rightarrow \mathbf{x} \geq \mathbf{z}]$ (transitivity).

Preference (and strict preference defined below) relations use the same symbols as arithmetic relations (greater than, and greater than or equal to). This, however, shall not lead to any doubts, since it will always be clear from the context whether the formula " $\mathbf{x} \geq \mathbf{y}$ " means "x is preferred over y" or "x is greater than or equal to y". If x and y are consumption alternatives (bundles), i.e. $\mathbf{x}, \mathbf{y} \in X$ then the formula reads "x is preferred over y", but if x and y are numbers, i.e. $x, y \in \Re$ then the formula reads "x is greater than or equal to y".

Preference relation $\geq$ allows to define several related concepts, including – most importantly – strict preference and indifference. Their mathematical definitions are:

Strict preference, $>$: $\mathbf{x} > \mathbf{y} \Leftrightarrow (\mathbf{x} \geq \mathbf{y} \ \& \ \neg \mathbf{y} \geq \mathbf{x})$, and indifference, $\sim$: $\mathbf{x} \sim \mathbf{y} \Leftrightarrow (\mathbf{x} \geq \mathbf{y} \ \& \ \mathbf{y} \geq \mathbf{x})$

Mathematical symbols used are perhaps well known by most students. Let me explain them just in case. $\forall \mathbf{x} \in X$ means "for every x in X"; $\vee$ stands for alternative ("or"); $\&$ stands for conjunction ("and"); $\Rightarrow$ stands for implication ("if ... then"); $\Leftrightarrow$ stands for equivalence ("if and only if"); and $\neg$ stands for negation ("not").

The following theorem summarises several easy consequences of the definitions introduced. Their proofs are trivial, as seen from proving "1". Proof of the first item is extremely simple. By the completeness property we know that $\forall \mathbf{x}, \mathbf{y} \in X [\mathbf{x} \geq \mathbf{y} \vee \mathbf{y} \geq \mathbf{x}]$. An alternative is true if one

of its elements is true. If $\mathbf{x}$ is substituted for $\mathbf{y}$, we get $\forall \mathbf{x} \in X [\mathbf{x} \geq \mathbf{x} \vee \mathbf{x} \geq \mathbf{x}]$ which completes the proof.

Theorem: If $\geq$ is rational, then:

1. $\forall \mathbf{x} \in X [\mathbf{x} \geq \mathbf{x}]$ (reflexivity of $\geq$)
2. $\forall \mathbf{x} \in X [\neg \mathbf{x} > \mathbf{x}]$ (counter-reflexivity of $>$)
3. $\forall \mathbf{x}, \mathbf{y}, \mathbf{z} \in X [(\mathbf{x} > \mathbf{y} \ \& \ \mathbf{y} > \mathbf{z}) \Rightarrow \mathbf{x} > \mathbf{z}]$ (transitivity of $>$)
4. $\forall \mathbf{x} \in X [\mathbf{x} \sim \mathbf{x}]$ (reflexivity of $\sim$)
5. $\forall \mathbf{x}, \mathbf{y}, \mathbf{z} \in X [(\mathbf{x} \sim \mathbf{y} \ \& \ \mathbf{y} \sim \mathbf{z}) \Rightarrow \mathbf{x} \sim \mathbf{z}]$ (transitivity of $\sim$)
6. $\forall \mathbf{x}, \mathbf{y} \in X [\mathbf{x} \sim \mathbf{y} \Rightarrow \mathbf{y} \sim \mathbf{x}]$ (symmetry of $\sim$)
7. $\forall \mathbf{x}, \mathbf{y}, \mathbf{z} \in X [(\mathbf{x} > \mathbf{y} \ \& \ \mathbf{y} \geq \mathbf{z}) \Rightarrow \mathbf{x} > \mathbf{z}]$

Completeness implies reflexivity of preferences (1). Hence rational preferences must be reflexive by definition. Some textbooks consider reflexivity so important, that they add it to the definition of rationality which – as we see – is not necessary. Students who like mathematics may note that items 4-6 (and completeness) imply that $\sim$ is an equivalence on X ; the set of all bundles X can be partitioned into equivalence classes and bundles belonging to different classes are not equivalent to each other. Those who do not like mathematics may disregard this information.

Preferences are often represented by so-called indifference curves. These are defined as sets of bundles that a given consumer is indifferent about. Their mathematical definition is very simple. Indifference curve $I(\mathbf{x}) = \{\mathbf{y} \in X: \mathbf{y} \sim \mathbf{x}\}$. It is easy to prove the following theorem about indifference curves.

Theorem: If $\geq$ is rational, then every two indifference curves are either identical or disjoint.

Proof:

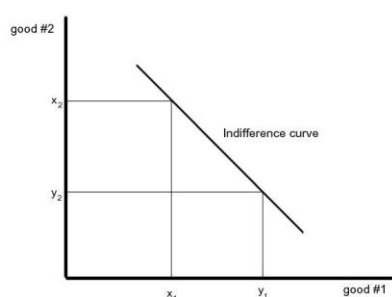
Let us assume that there are two indifference curves I_1 and I_2 , and $\mathbf{x}$ belongs to both of them ($\mathbf{x} \in I_1$, and $\mathbf{x} \in I_2$). This means that I_1 consists of all bundles $\mathbf{y}$ such that $\mathbf{y}$ is indifferent with $\mathbf{x}$, and I_2 consists of all bundles $\mathbf{y}$ such that $\mathbf{y}$ is indifferent with $\mathbf{x}$. Hence they must be the same.

The definition above is general, but the most popular examples refer to two-dimensional bundles ($X \subset \mathbb{R}^2$). If $n=2$, then for perfect substitutes (goods that can be exchanged for each other in fixed proportions) indifference curve are segments of a straight line; for perfect complements (goods that have to be consumed together in fixed proportions) indifference curves are L-shaped. Swiss francs (CHF), and British pounds (GBP) are examples of perfect substitutes; in September 2020 they could be exchanged according to the rate of 1.17 (1 GBP went for 1.17 CHF or 1 CHF went for 0.85 GBP). Microeconomic textbooks refer to coffee and sugar as an example of perfect complements; for many people one cup of coffee goes with two spoons of sugar (the two goods make sense only together).

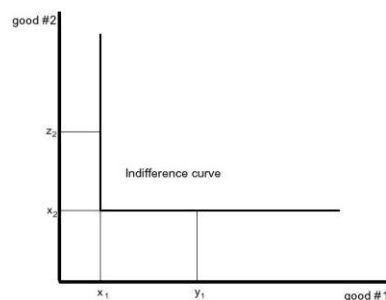
Pictures below illustrate indifference curves for perfect substitutes and perfect complements.

Picture on the left illustrates preferences with respect to perfect substitutes, like CHF and GBP. For instance (x_1, x_2) may be interpreted as 5 CHF and 8 GBP, while (y_1, y_2) as 9 CHF and 4.44 GBP. Both combinations are worth the same (up to rounding errors): 14.46 CHF or 12.27 GBP (13.45 €). Picture on the right illustrates preferences with respect to perfect

complements, like coffee and sugar. For instance (x_1, x_2) can be interpreted as one cup of coffee accompanied by two spoons of sugar. Availability of more sugar, say four (z_2) spoons, does not improve the situation of a consumer who has only one cup of coffee ((x_1, z_2) is on the same indifference curve as (x_1, x_2)). Likewise, availability of more coffee, say three (y_1) cups, does not improve the situation of a consumer who has only two spoons of sugar ((y_1, x_2) is on the same indifference curve as (x_1, x_2)).



Perfect substitutes



Perfect complements

Sometimes preferences are assumed to comply with specific axioms. Three of them are used fairly frequently.

- Monotonicity: If $\mathbf{x}=(x_1, \dots, x_n)$, $\mathbf{y}=(y_1, x_2, \dots, x_n)$, and $x_1 \geq y_1$, then $\mathbf{x} \geq \mathbf{y}$; likewise for goods 2, 3, ..., n
- Convexity: If $I(\mathbf{x})=I(\mathbf{z})$, and $\lambda \in [0, 1]$, then $\lambda \mathbf{x} + (1-\lambda)\mathbf{z} \geq \mathbf{x}$ (indifference curves for bundles consisting of two goods are graphs of convex functions)
- Continuity: $\forall k=1, 2, \dots \forall \mathbf{x}_k, \mathbf{y}_k \in \mathcal{R}^n [(x_k \geq y_k \ \& \ \mathbf{x} = \lim_{n \rightarrow \infty} \mathbf{x}_n \ \& \ \mathbf{y} = \lim_{n \rightarrow \infty} \mathbf{y}_n) \Rightarrow \mathbf{x} \geq \mathbf{y}]$

The first two ones are obvious. The third one requires a separate discussion. High-school graduates remember two definitions of continuity: The Cauchy (1789-1857) one, and the Heine (1821-1881) one. For all practical purposes they are equivalent. The latter is used in the definition of preference continuity above. So-called lexicographic preferences will be introduced in order to illustrate what may happen when preferences are not continuous.

As the name suggests, lexicographic preferences resemble procedures used by us when we look up for a word in a dictionary. We know that the word "bear" is before "beat" since the earliest letter that these words differ in is "r" in the first one, and "t" in the second one; moreover, we know that in the alphabet "r" comes before "t". Lexicographic preferences apply this logic. We will define lexicographic preferences for $X \subset \mathcal{R}^2$ only, but the definition can be easily generalised for $X \subset \mathcal{R}^n$ for $n > 2$ too.

Let $X \subset \mathcal{R}^2$. Then lexicographic preferences are defined in the following way: $(x_1, x_2) \geq_L (y_1, y_2)$ whenever $x_1 > y_1$; or $x_1 = y_1$ and $x_2 \geq y_2$.

For instance, $(1, 2) \geq_L (0, 3)$, but $(1, 2) \not\leq_L (1, 3)$.

Lexicographic preferences are not continuous. To see this, we can compare two sequences: $\mathbf{x}_n = (1/n, 0)$, and $\mathbf{y}_n = (0, 1)$. For any $n=2, 3, \dots$ $(1/n, 0) \geq_L (0, 1)$, that is $\mathbf{x}_n \geq_L \mathbf{y}_n$. $\mathbf{x}_0 = \lim_{n \rightarrow \infty} \mathbf{x}_n = (0, 0)$,

and $y_0 = \lim_{n \rightarrow \infty} y_n = (0, 1)$, but $x_0 <_L y_0$ (in the limit the ordering reverses). This is not only a theoretical curio. Some econometric methods to estimate demand functions assume that preferences are continuous. If there are not, econometric estimates are inadequate. Before applying such methods, researchers check whether the continuity assumption is justified.

Lexicographic preferences have to be referred to whenever consumers look at several criteria in certain order. Choosing holiday destination provides a classical example. Let us assume that before choosing a holiday destination, a consumer looks at the number of sunny days per month, and the lake temperature. Two destinations can be characterised as $x = (26, 24)$, and $y = (28, 21)$. $y \geq_L x$, because the number of sunny days is checked first, and the lake temperature afterwards. Please note that if the order was reversed, the second destination would turn out to be more attractive.

19th century economists were impressed by the development of physics, where everything could be quantified. Movement of objects were characterised by trajectories, speed, acceleration, and so on. They thought that if well-being was quantified, then the success of physics could be replicated by economics. Mere preference relationships are not sufficient, because they are not represented by numbers. A "utility" was considered a measure of well-being that can represent preferences. In mathematical language the definition reads: a utility function $u: X \rightarrow \mathbb{R}$ representing $\geq$ is any function such that $\forall x, y \in X [x \geq y \Leftrightarrow u(x) \geq u(y)]$. Several theorems can be proved for utilities easily.

Theorem: If there is a utility function representing $\geq$, then $\geq$ is rational.

Proof:

We need to prove completeness and transitivity for bundles x , y , and z . But both axioms can be derived from completeness and transitivity of numbers $u(x)$, $u(y)$, and $u(z)$.

Theorem: If u is a utility function representing $\geq$, and f is a strictly increasing function, then $f(u)$ is also a utility function representing $\geq$.

Proof:

We need to prove that $\forall x, y \in X [x \geq y \Leftrightarrow f(u(x)) \geq f(u(y))]$. The definition of utility u requires that $\forall x, y \in X [x \geq y \Leftrightarrow u(x) \geq u(y)]$. Strict (increasing) monotonicity of f guarantees that $\forall x, y \in X [f(u(x)) \geq f(u(y)) \Leftrightarrow u(x) \geq u(y)]$.

Theorem. Let u be a utility function representing $\geq$. If $x \in I(y)$ then $u(x) = u(y)$.

Proof:

By definition of indifference curves, if $x \in I(y)$ then $y \sim x$. By definition of indifference, it means that $y \geq x$ and $y \leq x$. By definition of utility, it means that $u(y) \geq u(x)$ and $u(y) \leq u(x)$ that is $u(y) = u(x)$.

It would be easy to prove that lexicographic preferences cannot be represented by a continuous utility function. In fact, they cannot be represented by any utility function, but such a theorem is somewhat more difficult to prove. This is not the only problem with

utilities. 19th century economists dreamt that one day it would be possible to measure utilities, as we measure temperature, tension or pulse. They dreamt that a meter could be constructed to indicate the utility of a bundle. Today we know, that this is not possible; utilities cannot be quantified in an unambiguous way. Nevertheless this term is used by economists routinely in order to characterise preference relationships. The concept of "Marginal Rate of Substitution" – to what extent one good can substitute for another one without making the consumer strictly better off or strictly worse off – serves as a useful link between preferences and utilities. The mathematical definition reads: Marginal Rate of Substitution (MRS) is the coefficient of the tangent line to the indifference curve. The following theorem relates utilities to indifference curves.

Theorem. Let u be a differentiable utility function representing $\succeq$ ($X \subset \mathbb{R}^2$). Then

$$MRS = -\partial u(x_1, x_2) / \partial x_1 : \partial u(x_1, x_2) / \partial x_2.$$

Proof:

An easy proof makes use of symbols "dy" and "dx". Economists (and engineers) apply the symbol "dy" in order to denote so-called differential of y . This notation was widely used by 17th and 18th century mathematicians who laid foundations for differential calculus. Later on it was used only in symbols such as "dy/dx", and examples were provided to prove that decoupling "dy" from "dx" can lead to a nonsense. Indeed it can. Nevertheless, when applied with caution, it can simplify formulae. For practical purposes, in order to calculate $df(x)$, one needs to calculate $df(x)/dx = g(x)$ and "multiply" both sides of the equation by dx .

Our proof starts by observing that a two dimensional indifference curve (we use the assumption that $X \subset \mathbb{R}^2$) gives x_2 as a function of x_1 . The tangent is thus the derivative of this function. By the definition of an indifference curve, the utility does not change, if one considers bundles located along this curve. In terms of differentials, it means that $du=0$. The so-called "complete differentiation" formula (applied to calculating derivatives of functions with many variables) allows us to calculate $du = (\partial u(x_1, x_2) / \partial x_1) dx_1 + (\partial u(x_1, x_2) / \partial x_2) dx_2$. But – as observed earlier – along an indifference curve this number must be zero. Thus

$$(\partial u(x_1, x_2) / \partial x_1) dx_1 + (\partial u(x_1, x_2) / \partial x_2) dx_2 = 0, \text{ that is } (\partial u(x_1, x_2) / \partial x_1) dx_1 = -(\partial u(x_1, x_2) / \partial x_2) dx_2.$$

By "dividing" both sides into dx_1 , one gets $(\partial u(x_1, x_2) / \partial x_1) = -(\partial u(x_1, x_2) / \partial x_2) dx_2 / dx_1$. By dividing this equation into $-(\partial u(x_1, x_2) / \partial x_2)$ we get the formula for MRS.

Questions and answers to lecture 2

2.1 If $x \sim y$ when $x \neq y$ does this mean that preferences are incomplete?

No. The indifference relationship, $x \sim y$ means that $x \succeq y$ & $y \succeq x$ (a consumer whose preferences are analysed here can compare bundles x and y). Incompleteness means that the consumer cannot compare these bundles. Completeness is a fairly strong assumption; it implies that no matter how strange bundles are to be compared (for instance x includes a stereo and no bicycle, and y includes no stereo and a bicycle) the consumer can always indicate the

preferred one (sometimes both bundles are preferred). Indifference does not mean that the bundles cannot be compared; they can, but there is no strict preference of one over the other.

2.2 Does item 7 in the theorem on page 7 state the transitivity of strict preferences?

Please note that item 7 reads $\forall x, y, z \in X [(x > y \ \& \ y \geq z) \Rightarrow x > z]$, while the transitivity of $>$ would read $\forall x, y, z \in X [(x > y \ \& \ y > z) \Rightarrow x > z]$. This is not the same, but one can observe that $y > z$ implies $y \geq z$. Hence the transitivity of strict preferences implies item 7, but not necessarily *vice versa*. Item 7 is in fact stronger than strict transitivity of $>$.

2.3 Can clever people have intransitive preferences?

Intransitivity of preferences is observed empirically sometimes. Rationality does not allow intransitivity and some analysts conclude that people who have such preferences must be stupid. Indeed, it is difficult to think of someone who confirms that $(x \geq y \ \& \ y \geq z)$ and at the same time states that $x < z$. Lexicographic preferences help to interpret such a strange combination.

Let us assume that a consumer has lexicographic preferences like explained in the class, but in addition he/she states that $x \geq y$ whenever $x_1 > y_1$ or $x_1 \in (y_1 - 0.5, y_1 + 0.5)$ and $x_2 \geq y_2$. In other words, x_1 does not have to be exactly equal to y_1 ; it can be slightly lower than y_1 or slightly higher than y_1 , and the ordering will be determined by the second coordinate. Let $x = (26, 24)$, $y = (25.7, 26)$, and $z = (25.4, 28)$. Hence $x \leq y$ and $y \leq z$ (in both comparisons the first coordinate was approximately the same, so the second coordinate was looked at in order to determine the preference. The difference between the first coordinate of x (26), and the first coordinate of z (25.4) started to exceed 0.5 thus $x > z$. I think that consumers who choose according to such a mechanism cannot be considered stupid.

2.4 Indifference curves for perfect complements (see picture on page 3) are not differentiable. Can MRS coefficients be defined in a reasonable way?

No. When you look at an L-shaped indifference curve you see that tangent lines are vertical along vertical parts of the curve, and they are horizontal along horizontal parts of the curve. In (x_1, x_2) , i.e. where the vertical part meets the horizontal part, it would be difficult to draw a tangent curve. In the first case, the slope coefficient of the tangent can be considered $-\infty$; in the second case 0; and it cannot be determined where the two parts meet. Intuitive interpretation of this in-determination results from the fact that the two goods cannot be substituted for each other. No matter how many units of the good #2 are added, if the number of units of the good #1 is unchanged, the new bundle is not better than the old one (both bundles are in the vertical part of the indifference curve). And *vice versa*. No matter how many units of the good #1 are added, if the number of units of the good #2 is unchanged, the new bundle is not better than the old one (both bundles are in the horizontal part of the indifference curve).

2.5 Indifference curves for perfect substitutes are segments of a straight line (see picture on page 3). What does this mean for MRS?

The MRS is fixed (it is the same irrespective of the bundle). It indicates how many units of the good #2 can be removed if one unit of the good #1 is added. This is precisely what the definition of perfect substitution includes.

2.6 Local satiation principle is defined in the following way:

$$\exists \mathbf{x} \in \mathbb{R}_+^n \exists \varepsilon > 0 \forall \mathbf{y} \in \mathbb{R}_+^n [|\mathbf{y} - \mathbf{x}| < \varepsilon \ \& \ \mathbf{y} < \mathbf{x}].$$

Prove that strictly monotonic preferences (strict monotonicity is defined by $>$ rather than $\geq$) cannot satisfy the local satiation principle.

Microeconomic textbooks quote ice cream balls in order to illustrate the local satiation principle. Many people like ice cream, and the more balls they eat the better they feel – but up to some limit. Two ice cream balls are preferred to one ice cream ball, three balls are preferred to two, four are preferred to three, and five are preferred to four. Yet, if a consumer is asked about the sixth ice cream ball (all the other goods consumed at the same level), he/she will probably prefer to stick to the fifth one. In this example consumer preferences are "satiated" at the level of five ice cream balls; if the consumer is offered more, then such a more abundant bundle is not preferred.

Now let us assume that consumer's preferences have a local satiation level $\mathbf{x} = (x_1, x_2, \dots, x_n)$. Let the first coordinate is where the strict monotonicity takes place, $\mathbf{y} = (y_1, x_2, x_3, \dots, x_n)$, $y_1 > x_1$ and hence $\mathbf{y} > \mathbf{x}$. This, however, contradicts the definition of the local satiation (if y_1 is close enough to x_1).

2.7 Can intransitive preferences be represented by a utility function?

No. According to the theorem proved on page 11, only rational preferences can be represented by a utility function. Intransitive preferences fail to be rational.

2.8 Utility function defined on page 11 is sometimes called "ordinal utility". Can it reflect the strength of preferences?

No. The definition of utility refers to one condition only: $u(\mathbf{x}) \leq u(\mathbf{y})$ if and only if $\mathbf{x} \leq \mathbf{y}$. The utility function u is called "ordinal" (indicating ordering only) if no additional assumptions are adopted. If $u(\mathbf{x}) = 2u(\mathbf{y})$ (the utility enjoyed by consuming $\mathbf{x}$ is twice as high as when consuming $\mathbf{y}$) one cannot say that "the bundle $\mathbf{x}$ is preferred twice as much as the bundle $\mathbf{y}$ "; one can only say that the bundle $\mathbf{x}$ is more preferred than the bundle $\mathbf{y}$. When some additional assumptions (not listed on page 11) are adopted about the function u (u becomes a "cardinal" utility), one can make statements not only about the ordering of bundles, but also about the strength of preferences. Adding utilities and multiplying them by some numbers becomes justified. Such "cardinal" utilities are used in so-called expected utility theory discussed in lecture 5 (QF-5).

3 – Choice under budgetary constraint)

Previous lecture was devoted to consumers' preferences, as if there was no money. In other words, preferences with respect to bundles were analysed irrespective of the cost necessary to buy them. Now we acknowledge the fact that goods included in bundles have some prices.

Consumers may have to sacrifice some goods if they want to enjoy some other goods. Several definitions have to be introduced in order to proceed.

Let $\mathbf{p}=(p_1,...,p_n)$ be the vector of prices of goods from the consumer's bundle and let m be the consumer's income (money to spent on the bundle). Then the consumer's budgetary constraint is:

- $p_1x_1+...+p_nx_n\leq m$,

and

- $p_1x_1+...+p_nx_n=m$

is called budget line. The set

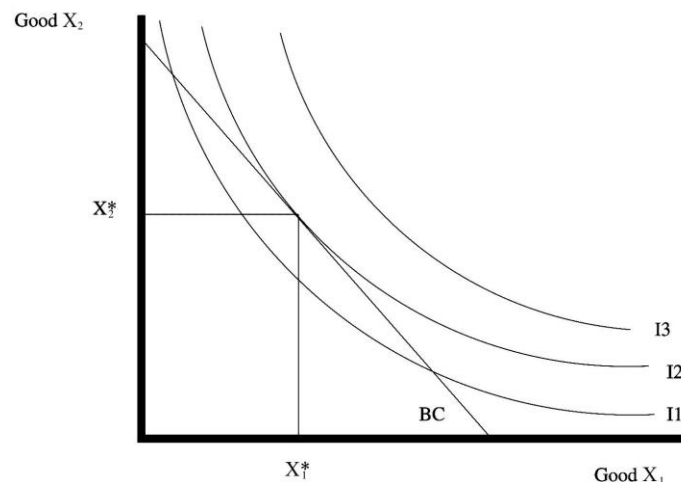
- $\{\mathbf{x}\in X: p_1x_1+...+p_nx_n\leq m\}$

is called budget set.

Now we can define the topic of the lecture. Consumer's choice under budgetary constraint is to:

- select a bundle in the budget set located at the highest indifference curve; or
- select a bundle in the budget set yielding the highest utility.

The following picture illustrates the choice in the case of bundles consisting of two goods ($n=2$).



Picture explains how a rationally behaving consumer chooses between two goods: X_1 and X_2 . It is assumed that preferences with respect to these goods are reflected by utilities, that is some real functions. It would be too complicated to draw everything in a two-dimensional plane, so utilities are represented by so-called "indifference curves", that is lines linking points that give the consumer the same utility (make the consumer indifferent with respect to such bundles). Three indifference curves – I_1 , I_2 , and I_3 – are drafted above. I_3 represents a higher utility than I_2 , and I_2 represents a higher utility than I_1 . Hence a rationally behaving consumer would prefer to consume any combination of goods 1 and 2 lying on I_3 rather than on I_2 , and anything lying on I_2 rather than on I_1 .

But the consumer cannot afford an arbitrary combination of goods. He (or she) is constrained by the amount of available money. If the prices of goods 1 and 2 are p_1 and p_2 , respectively, the bundle (x_1,x_2) costs $p_1x_1+p_2x_2$. This expenditure cannot be higher than the money available m , i.e. the budgetary constraint ($p_1x_1+p_2x_2\leq m$; BC in the picture above). Prices p_1

and p_2 determine the slope of BC. If the prices are equal to each other then BC has a (negative) slope of 45° . If p_1 is higher than p_2 then BC is steeper. If p_1 is lower than p_2 then BC is flatter. The proportion of prices determines the slope of BC, and the amount of money available – m – determines whether BC is farther or closer to the origin. If m is larger, then BC is farther from the origin which means that a consumer can afford more.

There are three indifference curves drawn in the picture: I_1 , I_2 , and I_3 . Bundles lying on I_3 are more preferred by the consumer than those lying on I_2 and I_1 . However he (or she) does not have money to afford anything that lays on I_3 . Only I_2 and I_1 are consistent with BC. Any bundle lying along I_1 is affordable, but if the consumer is rational, then he (or she) will prefer to buy a bundle lying on I_2 , since this is the highest indifference curve intersecting with BC. The combination (x_1^*, x_2^*) is the optimum choice. Students who recall the high school mathematics note that this optimum choice is where the BC is tangent to an indifference curve (as observed by one of the class participants today).

We come back to definitions of consumer's choice (select a bundle in the budget set located at the highest indifference curve; or select a bundle in the budget set yielding the highest utility). The first definition is a more general one, but both approaches are equivalent (if a utility function exists). The second problem

$$\max_{\mathbf{x}} \{u(\mathbf{x}): \mathbf{x} \in X \text{ and } p_1x_1 + \dots + p_nx_n \leq m\}$$

is called the Utility Maximization Problem (UMP). The following theorem states that UMP can be solved.

Theorem: If u is a continuous function, X is a closed set, and all prices are positive ($p_i > 0$ for $i=1, \dots, n$) then UMP has a solution $(x_1^*, \dots, x_n^*) = \mathbf{x}^*(\mathbf{p}, m)$.

Proof:

The Weierstrass theorem (a continuous function over a compact set achieves its maximum, and its minimum) can be applied in order to prove our theorem. The continuity of u is assumed. Compactness is therefore what needs to be checked. Any closed and constrained set is a compact set. Closedness is guaranteed by the budgetary constraint ($p_1x_1 + \dots + p_nx_n \leq m$); adding and multiplying are continuous operations. If prices are positive, then the set is constrained since for any good i , $x_i \leq m/p_i$ (the number at the right hand side is finite if $p_i > 0$).

A solution $\mathbf{x}^*(\mathbf{p}, m)$ to an UMP is called (primary) demand, also Marshallian or Walrasian demand; please note that optimisation outcomes do not have to be unique. It also defines so-called indirect utility. This is denoted by $v(\mathbf{p}, m)$. By definition, $v(\mathbf{p}, m) = u(\mathbf{x}^*(\mathbf{p}, m))$; if $\mathbf{x}^*(\mathbf{p}, m)$ is not unique then $u(\mathbf{x}^*(\mathbf{p}, m))$ is understood as the utility of any bundle which solves the UMP. Adjective 'primary' alludes to the fact that consumers are more important than firms. Also firms reveal demand which will be called 'secondary' (see lecture 6, QF-6).

As far as numbers go, indirect utility is the same as the utility defined earlier. The difference is in how it is computed. The earlier defined utility reads $u(\mathbf{x}^*)$. Indirect utility reads $v(\mathbf{p}, m)$. In other words, in order to calculate indirect utility, one needs to find $\mathbf{x}^*(\mathbf{p}, m)$, i.e. to solve an UMP, and then to calculate $u(\mathbf{x}^*(\mathbf{p}, m))$.

In advanced microeconomics textbooks the following Roy's identity is proved in order to find a relationship between indirect utility and Marshallian (Walrasian) demand:

$$x_i^*(\mathbf{p}, m) = -(\partial v(\mathbf{p}, m) / \partial p_i) : (\partial v(\mathbf{p}, m) / \partial m).$$

The identity can be proved for a continuous utility function representing locally non-satiated and strictly convex preferences. It is remarkable for practitioners, because it relates observable quantities ($x_i^*(\mathbf{p}, m)$) to unobservable ones ($v(\mathbf{p}, m)$), but it will not be used in this course.

Instead, we will analyse how to measure welfare changes. The most obvious way to proceed is to calculate so-called consumer surplus. It is defined in two versions: gross and net. Gross consumer surplus is the sum or integral of marginal benefits (measured by reservation prices, that is prices that people are willing to pay or willing to accept – not necessarily what they are exposed to in the market) from the bundle of the goods consumed. (Net) consumer surplus (CS) is the gross consumer surplus after subtracting the cost of purchasing the bundle. Calculating CS for welfare changes caused by price changes is difficult since – in general – changing prices imply changing consumer's relative wealth and thus his or her optimum choices.

More adequate measures of welfare changes require defining money metric utility. Let $n=2$, and let the good $i=2$ represent aggregate consumption of everything except for the good $i=1$. This aggregate consumption is measured in money (and thus $p_2=1$). Then x_2 in utility $u(x_1, x_2)$ can be interpreted as the income left for purchasing all goods except for x_1 (once the good $i=1$ has been purchased), i.e. $x_2 = m - p_1 x_1$. If we assume that $p_2=1$, then consumer's choice depends on the price of good $i=1$ only. The price of this good will be denoted by p (without a subscript).

Two definitions of welfare changes can be contemplated once money-metric utility is established: compensating variation, and equivalent variation. In general they can be different. We will start with the former.

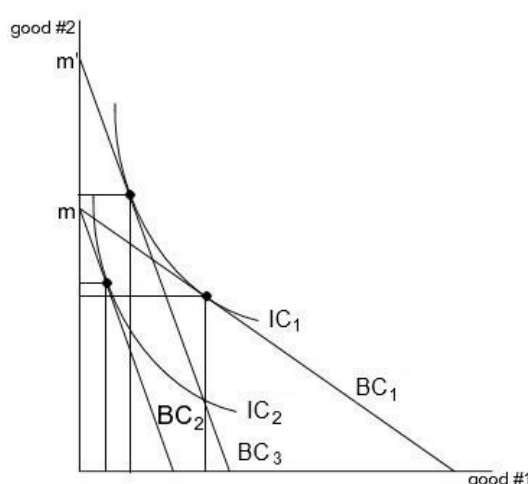
Compensating Variation (CV) implied by changing prices from $\mathbf{p}^0$ to $\mathbf{p}^1$ reflects the value of income change (decrease if $p^0 > p^1$ or increase if $p^0 < p^1$) necessary in order to keep the money metric utility at the original level despite the price change. CV is computed from the equation $CV = m - m'$, where m (the old original income) and m' (the hypothetical new income) satisfy the equation:

$$u(x_1(p^0, m), m - p^0 x_1(p^0, m)) = u(x_1(p^1, m'), m' - p^1 x_1(p^1, m')).$$

In other words, CV informs to what extent the price change modifies the *de facto* income ($-CV$ is the minimum amount the consumer should be paid in order to voluntarily accept such a price change; if $CV > 0$, the consumer should be willing to pay for the change).

The definition looks complicated but a picture (below) explains how m' is calculated (assuming that the price increased: $p^0 < p^1$). The original budget constraint was BC_1 (the total amount of money was m). It allowed the consumer to be at the indifference curve IC_1 , enjoying the utility $u(x_1(p^0, m), m - p^0 x_1(p^0, m))$. The price increase made the budget constraint steeper (BC_2). It allows the consumer to be at a lower indifference curve (IC_2). There is no doubt that his/her welfare decreased. Can we estimate by how much? This is the essence of the formula above and this is how the hypothetical amount of m' can be calculated to answer

the question. The new budget constraint is steeper which reflects the fact that the new price is higher and the amount of money the consumer has does not allow to enjoy the earlier consumption. Not necessarily the earlier consumption of both goods (their relative prices changed), but the earlier indifference curve can be achieved if the consumer was given more money. If you look at the picture, you will see that a hypothetical budget constraint BC_3 is parallel to BC_2 (thus it reflects the new price ratio), but it is pushed up in order to touch the original indifference curve IC_1 . The tangency point is different than the original one (the consumer is likely to choose a different combination of good #1 and good #2), but the level of well-being is as before. Thus the price increase was compensated by allowing the consumer to spend m' on the optimal bundle. The difference $m-m'$ – the CV – measures the amount of money the consumer should be paid in order to enjoy the earlier level of welfare.



Compensating variation

The next welfare change measure is called equivalent variation (EV). While CV looked at old utility and new prices, EV looks at new utility and old prices.

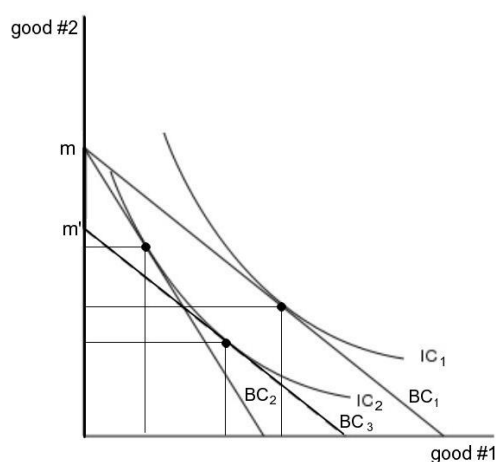
Equivalent Variation (EV) implied by changing prices from p^0 to p^1 reflects the value of income change (increase if $p^0 > p^1$ or decrease if $p^0 < p^1$) necessary in order to bring the money metric utility from the initial level to the final one without the price change. EV is computed from the equation $EV = m' - m$, where m (the old original income) and m' (the hypothetical new income) satisfy the equation:

$$u(x_1(p^1, m), m - p^1 x_1(p^1, m)) = u(x_1(p^0, m'), m' - p^0 x_1(p^0, m'))$$

In other words, EV informs to what extent the price change could be substituted by modifying the consumer's income ($-EV$ is the maximum amount the consumer would be willing to pay to avoid the price change; if $EV > 0$, the consumer should be paid to avoid this voluntarily).

As before, a picture lets explain how m' is calculated (assuming that the price increased: $p^0 < p^1$). The original budget constraint was BC_1 (the total amount of money was m). It allowed the consumer to be at the indifference curve IC_1 . The price increase made the budget constraint steeper (BC_2). It allows the consumer to be at a lower indifference curve (IC_2). There is no doubt that his/her welfare decreased. Can we estimate by how much? This is the essence of the formula above and this is how the hypothetical amount of m' can be calculated to answer the question. The new budget constraint is steeper which reflects the fact that the new

price is higher and the amount of money the consumer has does not allow to enjoy the earlier consumption. The change makes the consumer feel to have less money. Under the old prices it would feel like lowering the amount of money available. This hypothetical situation can be simulated by pushing the original budget constraint BC_1 down to BC_3 , i.e. until it touches the new indifference curve IC_2 . This hypothetical budget constraint corresponds to the amount of money m' . The difference $m'-m$ – the EV – is the amount of money the consumer is *de facto* deprived if the price increases.



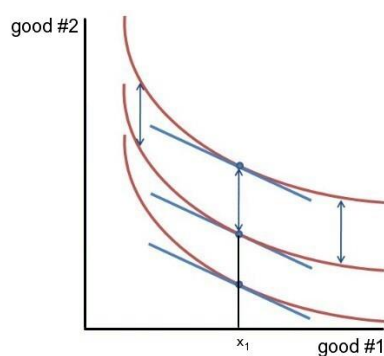
Equivalent variation

In general CV and EV can be different. There is an important special case when they are exactly the same. In order to see how this may happen, one more concept has to be introduced. A money-metric utility function u is called quasi-linear with respect to the good number 1, if there exists a function g such that $u(x_1, x_2) = g(x_1) + x_2$

Corollary

Isoquants of a quasi-linear utility function have identical shape and they are simply translations of any given one along the line $x_1 = \text{const.}$

The definition seems to be obscure, but the following picture clarifies it.



Quasi-linear preferences

The picture demonstrates three (red) indifference curves. Their shapes are identical. They are obtained by vertical transfers. Please note that the vertical distance between any two of them is exactly the same (compare blue arrows) no matter what x_1 is looked at; for a given x_1 they

differ only by the amount of x_2 (this is how a quasi-linear function was defined). Moreover, for a specific x_1 – no matter what x_2 is – the (blue) tangents are parallel to each other. This observation leads to two conclusions. First, the MRS_{12} depends on x_1 , not on x_2 . Second the Marshallian (Walrasian) demand – that is a solution to UMP given specific relative prices – is the same for the good #1. The amount of money available to the consumer has impact on the demand for the good #2 (the higher the amount of money the consumer has the higher the demand for this good, but not for the good #1).

Quasi-linear utilities allow economists to measure welfare changes in a non-ambiguous way, as the following theorem states.

Theorem

For a quasi-linear utility CV and EV are identical and they are equal to:

$$CV = EV = g(x_1(p^1, m)) - g(x_1(p^0, m)) - (p^1 x_1(p^1, m) - p^0 x_1(p^0, m))$$

Proof:

It results from definitions of CV and EV once utilities are substituted by indirect utilities.

Questions and answers to lecture 3

3.1 What is the slope of the budget line (see picture on page 15) if one of the goods does not cost anything (say, $p_1=0$)?

It is either horizontal (when $p_1=0$) or vertical (when $p_2=0$). If $p_1=0$ and $p_2>0$ then the only expenditure the consumer needs to make is to buy the second good. If he (or she) spends all the money m , then he (or she) can buy m/p_2 units of the second good. The horizontal line with such a coordinate is the budget line (any amount of the first good can be combined with m/p_2 units of the second good). If $p_1>0$ and $p_2=0$ it is the other way around. The only expenditure the consumer needs to make is to buy the first good. If he (or she) spends all the money m , then he (or she) can buy m/p_1 units of the first good. The vertical line with such a coordinate is the budget line (any amount of the second good can be combined with m/p_1 units of the first good). If both prices are positive then the budget line determines a triangle (like in the picture on page 15).

3.2 Let us assume that the consumer has a utility function $u(x_1, x_2) = 2x_1x_2^3$, the prices are $p_1=1$, $p_2=3$, and the money available is $m=12$. Please solve the UMP.

There are several ways to calculate the answer. First of all, let us observe that the utility function $u(x_1, x_2) = 2x_1x_2^3$ – called Cobb-Douglas function – allows internal solutions only (otherwise the utility is zero: $u(x_1, 0) = u(0, x_2) = 0$).

The budget line is $x_1 + 3x_2 = 12$. Assuming that the consumer spends all the money on buying both goods, $x_1 = 12 - 3x_2$. Substituting this to the utility function one gets $u(x_1, x_2) = 2(12 - 3x_2)x_2^3 = 24x_2^3 - 6x_2^4$. To maximise this function of one variable one needs to check where its derivative vanishes: $72x_2^2 - 24x_2^3 = 0$. One solution is $x_2=0$ (to be rejected), and the other one is $x_2=3$ which lets calculate $x_1 = 12 - 3x_2 = 3$. To make sure that this is maximum

indeed, one needs to check the second derivative, i.e. $144x_2 - 72x_2^2$. For $x_2=3$ this expression is negative (-216) which confirms the answer: $\mathbf{x}^*=(3,3)$.

Let us apply the Lagrange method. $L(x_1, x_2, \lambda) = 2x_1x_2^3 - \lambda(x_1 + 3x_2 - 12)$. The first order conditions read: $2x_2^3 - \lambda = 0$ and $6x_1x_2^2 - 3\lambda = 0$. Trying $\lambda=0$ implies $x_2=0$ which has to be rejected (a solution must be an internal one). Hence $\lambda > 0$ and $x_1 + 3x_2 = 12$. We have three equations with three unknowns. By substituting $x_1 = 12 - 3x_2$ (using the third equation) we have

- $2x_2^3 - \lambda = 0$, and
- $-18x_2^3 + 72x_2^2 - 3\lambda = 0$.

By multiplying the first equation by 9 and adding to the second one we get $72x_2^2 - 12\lambda = 0$. From this $\lambda = 6x_2^2$. Substituting it to the first equation we get $2x_2^3 - 6x_2^2 = 0$ which has two roots: $x_2=0$ (to be rejected) and $x_2=3$. Hence $\mathbf{x}^*=(3,3)$ as before. The Lagrange method is somewhat more complicated, but it lets calculate $\lambda^*=54$ which informs about the sensitivity of u function to relaxing the constraint (i.e. allowing the consumer to spend more than 12).

The last method uses the fact that UMP for a Cobb-Douglas utility function has a special property. Namely, if $u(x_1, x_2) = Ax_1^\alpha x_2^\beta$ then the proportion of expenditures on the first good and on the second good in a solution to the UMP is like $\alpha:\beta$. In our case $\alpha=1$ and $\beta=3$. Hence $p_1x_1/p_2x_2=1/3$, that is $x_1/(3x_2)=1/3$ or simply $x_1=x_2$. This combined with the budget line $x_1+3x_2=12$ yields the answer (3,3).

There are many ways to answer questions of that sort, but each has to be supported by an appropriate method.

3.3 What is the Marshallian (Walrasian) demand for the goods #1 and #2 revealed by the consumer from question 3.2?

The Marshallian (Walrasian) demand is the solution of the UMP. As calculated above, $\mathbf{x}^*=(3,3)$, and this is the demand revealed by this consumer. To be exact, one should write $\mathbf{x}^*(\mathbf{p}, m) = \mathbf{x}^*(1, 3, 12) = (3, 3)$. This demand is revealed when the price of the first good is 1, the price of the second good is 3, and when the amount of money the consumer has to spend on the bundle is 12. Changing any of these may imply a change in the demand.

3.4 A household of retirees used to spend its entire monthly income on buying bread ($p_1=4$, $x_1=25$) and cross-word puzzles ($p_2=2$, $x_2=40$). The prices changed, but as a result of the revalorisation, the monthly income of 180 increased to 199 (by more than 10%). The new prices are $p'_1=5$, and $p'_2=1.5$. The new quantities bought are $x'_1=18$ and $x'_2=66$. Is the household better off now?

The new bundle (18,72) was unaffordable earlier, because the cost of its purchase was higher than what was spent on the old one ($4 \times 18 + 2 \times 66 = 204 > 180 = 4 \times 25 + 2 \times 40$). At the same time the old bundle (25,40) is affordable now ($5 \times 25 + 1.5 \times 40 = 185 < 199 = 5 \times 18 + 1.5 \times 66$). Yet it was not chosen by the household. Assuming that the choice was an optimal one (solving the UMP), the household is not worse off (it is better off, or it enjoys the same utility as before).

3.5 For a quasi-linear utility, if the demand for the good number 1 is derived from the FOC, then it depends only on the price (not on the income). Please demonstrate that the assumption on FOC cannot be left out.

This question – as E-3 – is included as an open ended exercise to the overheads for this course. An example referred to in the exercise should demonstrate that if the demand for the good number 1 is not derived from the FOC, then it may depend on income. Indeed, for quasi-linear preferences, i.e. preferences represented by a utility function $u(x_1, x_2) = v(x_1) + x_2$ the demand for x_1 is derived by maximizing $v(x_1) + x_2$ subject to $px_1 + x_2 = m$. One can rephrase this as a one-dimensional problem by substituting $x_2 = m - px_1$. The problem is then to maximize the function $v(x_1) + m - px_1$. If the maximum is where the derivative vanishes, the FOC condition reads $MU_1 = p_1$ (MU stands for marginal utility), i.e. it is independent of m . However, the maximum of a one-dimensional function can be attained at the border of its domain (not where its derivative vanishes). In such a case, the demand may depend on m .

4 – Intertemporal choice

If I asked you to give me 100 euro today, and promised to give you 100 euro back one year from now, most of you would say "no". But if I ask you "why did you disagree", many students explain that they are afraid of inflation: 100 euro a year from now – because of inflation – is worth less. They are polite, so they do not raise another reason often, but there is also a risk that I will die within the next year, and consequently they will never get their money back. There is also yet another reason of uncertainty. What if I survive, but I turn out to be a thief who denies any promises to give the money back?

So let us assume that there is no risk. If you give me the money, you will get it back for sure. Also inflation can be excluded. If there is inflation, say, 2%, I will give you 102 euro so that the money you get has exactly the same value as when you gave it to me. Despite that, most of us would not agree to such a deal.

If I changed my question: "you give me 100 euro today, and I give you 130 euro next year (no inflation, no risk of losing the money); do you agree?" Some of you will probably say "yes". Let us then bargain. I ask: "you give me 100 euro today, and I give you 101 euro next year". You say "no". I say "you give me 100 euro today, and I give you 125 euro next year". You say "yes". I say "you give me 100 euro today, and I give you 102 euro next year". You say "no". And the negotiation process can go on. If there is an amount, say, 105 euro to be paid back after a year, and you say "I am indifferent – I can agree to such a deal, but I do not insist; I can do without such a deal either", then we say that 5% is the discount rate applied by the person who agreed.

The discount rates we apply can be very different. As a rule, they are positive (nobody has a zero or a negative discount rate), but they depend on many circumstances. If somebody has very attractive investment opportunities, then his (or her) discount rate is high. If somebody does not have such opportunities, then the discount rate is lower. Also, if the sacrifice one has to envisage in order to decrease the consumption today is large, then the discount rate is high. If this sacrifice is not large (for instance, if you come from a wealthy country) then the discount rate is low.

Discount rates are very often identified with bank interest rates. This is an oversimplification. One reason is that bank interest rate is typically combined with a number of financial services (such as e.g. credit cards), and therefore it is difficult to claim that you agree to, say, 3% because this is your discount rate; perhaps your discount rate is higher, but you agree to 3%

because of the other benefits the bank offers you. Or it can be the other way around. You take advantage of the 7% rate of interest not because your discount rate is that high. The bank is threatened by bankruptcy, you are aware of the risk, and you require that the bank offers you some extra premium for the risk.

Interest rate on Swiss government bonds is considered a useful proxy for a discount rate (at least for those who buy such bonds). They are considered safe (the Swiss government is unlikely to go bankrupt), and they do not offer any benefits (like credit cards or so). The interest rate on Swiss government bonds is around 3%, and this is sometimes considered a useful reference for what our discount rate can be.

The formula for finding equivalents of money amounts in different time moments can be iterated. In general, it is:

$$X_t = X_0(1+r)^t, \text{ or } X_0 = X_t/(1+r)^t,$$

where r is a discount rate. Please note that in my examples $t=1$ (the time span between "now" and the "future" is one year). The formula then reads:

$$X_1 = X_0(1+r)^1, \text{ or } X_0 = X_1/(1+r)^1,$$

For two years it will be

$$X_2 = X_0(1+r)^1(1+r)^1 = X_0(1+r)^2, \text{ or } X_0 = X_2/(1+r)^2,$$

and so on. The formula linking X_t with X_0 is called the "Present Value" (PV) of X_t in moment $t=0$. If it is applied to a series X_1, X_2, X_3 , and so on, the formula of PV reads:

$$NPV = X_0/(1+r)^0 + X_1/(1+r)^1 + X_2/(1+r)^2 + \dots + X_T/(1+r)^T,$$

where T is the last year that the decision (project) is expected to imply a cost or a benefit (costs are entered with "minus" signs, and benefits are entered with "plus" signs). The letter "N" stands for "net", and consequently the acronym NPV reads "Net Present Value". The word "net" refers to the basic philosophy of Cost-Benefit Analysis: we are interested in benefits net of costs (i.e. benefits minus costs). If benefits and costs belong to different time periods, then we need discounting to make them comparable. Namely, we recalculate everything in present value terms, as if everything was to happen in the beginning (at the time of the analysis), in $t=0$.

If $X_t=X=\text{const}$ – then the formula for the present value simplifies:

$$NPV = X/(1+r)^0 + X/(1+r)^1 + X/(1+r)^2 + \dots + X/(1+r)^T = X(1 + 1/(1+r) + \dots + 1/(1+r)^T)$$

It is good to know, that there is such a simplification, but we will not use this formula extensively (I used it when I calculated the tables you will see in a moment).

Formulae of how to calculate NPV are used to define the so-called Internal Rate of Return (IRR). Namely, IRR is a discount rate which makes the $NPV=0$. In other words, IRR is a discount rate r such that:

$$X_0/(1+r)^0 + X_1/(1+r)^1 + X_2/(1+r)^2 + \dots + X_T/(1+r)^T = 0$$

If $X_0, X_1, \dots, X_{T-1} < 0$, and $X_T, X_{T+1}, \dots, X_T > 0$, then IRR is the only r which solves the equation above (IRR is unique). For typical projects this condition is satisfied (i.e. one needs to bear some investment cost in the beginning, and then one benefits from it).

IRR is a useful indication of profitability (it answers whether the project makes economic sense):

- If IRR is higher than the interest rate available for investor (to borrow the money in order to finance it), the project is efficient (and it is worth financing)
- If IRR is lower than the interest rate available for investor (to borrow the money in order to finance it), the project is inefficient (and it should be abandoned)
- If IRR is negative then the project does not make sense irrespective of the terms of availability of the money (unless you do not have to pay for the credit; the bank pays you to borrow the money).

Projects with low IRR need to be scrutinized carefully. Even if you do not have to borrow money (if you can finance the project yourself), low IRR indicates that the advantage of benefits over costs is small. If there are compelling non-economic reasons to carry out the project, it may be worth promoting, but please be careful. Perhaps there are better ways to spend the money.

Now let us go back to discounting. It is fairly easy to convince people about the need of discounting. If there is no discounting (or if the discount rate is zero) then you should accept the deal of giving me 100 euro now, and I will give it back to you a year from now. Everybody understands that this would be ridiculous, and therefore some discounting is inevitable. Problems arise when we try to envisage long-term consequences of discounting.

The present value of the future amount of $X_T=1,000,000$

	$T=1$	$T=5$	$T=10$	$T=20$	$T=50$	$T=100$	$T=200$
$r=0\%$	1,000,000	1,000,000	1,000,000	1,000,000	1,000,000	1,000,000	1,000,000
$r=1\%$	990,099	951,466	905,287	819,544	608,039	369,711	136,686
$r=4\%$	961,538	821,927	675,564	456,387	140,713	19,800	392
$r=8\%$	925,926	680,583	463,193	214,548	21,321	455	0.21
$r=12\%$	892,857	567,427	321,973	103,667	3,460	12	<0.01

Please look at the table above. It gives you the present value of a high amount which is going to be offered to you next year, 5 years from now, 10, 20, 50, 100, and 200. The present value depends on the discount rate applied. If the discount rate is $r=0\%$ (no discounting), then the present value is always the same and it is equal to the nominal value of 1,000,000. If a modest discount rate is introduced (say, 4%), the present value is less than a half of the original amount in less than 20 years. If the discount rate is $r=12\%$, then the present value almost disappears in 100 years. This is something that most people are not ready to acknowledge. Yet it is an unavoidable consequence of applying discount rates.

The next table reveals another uneasy result of discounting. Let us assume that we are to receive 100 annually forever. The present value of 100 obtained now is obviously 100. However, we are going to get this amount over and over again. If no discounting is applied (or if $r=0\%$), then the value of this payment for 10 years is 1,000, 5,000 in 50 years, 10,000 in 100 years, and infinity if there is no finite time horizon. The situation changes when some discounting is introduced. The present value of such a constant amount increases with number of years we look forward to, but it is finite even if $T=\infty$. For example, if the discount rate is $r=4\%$, and we are interested in the present value of a gift of 100 per annum we are going to receive for 10 years, it is not 1,000; it is only 811 (mere 81% of the product 100×10).

The present value of the flow of a constant amount of $X=100$

	$T=10$	$T=50$	$T=100$	$T=\infty$
$r=0\%$	1,000	5,000	10,000	∞
$r=1\%$	947	3,920	6,303	10,000
$r=4\%$	811	2,148	2,451	2,500
$r=8\%$	671	1,223	1,249	1,250
$r=12\%$	565	830	833	833

The calculation of all columns can be carried out in a spreadsheet. Only the last column can be calculated in an easier way (e.g. on your calculator). To see this, please recall the formula from page 2:

$$NPV = X(1/(1+r) + \dots + 1/(1+r)^T)$$

This is a geometric series with the quotient $q = 1/(1+r)$. Please note that the first payment is to be received by the end of the first period; thus its PV is $X/(1+r)$. The formula for the sum of such a geometric series (I assume that you remember your high school mathematics) is

$$a_1(1-q^T)/(1-q),$$

where a_1 is the first term of the geometric series, and q is its quotient. If $r>0$ then the quotient q is less than 1. If you go with $T \rightarrow \infty$, then q^T goes to zero (hence the numerator goes to a_1), and the denominator is constant (it is equal to $1-q$). This lets us calculate the last column as $(X/(1+r))/(1-q)$, that is $(X(1+r))/(r(1+r))$ or simply x/r .

In financial mathematics continuous discounting is applied. Interest rate contracts inform sometimes depositors about adding interest to the asset semiannually, monthly or even daily. If interest is added twice a year rather than once, it means that $X_1 = X_0(1+r/2)^2$, rather than $X_1 = X_0(1+r)^1$. If interest is added m times a year, the formula reads $X_1 = X_0(1+r/m)^m$. In their elementary Calculus course students are encouraged to prove that $\lim_{m \rightarrow \infty} (1+r/m)^m = e^r$. Thus the 'continuous time' version of the formula $X_t = X_0(1+r)^t$ reads $X_t = X_0 e^{rt}$ (or, *vice versa*, $X_0 = X_t e^{-rt}$).

The following case study will help us to see how the concepts of discounting can be applied in practice. We will calculate the IRR of a wind-mill project.

Let us assume that the project is characterized by the following parameters. The wind mill of 1 MW capacity

- works 2000 hours per annum
- costs 2 million euro
- sells electricity at 50 euro/MWh
- requires no maintenance costs for 30 years.

The assumptions are not far from reality. A typical wind-mill can work for 30 years. It is very expensive to build and to connect to the network, but once this is done, it operates almost for free. Annual maintenance costs are 2% of the investment cost (40,000 euro in this example), and they can be simply neglected; wind-mills do not require any repairs usually. The wind mill works 2000 hours per annum (on average less than 6 hours per day). Of course it can work longer, but with the capacity lower than the nominal 1 MW. Thus, every year, it produces 2,000 MWh. It sells electricity at the price of 50 euro/MWh (5 eurocent/kWh). The price you pay for electricity is much higher, but wind-mill sells at wholesale (not retail) prices. What you pay includes the cost of distribution, taxes etc. (By the way, students are often unaware of the price paid for electricity, because it is not very high.) Hence the annual revenue of the wind-mill is 100,000 euro (if you do not like the maintenance cost of 40,000 euro to be considered negligible, you may assume that the wholesale price of electricity is higher by this amount, and then you can subtract the maintenance cost; this will not change the calculations).

In other words, the project costs 2,000,000 euro to be paid on January 1 of the first year, and then it brings 100,000 euro every December 31 for 30 years. Its present value depends on the discount rate applied. If we want to calculate IRR, we shall answer the question what discount rate makes NPV=0. The question can be rephrased: what is the highest discount rate such that this project makes sense?

$$\begin{aligned} \text{NPV} &= -2,000k/(1+r)^0 + 100k/(1+r)^1 + 100k/(1+r)^2 + \dots + 100k/(1+r)^{30} = \\ &= -2,000k + 100k(1/(1+r) + 1/(1+r)^2 + \dots + 1/(1+r)^{30}) = \\ &= -2,000k + 100k((1-q^{30})/(1-q)), \text{ where } q=1/(1+r), \text{ if } q \neq 1, \text{ i.e. } r \neq 0 \end{aligned}$$

(like in the language of computer scientists, here k is understood as one thousand). In order to solve the equation $\text{NPV}=0$, one needs to solve an algebraic equation of the 30th degree. We know how to solve 1st degree (linear) equations and 2nd degree (quadratic) equations. Some people know how to solve the 3rd or the 4th degree equations, but there are no formulae to solve 5th degree and higher degree equations. They can only be solved numerically. By trial and error, we can check when the NPV is positive, and when it is negative. Then we can expect that somewhere "in between" it is equal to zero.

If you have access to spreadsheets, then finding an IRR (finding r or q such that $\text{NPV}=0$) can be made automatically, not by trial and error. I did it, and I found in this example that $\text{NPV}=0$ if and only if $q=0.97$, i.e. $r=0.03$. Please note that r is not simply $1-0.97$. r needs to be calculated from the formula: $q=1/(1+r)$. Hence $r=1/q-1$. Incidentally, it is 0.03.

IRR=0.03 for this project can be interpreted in a straightforward way. If the discount rate is higher than 0.03, then the investment will never pay back. If the discount rate is lower than 0.03 then the investment makes sense. It can also be phrased in the following way. If the money is available to the investor at the rate lower than 3%, then investment makes sense. If the money is available to him (or her) at the rate higher than 3%, then investment should be abandoned, since it will never pay back. Finally, please note that if $r=0$, then the investment will pay back after 20 years (because $2,000,000 = 20 \cdot 100,000$).

Please look at the first row in the table on page 4: if you have a steady flow of revenues, and if the project works for a sufficiently long time period, then it will repay its investment cost sooner or later. Discounting demonstrates that this argument is wrong.

Discounting is inevitable. If you do not discount, then it is not possible to make any reasonable economic analysis. However, several decades of my discussions with non-economists (especially environmentalists) demonstrate that they do not accept discounting (or – what is logically equivalent – they are convinced that the discount rate should be zero). There are two case studies which illustrate this.

The first is a wind-mill story, or any analysis involving renewable resources: you make a one-time investment, and then you get benefits "for free", year by year forever. If you receive benefits forever, then it is just a matter of a sufficiently long time horizon to see that their sum can be made arbitrarily high; in particular it can be made higher than the initial investment cost. Therefore the argument is: build wind-mills at any cost – they will pay back after some time. But all of a sudden there are economists who say (as in the example above), that NPV depends on a discount rate. Wind-mills make economic sense only under specific assumptions.

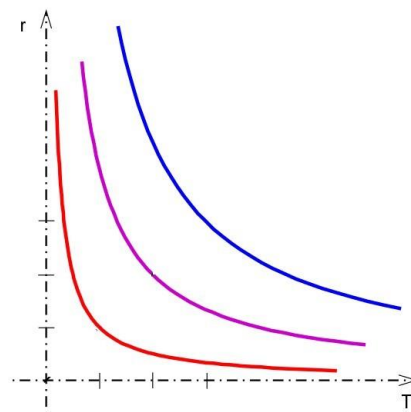
The second case is a nuclear controversy. Everybody knows that a nuclear power plant is excellent, and – as a rule – environmentally safe (Chernobyl and Fukushima are exceptions which do not contradict this general rule). There is only one weak point in this analysis. They require a huge decommissioning cost. There are no ideas what to do with scrapped reactors and radioactive waste. The cost of this can be quite high, perhaps the same as the cost of investment. Nevertheless it is to be borne after 60 years (since a typical nuclear power plant is likely to work for 60 years). The table on page 3 shows you, that if you apply a modest discount rate, say, 4%, the PV of this postponed cost is less than 14% of the nominal value. If you apply 8%, then the PV is less than 2%, i.e. almost nothing. The environmentalists say: nuclear power plants should not be built. If they are built, then this is because of discounting: only economists can ignore their huge and perhaps prohibitive decommissioning costs.

There is a pressure either to ignore discounting or to apply a much lower rate than implied by empirical research. The latter suggests that many of us apply a discount rate close to 4%. This is based on calculations, experiments, and observations on how we behave if we choose between now and the future. But the future we can capture is typically of one or several years ahead. It would be impossible to run an experiment extending over, say, a 50-year period. Therefore we do not know what are our preferences with respect to very long-term decisions. We simply extrapolate what we observe with respect to short-term behaviour to our hypothetical behaviour in the long run.

This is justified by so-called time consistency principle. This principle says that we can calculate NPV in two stages. Let us assume that we would like to calculate NPV of 1,000 to

come 10 years from now, if the discount rate is 4%. This is simply $1,000/1.04^{10} = 1,000/1.48 = 676$. In other words, the NPV of 1,000 in 2030 is 676 in 2020. We can arrive at this number by discounting it from 2030 to 2024 (six years), and then from 2024 to 2020 (four years). The "intermediate" value (NPV discounted from 2030 to 2024) will be $1,000/1.04^6 = 1,000/1.26 = 790$. When this amount is discounted from 2024 to 2020 we get $790/1.04^4 = 790/1.17 = 675$ (it is not 676 because of the rounding errors; if all calculations were carried out with higher accuracy, then the second number would be equal to the first one exactly).

The principle is based on a rule that everybody remembers from school: $x^{a+b} = x^a x^b$. If you want to calculate 2^5 (which is 32), you can calculate $2^3 = 8$, and $2^2 = 4$, and then multiply the two numbers. Until recently economists were convinced that the principle of time consistency is unquestionable. Hence, if you calculate NPV over a ten-year period you have to use the same discount rate when you discount over, say, 4 years, or 6 years. The first empirical doubts were voiced in the 1980s. Yet it was much later that economists provided a convincing explanation that – indeed – if you discount over a long period of time, you may apply a discount rate lower than appropriate for a shorter one. This idea is often called "hyperbolic discounting", as the following pictures explains.



The graph – displaying three examples of hyperboles $r=A/T$ (with three positive constants A) – illustrates the tendency that the longer the time horizon T , the lower the discount rate r . For instance, if we look at short time periods (say, one or three years) then the rate we discount with can be 4% or so. But if we look at a long time period (say, forty years) an appropriate discount rate can be much lower, perhaps 2% or so.

Discounting with variable discount rates violates the time consistency principle. For instance, if a long-term investment is planned (60 years in the case of a nuclear power plant) then a low discount rate may be appropriate. If a shorter-term investment is planned (30 years in the case of a wind-mill) then a higher discount rate may be appropriate. Calculations of NPV cannot be carried out in two stages like in the previous example (if the long period is divided into two shorter ones then discount rates applied do not have to be the same, and the formula we used above does not hold).

Questions and answers to lecture 4

4.1 Most of us think that if inflation rate is 3%, then the nominal discount rate of 8% corresponds to the real rate of 5%. This is not quite correct. Why?

Discount rate makes the two amounts – X_0 this year and $X_1 = X_0/(1+r_{\text{nominal}})$ next year – equivalent. But because of inflation the future amount is worth only $1/(1+\text{CPI})$ of its nominal value (where CPI – *Consumer Price Index* – is the inflation rate). We can thus write $X_1/(1+\text{CPI}) = X_0/(1+r_{\text{nominal}})$ or – equivalently – $X_1 = X_0(1+\text{CPI})/(1+r_{\text{nominal}})$, and we would like to compute r_{real} such that $X_1 = X_0/(1+r_{\text{real}})$. By combining both equations we get $X_0(1+\text{CPI})/(1+r_{\text{nominal}}) = X_0/(1+r_{\text{real}})$, that is $(1+\text{CPI})/(1+r_{\text{nominal}}) = 1/(1+r_{\text{real}})$, or $(1+r_{\text{nominal}}) = (1+\text{CPI})(1+r_{\text{real}})$. Hence $1+r_{\text{nominal}} = 1+\text{CPI}+r_{\text{real}}+\text{CPI}r_{\text{real}}$. The identity $r_{\text{nominal}} = \text{CPI}+r_{\text{real}}$ neglects the term $\text{CPI}r_{\text{real}}$. This term is small indeed if both CPI and r_{real} are few percent. In our example $\text{CPI}=3\%$, and $r_{\text{real}}=5\%$, so their product is equal to 0.0015, and can be neglected. Nevertheless, the general rule that in order to arrive at real values you simply subtract inflation rate from nominal values is not precise.

4.2 So-called Ramsey formula reads $r = \rho + \eta * g$, where r is the discount rate, ρ is the pure time preference, η is the elasticity of the marginal utility of money, and g is the growth rate. Analysed in the class discount rate r is higher than "pure time preference" ρ . Why?

Educated at the University of Cambridge, Frank Ramsey (1903-1930) made important contributions to mathematics and economics. What we call "the Ramsey formula" ($r = \rho + \eta * g$) is perhaps the best known of these. The formula was developed in order to solve an optimisation problem of how to split production between consumption (which pleases us immediately) and investment (which will satisfy our needs in the future). In trying to determine r in the class discussion today, we probably had in mind ρ , η , and g as well. Let us explain how these parameters may influence our choices.

The pure time preference ρ is perhaps the most obvious motive of our desire to have now rather than in the future. Yet we expect to be wealthier in the future. If the growth rate is, say, 4%, then we expect that instead of having, say, 10,000 € today we will have 10,400 € next year. Will this additional money make us better off? Surely it will: we will have 400 € more. But how much better off will we feel. Economists say, that the marginal utility of money (whatever it means) increases somewhat slower than the money. This means that $\eta < 1$. If $\eta = 1/2$, then these additional 4% correspond to only 2% as far as our wellbeing is concerned. Thus $r = \rho + 2$. Hence if we discovered that our discount rate was, say, 5%, this number consisted of the pure time preference of 3% and 2% corresponding to $\eta * g$ (if we expected that our wealth would grow by 4%, and the marginal utility of money was estimated at $1/2$). If we expect our income to grow, and additional money makes us happier, then $\eta * g > 0$. Consequently $r > \rho$.

4.3 IRR of the project #1 is lower than IRR of the project #2. Does this mean that #1 is worse than #2?

If the projects have the same duration, it does. But if the second project has a longer duration than the first one, then it does not have to be worse. In order to compare projects with different time spans economists developed the concept of so-called levelised costs. This term is especially popular among electricity specialists.

The *Levelised Cost of Electricity* (LCOE) tries to capture not only the distinction between investment and maintenance cost, but also tries to find a comparability between projects characterised by different time scales. The most relevant dilemma to be addressed in this

context is to compare windmills and nuclear power plants. The former are typically built for 25-30 years. The latter can work 50-60 years. As expected, discounting proves to have a crucial role in determining the adequacy of either technology, but – unfortunately – there are no convincing arguments to advocate for a discount rate of, say, 4% or 6%. Arguments for using variable rates (perhaps hyperbolic discounting for very long-term projects) seem to be appropriate, but this obscures the analysis even further.

The definition of LCOE is fairly straightforward:

$LCOE = (I_0 + \sum_i M_i/(1+r)^i)/(\sum_i E_i/(1+r)^i)$, the summation is for $i=1, \dots, n$,
where

- n is the lifetime of a project,
- I_0 is the investment cost,
- M_i is the maintenance cost in year i ,
- E_i is the electricity production in year i ,
- r is a discount rate

LCOE (e.g. measured in € per kWh) takes into account discounting (like IRR), but also the fact that various projects may have different time spans. No matter what their respective IRRs are, electricity generating projects can be compared by looking at their LCOE numbers.

4.4 As the winner of a lottery you have a choice of getting your prize, i.e. \$ 1,000,000, immediately. Alternatively you can wait one year to receive \$ 1,030,000. What will you choose?

The answer is simple. You check whether $1,030,000/(1+r)$ is greater than 1,000,000 or not. If your discount rate is lower than 3% you will wait. If your discount rate is higher than 3% you will claim your prize immediately. But please see also 4.5.

4.5 As the winner of a lottery you can get your \$ 1,000,000 immediately, or in two instalments: \$ 500,000 now, and \$ 500,000 a year from now. What will you choose?

Unless your discount rate is 0%, you should claim your prize immediately, because the present value of the second instalment is lower than its nominal value. Nevertheless, shrewd investors calculate their tax obligations too. It is possible that your income of \$ 1,000,000 is subject to a higher tax rate than \$ 500,000. If this happens then in order to choose between the two variants you should compare the hypothetical loss of decreased present value with the hypothetical gain from lower taxation.

4.6 Please refer to the table on page 24. The present value of 1,000,000 to be received 5 years from now is 821,927, and to be received 10 years from now is 675,564 when the 4% discount rate is applied. If the waiting time doubles from 50 years to 100 years, the present values shrink from 140,713 to 19,800. Why are these proportions different?

The proportions should be the same if the numbers change linearly, but they do not.
 $X_{10}/X_5 = (1+r)^5/(1+r)^{10} = (1+r)^{-5}$ and $X_{100}/X_{50} = (1+r)^{50}/(1+r)^{100} = (1+r)^{-50}$ The second number is smaller than the first one (unless $r=0$).

4.7 Please refer to the table on page 25. If the discount rate is 12%, then the present value of the annual payment of 100 over 50 years is 830 (not 5,000) It almost does not grow afterwards. Why?

The present value of 100 to be received in the 50th year is 0.35, that is a negligible amount. Later on it is even smaller (for instance, the present value of what you receive in the 51st year is 0.31). If you keep adding these smaller and smaller numbers the sum increases only slightly. The present value of 100 to be received in the 100th year is 0.0012, that is virtually zero.

4.8 The prince of Milan, Ludovico Sforza, gave a vineyard near the village of Fiesole, instead of paying Leonardo da Vinci what they agreed upon. How should this transaction be judged?

Leonardo as a painter was a genius, but he had problems with arithmetic. He probably agreed to this transaction without deep analyses. The prince of Milan would not be able to carry out necessary calculations either, but his officials were probably sufficiently competent to prepare the transaction. Its adequacy should be judged by looking at the table on page 4. If the vineyard gives an annual revenue of 100, the NPV of such a permanent flow depends on the time horizon, and on the discount rate. It can be assumed that a vineyard can live 50 years, and the discount rate is 1% (this is what it could have been in the 15th century Lombardy). Under these assumptions, the market value of such a vineyard was 3920. If this is what the prince owed to Leonardo then the transaction was adequate. If he owed more, then he cheated. If he owed less, then he turned out to be generous.

4.9 IRR=3% for a wind mill means that if money can be borrowed at 4%, then the investment is not economically justified. The availability of money is often even more difficult, and you see hundreds of wind mills being constructed around. How is it possible?

Please look at the assumptions of the class example. The wind mill is supposed to work 2000 hours per year, and the retail price of electricity is 50 €/MWh. These numbers can be changed. An average year has 8766 hours. It was assumed that a wind mill operates only 2000 hours, that is less than 25% of the time. There are maps of wind velocity. They show that the average wind velocity varies quite a lot. If the wind mill is located in a windy place, then it can operate more than 2000 hours per year. It was also assumed that the retail price of electricity is 50 €/MWh. But the government has instruments to make it higher. For instance, electricity from renewable sources does not need carbon dioxide emission permits. If other power plants have to buy such permits, the retail price of electricity may go up. It is a matter of environmental policy to regulate the market for carbon dioxide emission permits. If the number of operation hours, and if the retail price of electricity are higher than assumed in our calculations, then the IRR of a wind mill may exceed the rate the money is available at.

4.10 Why did economists consider the time consistency principle as an unquestionable thing?

We are so used to the rule " $x^{a+b}=x^a x^b$ " that the time consistency principle seems obvious. It allows for easy calculations of NPV. It was unimaginable that an analysis can be performed for variable discount rates; using variable discount rates would not allow to combine results carried out in shorter time intervals. Yet empirical observations – suggesting that people use higher discount rates for shorter periods – cast some doubts in the 1980s. Nevertheless, the

principle was rejected fairly recently, when a reasonable explanation was developed whether a project can be divided into smaller segments. By the way, whether projects are "divisible" or not depends on analyses of how a newly acquired information (as a result of technological progress) can be used. For instance, a nuclear power plant will use an old technology for 60 years even if a new one became available after 30 years of its operations. But this is just a curio for people who are interested in more sophisticated principles of economic analysis.

4.11 Why does the World Bank recommend using the discount rate of 8%?

The discount rate of 8% seems to be high – higher than many of us consider reasonable. Nevertheless, as one of the students observed correctly, the World Bank operates mainly in low income countries. If citizens of such countries are to make an investment in order to enjoy higher consumption in the future, they need to be convinced that their sacrifice is very rewarding. Otherwise, they would hesitate to lower their present consumption in order to gain something in the future. This dilemma is not that acute in high income countries; if citizens are fairly well off and they have their basic needs met, they are not afraid of setting aside funds to enjoy something in the future. Thus in a high income country even a low (but positive) IRR may be tolerated. In a low income country – where money is less abundant – any investment (which subtracts from the present consumption) has to be justified more convincingly.

5 – Uncertainty

Uncertainty is something that adds to the complexity of economic decisions. No matter what we do, we can never be certain about the result. For instance, we may protect human health against certain injury, and in fact we expose people to an even greater damage.

We often refer to "uncertainty" and "risk" as synonyms, but in 1921, Frank Knight made an important distinction between the two. The former is a broader concept, and its meaning is consistent with what we attach to the word in our everyday language. In contrast, the latter is a special case of uncertainty, when we do not know what will happen, but we can estimate probabilities of alternative outcomes. Unlike uncertainty in general, risk can be insured. For instance, if someone's real estate is located by the river, the risk of flooding can be calculated by hydrologists. If it is estimated at, say 2%, and the owner would like to get a compensation of 10,000 € if flood occurs, then he (or she) can buy insurance by paying $10,000 \times 2/100$ €, that is 200 €. Nobody knows when the flood comes, but if it comes, the owner of the real estate can be paid what his (or her) insurance policy states.

Not all uncertainty is insurable. For instance, an entrepreneur who invests in a project cannot be sure whether the investment will pay back. Nevertheless it is impossible to buy insurance against the failure, because its probability cannot be calculated in a meaningful way. There are a number of uncertain situations that do not allow to buy effective insurance. Some economists claim that this is the essence of business activity, and the source of entrepreneurial profits.

The most important extension of the microeconomic consumer choice model outlined in previous lectures is to interpret bundles as lotteries. It is easy to make this step if you imagine that whatever you do does not have a fully predictable outcome. If you buy a bottle of milk,

you expect to enjoy a nice drink, but the milk can turn out to be bad, and you are surprised unpleasantly. If you decide to walk through a park in order to enjoy a nice weather, you can be attacked by a thief or by a drug addict and deprived of your wallet. If you kiss a frog, the frog may turn out to be a beautiful princess who wishes to marry you (perhaps in this case boys can be surprised pleasantly). These observations can be formalised using some mathematical definitions and theorems. Some of them are easy to prove (you are encouraged to check it). If they are very difficult, they will not be proved here; students are encouraged to look into more advanced textbooks.

Elements of the set X can be interpreted as lotteries $L=(p_1,...,p_N)$, where $p_1+...+p_N=1$ (non-negative $p_1,...,p_N$ are interpreted as probabilities); $\mathcal{L}$ is the set of such lotteries; their outcomes – numbered $1,...,N$ – are predetermined.

Convexity is a key assumption referred to in economic analyses. Your high school mathematics used this concept as well. In economics jargon a convex combination means simply that you take a sum of something with weights that add up to 1. Formally, a convex combination of $y_1,...,y_k$ is $\alpha_1y_1+...+\alpha_ky_k$, where $\alpha_1+...+\alpha_k=1$.

I try to be fair with respects to students, but I heard that some professors treat exams as lotteries: for instance, they toss a coin or throw a dice and based on the outcome they either pass a student or fail. Let us assume that there are two lotteries: $L=(1/6, 5/6)$, and $L'=(1/2,1/2)$. The first outcome is "fail", and the second outcome is "pass". L can be interpreted as throwing dice: if it is "1" you fail, and if it is "2", or "3", or "4", or "5" or "6" you pass. L' can be interpreted as tossing a coin: if it is "head" you fail, and if it is "tail" you pass. Of course, students prefer L over L' . How about a convex combination of L and L' with weights $\alpha_1=3/4$ and $\alpha_2=1/4$? In this combined lottery the probability of failure is $1/4$ (because $3/4 \times 1/6 + 1/4 \times 1/2 = 1/4$), and the probability of passing is $3/4$ (because $3/4 \times 5/6 + 1/4 \times 1/2 = 3/4$). This combination is better than L' but worse than L for a typical student.

The following theorem formalises the concept of convex combinations of lotteries. It is extremely easy to prove.

Theorem A convex combination of lotteries is also a lottery (with probabilities calculated as convex combinations of probabilities from original lotteries).

We are now ready to introduce the main concept applied in consumer choice theory with respect to lotteries.

The von Neumann-Morgenstern (vNM) expected utility function, U has the "expected utility form", when

$$\exists u_1,...,u_N \in \mathfrak{R} \quad \forall L=(p_1,...,p_N) \in \mathcal{L} \quad [U(L) = u_1p_1+...+u_Np_N]$$

It is assumed in this definition that a consumer can measure his (or her) preferences with respect to a lottery in the following way. Every outcome of the lottery (numbered $1, 2, ..., N$) is characterised by the utility $u_1, u_2, ..., u_N$. The utility of the lottery (defined by probabilities $p_1,p_2,...,p_N$ that these outcomes happen) is then the average (expected) value of these numbers.

We usually think of lotteries as something uncertain. How about a lottery $L=(1, ..., 0)$? The first outcome happens for sure, and the other ones do not happen at all. Is this a lottery?

Despite doubts we may have (as some students raised), this is a lottery from a formal point of view since the requirement that probabilities are less than one or equal to one is satisfied. But if the probability of something is one, then we know that this will happen for sure. Mathematicians call such lotteries "degenerated" ones.

Numbers u_i from the previous definition – interpreted as utilities of "degenerated lotteries" $L^1=(1,0,...,0),...,L^N=(0,...,0,1)$ – are called Bernoulli utilities.

The following theorem says that the "expected utility form" is satisfied when convex combinations of lotteries are taken into account.

Theorem. A utility function $U:L \rightarrow \mathfrak{R}$ has the "expected utility form" if and only if

$$\forall K=1,2,... \forall L_1,...,L_K \in L \forall \alpha_1,...,\alpha_K > 0 [\alpha_1+...+\alpha_K=1 \Rightarrow U(\alpha_1 L_1+...+\alpha_K L_K) = \alpha_1 U(L_1)+...+\alpha_K U(L_K)]$$

Proof

The theorem has the "if and only if" (equivalence) form, and it will be proved separately for "only if" ($\Leftarrow$) and "if" ($\Rightarrow$) parts.

$\Leftarrow$

Let $L=(p_1,...,p_N)$. We define degenerated lotteries $L^1,...,L^N$ such that $L^i=(0,...,0,1,0,...,0)$; the i th probability is equal to 1. Then $L=p_1 L^1+...+p_N L^N$ and $U(L) = U(p_1 L^1+...+p_N L^N) = p_1 U(L^1)+...+p_N U(L^N) = p_1 u_1+...+p_N u_N$, where the second to the last equality holds by the assumption.

$\Rightarrow$

Let us consider a convex combination of lotteries $L_1,...,L_k$; with weights $\alpha_1,...,\alpha_k$, where $L_k=(p_1^k,...,p_N^k)$. Let $L'=\alpha_1 L_1+...+\alpha_k L_k$. Hence it can be calculated that: $U(L') = U(\alpha_1 L_1+...+\alpha_k L_k) = \alpha_1(u_1 p_1^1+...+u_N p_N^1)+...+\alpha_k(u_1 p_1^k+...+u_N p_N^k) = \alpha_1 U(L_1)+...+\alpha_k U(L_k)$. The first equality results from the definition of L' . The second – from the assumption. The last one – from the definitions of $L_1,...,L_k$.

From now on it will be assumed that all lotteries have monetary payoffs. A formal definition is obvious and it reads:

The lottery has monetary payments $x_1,...,x_N$, and Bernoulli utilities are a function $u: \mathfrak{R} \rightarrow \mathfrak{R}$ of these payments: $u(x_1),...,u(x_N)$.

The von Neumann-Morgenstern theory of expected utility (vNM) has been convincing to a lot of people, and from the very beginning (i.e. from the middle of the 20th century) it has become a standard reference for decision making in the presence of uncertainty. Nevertheless, from the very beginning too, several paradoxes were demonstrated in order to reveal that it does not explain everything; it fails to explain how people choose in some circumstances.

The first warning was formulated by Maurice Allais in 1953. He analysed lotteries with the following (monetary) outcomes:

- $x_1=0$,
- $x_2=1$ million USD, and

- $x_3=5$ million USD

It is important that the first outcome is zero, and the second and third ones represent large amounts of money. He ran two experiments. In the first one he asked people to choose between two lotteries defined as $L_1=(0,1,0)$ and $L_2=(0.01,0.89,0.1)$. Most people preferred L_1 . Then he asked the same people to choose between two other lotteries: $L_3=(0.89,0.11,0)$ and $L_4=(0.9,0,0.1)$. Most people preferred L_4 . Allais' experiments were replicated many times over the last 67 years – always with the same conclusion: if confronted with L_1 and L_2 people choose L_1 , and if confronted with L_3 and L_4 they choose L_4 .

These results are hardly surprising. In the first lottery (which happens to be a degenerated one) people are offered 1 million USD for sure. The second one gives a higher reward on average; 1.39 million USD rather than 1 million USD. But there is a small probability (just 1%) to get nothing. This small risk of getting nothing *vis-à-vis* getting 1 million USD for sure makes the first lottery more attractive. Also the second experiment has an easy intuitive interpretation. The lottery L_4 differs from L_3 by two things: the zero amount is to be received either with 89% probability (in L_3), or 90% probability (in L_4). These numbers are almost the same; it is very difficult to see the difference between 89% and 90%. Another difference between L_3 and L_4 is that the former attaches zero probability to the amount of 5 million USD, while the latter attaches zero probability to 1 million USD. In other words, in L_4 one can win 5 million USD, and in L_3 one can win 1 million USD; probabilities are almost the same (10%, and 11% respectively). Hence it is not surprising that almost everybody prefers L_4 over L_3 . The paradox discovered by Allais is that people who make such choices do not comply with vNM theory, as the following theorem states.

Theorem People who choose L_1 in the first experiment and L_4 in the second one do not comply with the vNM theory.

Proof:

If the vNM theory was followed, then some Bernoulli's utilities would have been applied:

- $u(x_1)=u_1$,
- $u(x_2)=u_2$, and
- $u(x_3)=u_3$.

Please note that we do not know these numbers, and the proof works irrespective of what they are. The outcome of the first experiment implies that:

$$u_2 > 0.01u_1 + 0.89u_2 + 0.1u_3.$$

The outcome of the second one implies that:

$$0.9u_1 + 0.1u_3 > 0.89u_1 + 0.11u_2.$$

These two inequalities contradict each other, since the second one can be rewritten as:

$$0.01u_1 + 0.1u_3 > 0.11u_2, \text{ and consequently}$$

$$0.01u_1 + 0.1u_3 > u_2 - 0.89u_2 \text{ (because } 0.11u_2 = u_2 - 0.89u_2), \text{ and finally}$$

$$0.01u_1 + 0.89u_2 + 0.1u_3 > u_2 \text{ (if } -0.89u_2 \text{ is moved to the left),}$$

which contradicts the first one (characterising the first experiment).

The overall conclusion is that expected utility theory may be insufficient to model people's behaviour in some circumstances.

One more concept will be introduced in this lecture. By referring to people's attitudes towards lotteries, we can define risk aversion and risk neutrality. People who are not risk averse are called 'risk loving' or 'risk prone'.

Our attitudes towards lotteries are affected not only by expected utilities, but also by a specific distribution of various outcomes. Let me define the following lottery. You win 10 € with the probability 50%, or you lose 8 € with the probability of 50%. On average you win 1 €, so some of us will probably agree to take part in such a lottery. Yet some of us will not. Even though the expected payoff of 1 € provides an incentive to take part, there is a probability of 50%, that a participant loses and has to pay 8 €. If a prospective loss of 8 € is more meaningful than a prospective gain of 10 €, then the lottery is not attractive.

Let me define another lottery. You win 8 € with the probability of 50%, or you lose 10 € with the probability of 50%. On average you lose 1 €, so probably many of us will refuse to take part in such a lottery. Yet for some of us the lottery can be attractive. If someone considers winning 8 € more meaningful than losing 10 €, the lottery can be attractive.

Finally, let me define yet another lottery. You either win 10 € or lose 10 € with the same probability of 50%. On average you receive zero, that is you neither win nor lose. Some people will be indifferent with respect to participating in such a lottery: they can either take part or not, and this makes no difference for them. Those who refused to take part in first lottery will probably refuse to take part in this one too. Those who wanted to take part in the second lottery will probably want to take part in this one too. Our attitudes depend on how we look at risk. The following definition formalises this idea.

Risk aversion and risk neutrality implied by Bernoulli utilities

Aversion: $\forall L=(p_1,...,p_N) \in L [u(x_1)p_1+...+u(x_N)p_N \leq u(x_1p_1+...+x_Np_N)]$

Neutrality: $\forall L=(p_1,...,p_N) \in L [u(x_1)p_1+...+u(x_N)p_N = u(x_1p_1+...+x_Np_N)]$

The definition says that we need to compare two numbers: $u(x_1)p_1+...+u(x_N)p_N$, and $u(x_1p_1+...+x_Np_N)$. They look almost the same. The first one calculates the utility of every outcome separately ($u(x_i)$), and then takes the average of them. The second one calculates the average outcome $x_1p_1+...+x_Np_N$ and then takes the utility of it. The difference is thus what comes first, and what comes after it; in both definitions we take the utilities and then their average on left, and we take the average, and then calculate its utility on the right. For risk neutral people the sequence does not make any difference. For risk averse people, the average utility of outcomes can be lower than the utility of the average outcome. People for whom the inequality reads $u(x_1)p_1+...+u(x_N)p_N \geq u(x_1p_1+...+x_Np_N)$ are called risk loving or risk prone.

This terminology allows us to characterise choices made with respect to the third lottery (with outcomes -10€ and +10€). Risk neutral consumers are indifferent (they can take part or not). Those who are risk averse will avoid such a lottery, and those who are risk loving may consider this lottery as attractive. Students who answered my questions regarding lotteries demonstrated that this lottery could be accepted by some of them.

The idea of attitudes towards risk can be also explained by the following definition.

Certainty equivalent of the lottery $L=(p_1,...,p_N)$ is a number $c(L,u)$ such that

$$u(c(L,u))=u(x_1)p_1+...+u(x_N)p_N$$

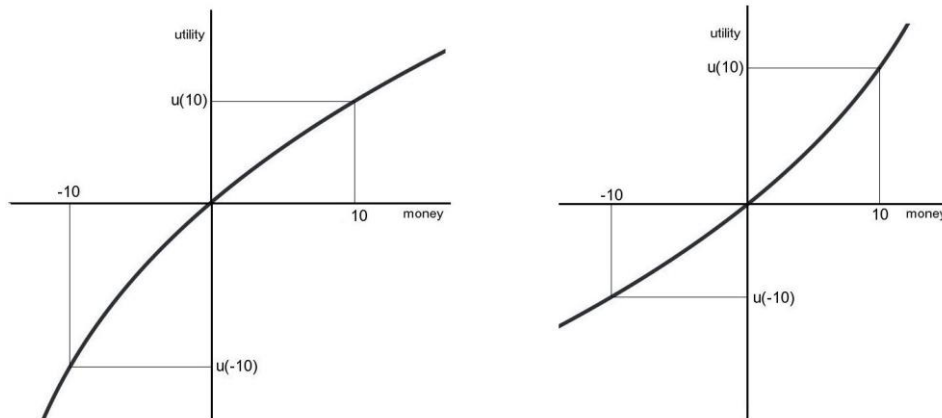
Once again let us look at examples of lotteries we referred to before. The first one was defined in the following way. You win 10 € with the probability 50%, or you lose 8 € with the probability of 50% (on average you win 1 €). Let us assume that prospective participants are offered cash which makes them indifferent with respect to participating or not participating in such a lottery. This is certainty equivalent. For a risk neutral person this will be the amount of money giving him (or her) the same utility as 1 € (the average outcome of the lottery). For a risk averse person (who did not want to participate in this lottery) the certainty equivalent is apparently less than that (less than 1 €). For a risk loving person it is the other way around: more than 1 €.

Picture below illustrates how a risk averse and a risk prone person feels. The graph on the left is characteristic for a risk averse consumer. The utility of loss of -10 is more meaningful than the utility of gain of +10:

$$|u(-10)| > u(10).$$

The graph on the right is characteristic for a risk prone consumer. The utility of loss of -10 is less meaningful than the utility of gain of +10

$$|u(-10)| < u(10).$$



We find very often that consumers are risk averse. This is how many people feel. Such people should not agree to take part in lotteries which – on average – make them lose more than gain. Yet this is not how they choose sometimes. In many countries popular lotteries announce that what they pay in prizes is less than what they collect for tickets. In other words, the average gain is negative (people pay more than what they receive in prizes). It can be argued that those who participate are not rational, and they do not understand probabilities and expected values. Of course some of them may behave irrationally. Nevertheless some of them are rational and consider themselves risk averse. Once again, economists have to admit that the vNM theory does not explain fully how we behave.

Picture above can be summarised in the following theorem (the equivalence of 1 and 2 is easy to prove):

Theorem The following conditions are equivalent:

1. A consumer is risk averse
2. Function u is concave

$$3. \forall L=(p_1,...,p_N) \in L [c(L,u) \leq x_1p_1+...+x_Np_N]$$

Questions and answers to lecture 5

5.1 Is there a possibility to buy an insurance against failing on an exam?

No. Some professors pass all the students, and some have a "failure rate" of 30% or so. These numbers can be fairly stable, so some students may claim that there is, say, 10% probability that they fail and therefore they may think of buying an insurance against failure. For example, if they value their "failure damage" at 500 €, they could expect to buy an appropriate insurance for 50 € (500 € is paid when they demonstrate the failing grade). I am afraid, however, that no insurance company would be willing to offer such a policy. The reason is that – unlike flood – failing on the exam is not a random process. Even though some probabilities can be attached, it is not random, because – to a large extent – students' behaviour may influence the result. If they study hard, the failure is less likely. If they do not, they are more likely to fail. An insurance company will not offer a policy if the mechanism of a damage is not random.

By the way, this is the reason why businessmen cannot buy insurance against bankruptcy. Even though statistics demonstrate that every year, say, 5% of registered companies turn out to be insolvent, the process is not random. Likewise, nobody can buy insurance against pregnancy, even though doctors say that the probability of success in the human fertilisation process is much less than 100%. Students may identify dozens of processes where probabilities exist, but there is no randomness.

5.2 Can measures aimed at risk reduction result in greater losses?

Yes, they can. In the 1930s and 1940s the American government spent a lot of money invested in flood protection infrastructure (walls, retention reservoirs, etc.). At the same time much higher flood damages were recorded. Some analysts were surprised, but the mechanism was fairly simple and predictable. Motivated by improved protection, investors were more eager to build homes, stores, and manufacturing plants close to the rivers. Irrespective of protection measures, floods occur sometimes anyway. Once they come, they imply higher damages if there are more assets to be hurt.

When undertaking anti-risk measures, one needs to expect that these may induce more risky behaviour. For example, some road safety measures do not lead to lower accident rates if they encourage people to drive less carefully.

5.3 Flyers attached to medical drugs list unwanted effects and their probabilities. One of the unwanted effects is the death of a patient. When we read that a death can be caused by using the drug, shall we reject the prescription?

The probability that a drug used for medical purposes can cause a patient's death is rather small. Nevertheless, it can be positive. Drugs known to cause patients' death (with very low probability) are administered usually when the alternative (i.e. no drug) leads to a higher probability of death. If I was prescribed a drug leading to unpleasant unwanted effects, I

would check if my chances of survival without the drug – according to what doctors say – are high enough. The prescription of the drug should not be rejected if unwanted effects show up with sufficiently low probabilities.

5.4 In microeconomic theory consumers may be characterised by different attitudes towards risk (for instance they can be risk averse, or risk prone). Yet it is assumed that firms are risk neutral. Why?

According to economic theories, consumers maximise utility while firms maximise profits. Of course, we can provide examples that sometimes they do not, but – as a rule – we assume an optimising behaviour, because otherwise little can be predicted. Firms' profits are denominated in money. Thus if a firm makes twice as high a profit, it doubled its objective. At the same time, if a consumer gets twice as much money, his (or her) utility goes up, although not necessarily in the same proportion (see graphs on page 37). The fact that firms measure things in monetary units, and consumers in utility units results in firms' risk neutrality; the loss and gain have the same meaning for a firm, while not necessarily they mean the same for a consumer.

5.5 Can various consumers have different certainty equivalents of the same lottery?

Yes. You can refer to the example used in the class (the lottery giving -10 or 10 with same probability of 50%). For a risk neutral consumer its certainty equivalent will be the amount of money giving the same utility as 0. For a risk averse person it will be less, and for a risk prone person – more. Please note that we look at the same lottery in each case.

5.6 Has the utility function of a given consumer to be always convex or always concave?

No. It may have a not constant curvature. In fact, one of the explanations of the Allais paradox is that consumers who choose among lotteries in a different way than vNM theory predicts, may have a different attitude towards risk when they lose and different when they win.

5.7 There is a lottery $L=(1/3,1/3,1/3)$ with three outcomes: -9 €, 0 €, +12 €. What is a fair price of this lottery ticket?

Fairness can be understood differently by different people. In microeconomics the fair price of a lottery ticket, $t(L)$, is equal to the expected payoff, i.e. if $L=(p_1,...,p_N)$, and the corresponding payoffs are $x_1,...,x_N$, then $t(L)=x_1p_1+...+x_Np_N$. By the way, for a risk neutral consumer $c(L,u)=t(L)$. Please note that the definition of a fair price of a lottery ticket depends on probabilities, not on how consumers perceive risks (in contrast, the certainty equivalent depends not only on what the lottery offers, but also on how a consumer perceives its risk). The expected payoff of this lottery is 1 €. Hence $t(L)=1$ €.

5.8 The price of the ticket to the lottery from 5.7 is 2 €. Who is likely to participate in it?

For a risk neutral consumer, the acceptable price of the ticket is 1 €. Thus a risk neutral consumer is not likely to participate in such a lottery. Risk averse consumers are even less likely to participate in it. The only consumers who may buy such a lottery ticket are those risk prone. But they need to love risk sufficiently strongly. Their certainty equivalent must satisfy

$c(L,u) \geq x_1p_1 + \dots + x_Np_N = -1$ (please note that the outcomes of the lottery are -11,-2,+10 rather than -9, 0,12, because from every prize one needs to subtract the price of the ticket, i.e. 2).

6 – General Equilibrium

In economics we distinguish between partial equilibrium, where demand and supply equate in a given (single) market, and general equilibrium, where demand and supply equate in all markets (simultaneously). The latter is more complicated, but it is more realistic in a sense that markets are linked to each other (if we spend money on something we have less to be spent on something else). Thus general equilibrium allows a deeper insight into what happens in an economy, although the notation is not that simple.

- k – the number of consumers
- n – the number of markets (products); one of the commodities can be labour – i.e. a resource owned by every consumer
- r – the number of firms
- Numbering of consumers: $i=1,\dots,k$
- Numbering of markets (products): $j=1,\dots,n$
- Numbering of firms: $h=1,\dots,r$

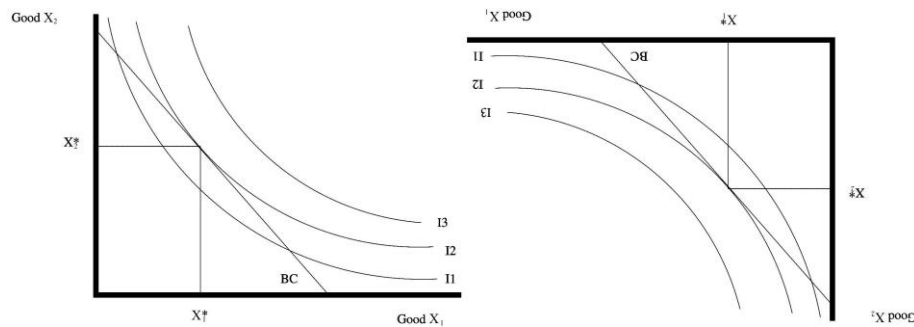
The numbers can be rather large. Even for a medium size economy (like in Poland) k is many millions (there are thirty-plus millions citizens, and almost twenty millions households). Likewise n can count in millions, unless we consider food a single market (product) rather than milk, beef, apples, etc. separately. Also the number of firms is very high. But in some analyses we make a simplifying assumption that economic agents do not produce anything; they merely exchange what is available. In this lecture we will concentrate on how prices are formed if the demand and supply are confronted in a static economy (that is without making decisions how much to invest).

A so-called pure exchange economy is an example of this simplifying assumption. In what follows we will confine to the case where two consumers hold two goods and contemplate whether to exchange some units of one good for some units of the other one (hence $r=0$, and $k=n=2$). They consider only allocations which do not make them worse off. The following terms will be used (the first subscript denotes the number of the consumer, and the second – the number of a good):

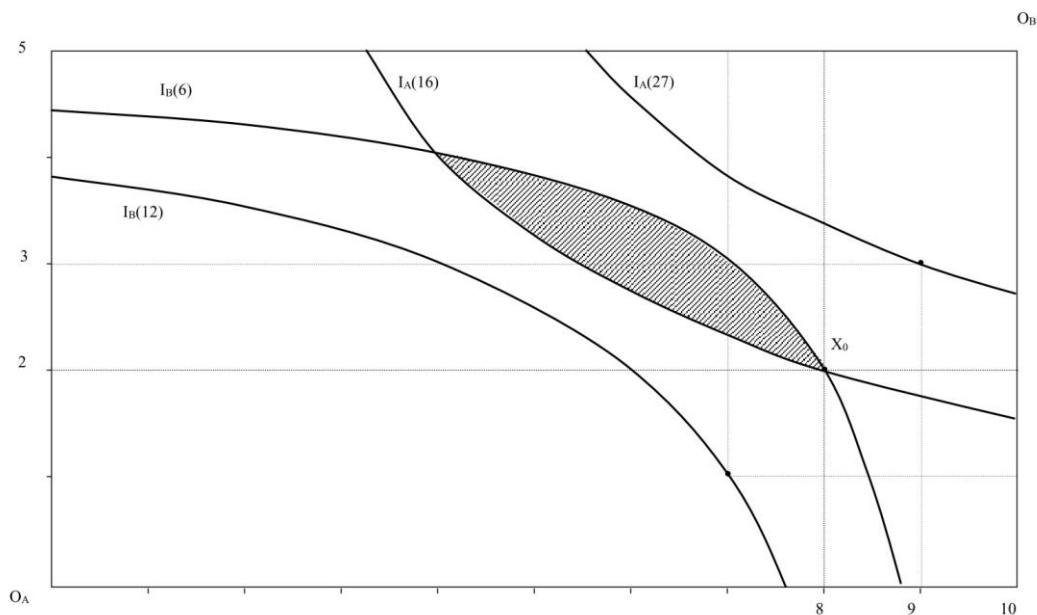
- Endowment (initial allocation) of the i th consumer: ω_{i1}, ω_{i2}
- Total endowment of the j th commodity: $\omega_j = \omega_{1j} + \omega_{2j}$
- Gross demand (final allocation) of the i th consumer: x_{i1}, x_{i2}
- Total demand for the j th commodity: $x_j = x_{1j} + x_{2j}$
- Excess demand of the i th consumer: $x_{i1} - \omega_{i1}, x_{i2} - \omega_{i2}$

In order to get insights into relationships between the supply, the demand and the prices, so-called Edgeworth box will be applied. This name refers to an English economist Francis Ysidro Edgeworth (1845-1926), who first applied this technique. This is a graphical analysis of feasible allocations in a pure exchange economy (superposition of two coordinate systems

for the analysis of a consumer's choice: the width of the rectangle = $\omega_{11} + \omega_{21}$, the height of the rectangle = $\omega_{12} + \omega_{22}$; the second system is rotated by 180°).



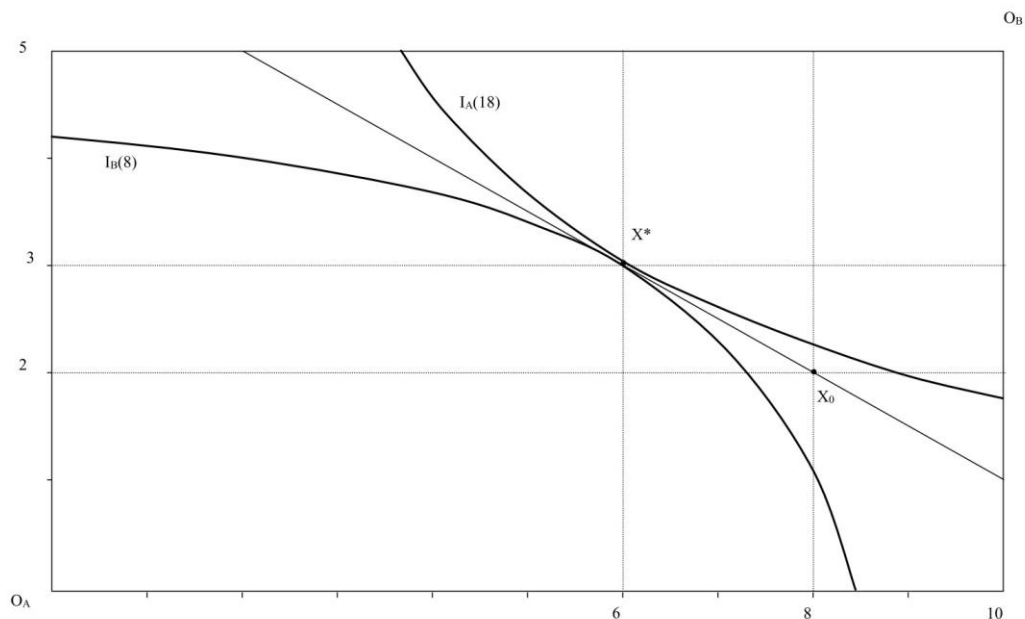
Please imagine the standard consumer choice model recalled in the third lecture (QF-3) – the left part of the picture above – with its image rotated by 180° – something that you see in the right part of the picture above. The good number 1 is measured horizontally, and the good number 2 is measured vertically. One detail you need to decide is how close to each other the coordinates should be. Edgeworth suggested that the "box" they make should be $\omega_{11} + \omega_{21}$ wide and the $\omega_{12} + \omega_{22}$ high. If we follow these suggestions, we can arrive at pictures like below (the shaded area indicates combinations of goods preferred by both consumers to what they hold initially).



Pictures refer to abstract goods number 1 and number 2. But let us be specific. The good number 1 is apples, and the good number 2 is oranges. The first person (A) brought 8 apples and 2 oranges, and the second person (B, whose axes were rotated by 180°) brought 2 apples and 3 oranges. Thus they have 10 apples and 5 oranges jointly. They can stick to what they have. However, they can also contemplate whether it would make sense to exchange some of them. The decisions may depend on relative prices of the goods they have. For example, person A can think: "if one apple goes for two oranges, I could exchange 2 of my apples for 4 oranges, but if two apples go for one orange I could also improve the utility from what I have". The person B can think: "if one apple goes for two oranges, I need to sacrifice two oranges to get an additional apple; this is not attractive, but – for instance – if one apple goes

for one orange, this would change my calculation". Having considered alternative relative prices, A and B may agree to exchange some of the goods they have in order to improve their utilities.

Edgeworth's trick allows us to have a correspondence of points in a plane with four numbers rather than two (as we usually have). Usually a point with coordinates (8,2) (8 apples and 2 oranges) corresponds to two numbers: 8 and 2. But in the Edgeworth box, like above, it corresponds to four numbers: 8 apples and 2 oranges owned by person A, and 2 apples and 3 oranges owned by the person B. We will see later, that in fact, it corresponds to five numbers. In addition to the numbers of apples and oranges owned by A and B, it shows the equilibrium price ratio, equal to the slope coefficient of the tangent line to both indifference curves in X^* (see picture below).



Now let us go back to the abstract terminology. Pictures illustrate the following situation. Indifference curves (i.e. the sets of points yielding the same utility) of the agent A, $I_A(\alpha)$ are given by the formula $x_{2A} = \alpha/x_{1A}$, while indifference curves of the agent B, $I_B(\beta)$ are given by $x_{2B} = \beta/x_{1B}$ ($\alpha, \beta > 0$ – parameters); additionally, we assume that the total quantity of the first good is 10, while that of the second – 5. Moreover, the diagram corresponds to the initial allocation of the first good 8:2, and to the initial allocation of the second one – 2:3 between consumers A and B (point X_0). There are two indifference curves containing this point: $x_{2A} = 16/x_{1A}$ ($\alpha = 16$) and $x_{2B} = 6/x_{1B}$ ($\beta = 6$). A would prefer to be on a higher indifference curve, say, in (9,3), i.e. on the curve $x_{2A} = 27/x_{1A}$ ($\alpha = 27$). At the same time, B would like to have more of everything too, i.e. to be in, say, (3,4), i.e. on the curve $x_{2B} = 12/x_{1B}$ ($\beta = 12$). It is impossible to satisfy these expectations at the same time. One solution which can place both agents in a jointly preferred point (one should solve a system of simultaneous equations) is: $x_{1A}^* = 6$, $x_{2A}^* = 3$, $x_{1B}^* = 4$, $x_{2B}^* = 2$, $\alpha = 18$, $\beta = 8$, $p = p_1/p_2 = 0,5$ (see the second picture). Agents A and B are on $I_A(18)$ and $I_B(8)$, respectively, and they are better off than in X_0 . One can see from the figure that they cannot improve their situations further simultaneously. In other words, (6,3) is a Pareto optimum (a discussion of this concept will be carried out in the next lecture, QF-7). Equilibrium prices which satisfy this solution are multiple, e.g. $p_1 = 1$, $p_2 = 2$, or $p_1 = 7$, $p_2 = 14$, or $p_1 = 0,5$, $p_2 = 1$ etc., as long as $p_1/p_2 = p = 0,5$.

Let us calculate the answer algebraically ($\mathbf{x}=[6,3,4,2]$ and $p_1=p_2/2$). In other words, we need to find x_{1A} , x_{2A} , x_{1B} , x_{2B} , α , β , and $p=p_1/p_2$. At the first glance it seems as if we had 7 unknowns. Nevertheless this number can easily be reduced to 5, since $x_{1B}=10-x_{1A}$ and $x_{2B}=5-x_{2A}$. Thus let $p=p_1/p_2$, $x_1=x_{1A}$ and $x_2=x_{2A}$. By the definition of the indifference curves:

$$\alpha=x_1x_2, \text{ and } \beta=(10-x_1)(5-x_2).$$

In the market equilibrium both indifference curves are tangent to the price ratio, i.e.

$$-\alpha/(x_1)^2=-p, \text{ and } -\beta/(10-x_1)^2=-p$$

(one needed to calculate the derivatives with respect to x_1 of functions

$$x_2=\alpha/x_1, \text{ and } x_2=5-\beta/(10-x_1))$$

which lets calculate α and β : $\alpha=p(x_1)^2$ and $\beta=p(10-x_1)^2$. Thus from the definitions of the indifference curves: $p(x_1)^2=x_1x_2$ and $p(10-x_1)^2=(10-x_1)(5-x_2)$; and after cancellations we get: $px_1=x_2$ and $p(10-x_1)=(5-x_2)$. These are 2 equations with 3 unknowns. We need an additional one. By the necessity of equating expenditures with revenues (please note that the same equation is derived irrespective of whether the balance reflects the consumer A or B) we get $p_1(4-x_1)=p_2(x_2-4)$, i.e. $p_1/p_2=(4-x_2)/(x_1-4)$ (with $x_1 \neq 4$). Hence the third equation (necessary to find three unknowns) reads

$$p=(4-x_2)/(x_1-4).$$

By substituting p into the previous equations we get:

$$x_1(4-x_2)=x_2(x_1-4) \text{ and } (10-x_1)(4-x_2)=(5-x_2)(x_1-4).$$

These are two equations with two unknowns: $4x_1-2x_1x_2=-4x_2$ and $60-9x_1+2x_1x_2=14x_2$. By adding the two we get $x_2=6-x_1/2$ which can be substituted into the first one, yielding

$$(x_1)^2-10x_1+24=0.$$

The last quadratic equation has two roots: $x_1=6$ or $x_1=4$ implying $x_2=3$ or $x_2=4$, respectively. The second solution has to be rejected (in order to calculate p properly; i.e. – as noted above – to avoid dividing into zero).

Hence the answer that could have been deduced from the picture, can be also calculated algebraically. Both persons are better off if they agree to the price ratio $p=p_1/p_2=1/2$ (one apple goes for two oranges), A sells two apples, and buys one orange, and B sells one orange and buys two apples. As a result of this exchange, A has 6 apples, and 3 oranges, while B has 4 apples, and 2 oranges. Both of them are now on higher indifference curves.

The picture on page 42 can be summarised in the following theorem.

In a perfectly competitive market with two consumers characterized by convex indifference curves, equilibrium in an Edgeworth box will be achieved in the point where indifference curves of these consumers are tangent to each other. The slope coefficient of the tangent line is equal (in absolute value terms) to the proportion of prices p_1^*/p_2^* .

What you saw on the picture on page 3 can be defined formally as a Walrasian (competitive) equilibrium. The concept originates from Marie-Esprit-Léon Walras (1834-1910) a French economist, who worked in the University of Lausanne. The adjective 'competitive' refers to the fact that the market equilibrium we talk about is competitive in a sense that buyers and sellers cannot manipulate prices (they take prices as given). The concept is formalised in the following definition.

X^* is a Walrasian (competitive) equilibrium in an Edgeworth box pure exchange economy if:

- $u_A(x_{1A}^*, x_{2A}^*) \geq u_A(x_{1A}, x_{2A})$ for all $(x_{1A}, x_{2A}) \in B_p(X_0)$,
- $u_B(x_{1B}^*, x_{2B}^*) \geq u_B(x_{1B}, x_{2B})$ for all $(x_{1B}, x_{2B}) \in B_p(X_0)$,
- $B_p(X_0) = \{(x_{1A}, x_{2A}, x_{1B}, x_{2B}) \in \mathbb{R}^4: p_1 x_{1A} + p_2 x_{2A} \leq p_1 \omega_{1A} + p_2 \omega_{2A}, \text{ and } p_1 x_{1B} + p_2 x_{2B} \leq p_1 \omega_{1B} + p_2 \omega_{2B}\}$

$B_p(X_0)$ is the budget set of consumers A and B, that is the set of bundles they can buy given the prices $p=[p_1, p_2]$. Please note that $p_1 \omega_{1A} + p_2 \omega_{2A}$ is the amount of money that A has if he/she decides to sell all his/her initial allocation. Likewise $p_1 \omega_{1B} + p_2 \omega_{2B}$ is the amount of money that B has if he/she decides to sell all his/her initial allocation. We call these numbers implicit incomes. Example below explains what in the Walras model is understood by money implicit in the initial allocation of goods.

An old lady without any revenues lives alone in a large villa worth 1 million € and starves. If she sold the villa (or rented it to somebody), she could meet all her needs like food, medication and a satisfactory housing. When asked about her preferences, she says that the villa is not for sale. Is the lady rich or poor? I asked this question in the class today, and students' opinions were not definite. The fact is that she has no cash. But on the other hand, her wealth is quite impressive; it is 1 million € not in cash, but in real estate. From the Walras model point of view this is irrelevant. Her wealth allows her to buy whatever we think that she needs for decent living. It is her preference that she wants to use it for living in the villa. If she sold it, she could have enough money to buy other things. But she does not want to.

The Walras model assumes that whatever we have has some implicit value. The value depends on market prices. We may prefer to keep it rather than sell and use the money to buy other goods, but this is our sovereign decision. In other words, "everything is for sale", even though sometimes we prefer to stick to what we have instead of using it to exchange for something else.

Moving from a pure exchange to the more realistic case with production requires somewhat more complicated notation. Let y_{hj} be the net supply of the j th commodity from the h th firm. Both negative and positive numbers are allowed. If $y_{hj} < 0$, then the h th firm uses more of the j th commodity than it produces. The number $-y_{hj} = |y_{hj}|$ is then the quantity of the (secondary) demand of the h th firm for the j th commodity.

Some students ask why there are no separate lists of inputs and outputs for firms. Instead there is a concept of a net supply; if it is positive, it can be understood as an output; otherwise it is a *de facto* input. But what is an output for some firms is an input for others. Moreover the same firm may use a product in both roles. For instance, IBM is considered a producer of computers. But at the same time it uses them in the production process. Thus computers for IBM are both inputs and outputs. The latter is higher than the former, so their net supply is positive according to this convention.

Definitions of the total market supply and demand are different than in a pure exchange economy model:

- $y_j = y_{1j} + \dots + y_{rj}$ – total net supply of the j th commodity
- $z_j = x_j - \omega_j - y_j$ – total excess demand for the j th commodity

Profit of the h th firm is defined in the following way: $\pi_h = p_1 y_{h1} + \dots + p_n y_{hn}$. The definition may look strange, because we are used to defining the profit as a difference between revenues and costs. But if you look at the expression $p_1 y_{h1} + \dots + p_n y_{hn}$ carefully, you will realise that some of the components $p_j y_{hj}$ are positive, and some are negative. Their sign depends on the sign of y_{hj} (because the sign of p_j is not negative). All the positive numbers refer to products that the firm sells rather than buys, and all the negative numbers refer to products that the firm buys rather than sells. Hence the expression includes the customary components of the firm's profit.

Share of the i th consumer in the h th firm's profit is defined as θ_{ih} . For many consumers these shares are zero. Nevertheless every firm must be owned by somebody: $\theta_{1h} + \dots + \theta_{kh} = 1$. If a firm makes a profit then these profits go to somebody. If a firm makes a loss then this loss must be financed by somebody. In the case of state owned firms, taxpayers have the right to claim their profits (or the obligation to finance their losses).

By a feasible allocation we understand a system which satisfies:

- $x_{11} + x_{21} + \dots + x_{k1} = \omega_{11} + \omega_{21} + \dots + \omega_{k1} + y_1,$
- $x_{12} + x_{22} + \dots + x_{k2} = \omega_{12} + \omega_{22} + \dots + \omega_{k2} + y_2,$
- ...
- $x_{1n} + x_{2n} + \dots + x_{kn} = \omega_{1n} + \omega_{2n} + \dots + \omega_{kn} + y_n,$

where $y_j = \sum_{h=1}^H y_{hj}$ (the total net supply, while net supplies of respective firms (y_{hj}) come from their production sets: $(y_{h1}, \dots, y_{hn}) \in Y_h$).

Everything is like in the pure exchange economy, except that in addition to what the consumers have, there are some outputs of firms. Consumer demand for products and their prices can be understood as vectors with n coordinates. It is customary to consider the demand as a column vector, and prices as a row vector. T denotes a transposition. The numbers ω_{ij} and θ_{ih} are parameters of the general equilibrium model. The numbers x_{ij} , y_{hj} , and therefore also z_j are functions of prices $\mathbf{p}$.

- $\mathbf{x} = (x_1, \dots, x_n)^T$ – demand vector (column)
- $\mathbf{p} = (p_1, \dots, p_n)$ – price vector (row)

Using this notation allows to define the Walras equilibrium in the following way:

Walras equilibrium is any pair $(\mathbf{p}^*, \mathbf{x}^*)$ such that for every commodity j we have:

$$x_j^* := x_{1j}(\mathbf{p}^*) + \dots + x_{kj}(\mathbf{p}^*) \leq \omega_{1j} + \dots + \omega_{kj} + y_{1j} + \dots + y_{hj}, \text{ i.e. } z_j(\mathbf{p}^*) \leq 0.$$

We have an inequality which states that the left hand side is smaller or equal to the right hand side. The demand is on the left, and the supply is on the right. The Walras equilibrium thus reads:

Demand $\leq$ Supply (in all markets).

Equilibrium is usually understood as an equality rather than inequality. Hence some people ask the question why it is assumed that for some markets (products) the demand can be lower than the supply. The reason is the logic of market exchange. Price mechanism is supposed to trigger adaptations of the demand and supply; if the demand is lower than the supply, the price should go down; if the demand is higher than the supply, the price should go up. However prices cannot be negative. If a price (of some product) is zero and the demand is still

lower than the supply, no additional adjustments can be expected – the Walras mechanism has been exhausted.

In a pure exchange economy the budget line of the i th consumer was $p_1x_{i1} + \dots + p_nx_{in} = p_1\omega_{i1} + \dots + p_n\omega_{in}$ (the right hand side was the implicit income – see page 5). In the third lecture (QF-3) income to be spent on purchasing goods was m (without analysing where it comes from). Here the equation reads: $p_1x_{i1} + \dots + p_nx_{in} = p_1\omega_{i1} + \dots + p_n\omega_{in} + \theta_{i1}\sum_j p_jy_{1j} + \dots + \theta_{ir}\sum_j p_jy_{rj}$. At the left hand side there is what the i th consumer spends, and at the right hand side there is what is available to the i th consumer (the implicit value of the endowment, $p_1\omega_{i1} + \dots + p_n\omega_{in}$, and the share in firms' profits, $\theta_{i1}\sum_j p_jy_{1j} + \dots + \theta_{ir}\sum_j p_jy_{rj}$). The equation is called the consumer's budget line. It is consistent with the definition of BL from the third lecture (QF-3), because it simply specifies where the available money m comes from.

These definitions allow to state the following Walras Law.

Theorem (Walras Law)

If all the consumers keep their budget lines then the value of excess demand is 0. In other words, $\sum_j p_j z_j = 0$.

Proof:

$$\begin{aligned} \sum_j p_j z_j & \\ & \stackrel{1=}{=} \sum_j p_j (x_j - \omega_j - y_j) = \\ & \stackrel{2=}{=} \sum_j p_j (\sum_i x_{ij} - \sum_i \omega_{ij} - \sum_h y_{hj}) = \\ & \stackrel{3=}{=} \sum_j p_j (\sum_i x_{ij} - \sum_i \omega_{ij} - \sum_h (\sum_i \theta_{ih}) y_{hj}) = \\ & \stackrel{4=}{=} \sum_j \sum_i (p_j x_{ij} - p_j \omega_{ij} - \sum_h \theta_{ih} p_j y_{hj}) = \\ & \stackrel{5=}{=} \sum_j \sum_i (p_j x_{ij} - \sum_j p_j \omega_{ij} - \sum_h \theta_{ih} \sum_j p_j y_{hj}) = \\ & \stackrel{6=}{=} \sum_j 0 = 0. \end{aligned}$$

Explanation of steps:

- 1 by definition of excess demand
- 2 definitions of x_j , ω_j and y_j were substituted
- 3 the factor $\theta_{1h} + \dots + \theta_{kh} = 1$ was inserted for every h
- 4 multiplication by p_j was carried out
- 5 the order of summation was changed
- 6 the expression to be summed is the difference between the left-hand-side and the right-hand side of a budget line.

There are two frequent mistakes that need to be avoided. First, the Walras Law has nothing to do with the definition of the Walras equilibrium; please note that these are completely different concepts. Second, the Walras Law holds for any system prices $\mathbf{p}$, not only for equilibrium prices.

The Walras Law implies several important corollaries. Four of them are listed below (without proofs, which are very easy).

- 1 If the value of excess demand in $n-1$ markets is equal to 0, then also in the remaining n th market the value of excess demand is 0.

- 2 In order to find the Walras equilibrium it is sufficient to solve a system of $n-1$ equilibrium prices in corresponding $n-1$ markets.
- 3 A free good is any good the excess demand for which is negative. In equilibrium the price of such a good is zero.
- 4 A desired good is any good whose excess demand at the zero price is positive. If all goods $1, \dots, n$ are desired, then in equilibrium the excess demand in all markets is 0.

Questions and answers to lecture 6

6.1 What is the difference between total demand and excess demand?

The total demand is what a consumer demands at given prices. The excess demand (sometimes called net demand) is the difference between what the consumer demands and what he (or she) holds as the initial allocation (endowment). In the class example, the first consumer brought 8 apples and 2 oranges ($\omega_{11}=8$, and $\omega_{12}=2$). His (or her) total demand (final allocation) was 6 apples and 3 oranges ($x_{11}=6$, and $x_{12}=3$). Thus his (or her) excess demand for apples was -2, and the excess demand for oranges was +1. These numbers indicate how many units of a good the consumer wants to buy. The negative value of excess demand means that the consumer wants to sell it rather than buy.

6.2 Calculate the Walrasian equilibrium (including the price ratio) in a pure exchange economy where the consumers have the allocations of the two goods (4,4) and (6,1), respectively, and their indifference curves are given by the following hyperboles: $x_{2A}=\alpha/x_{1A}$, and $x_{2B}=\beta/x_{1B}$, respectively ($\alpha, \beta > 0$ – parameters).

Some students may recognise that the Edgeworth box is identical to what we see in QF-17 (page 41 in this text), except that the initial allocation corresponds to the upper intersection of indifference curves $x_{2A}=16/x_{1A}$ ($\alpha=16$) and $x_{2B}=6/x_{1B}$ ($\beta=6$). We will solve the problem algebraically.

Obviously the answer is the same as for the class example ($\mathbf{x}=[6,3,4,2]$ and $p_1=p_2/2$). Let us find x_{1A} , x_{2A} , x_{1B} , x_{2B} , α , β , and $p=p_1/p_2$. As before, it seems as if we had 7 unknowns. This number can easily be reduced to 5, since $x_{1B}=10-x_{1A}$ and $x_{2B}=5-x_{2A}$. As before, let $p=p_1/p_2$, $x_1=x_{1A}$ and $x_2=x_{2A}$. By the definition of the indifference curves: $\alpha=x_1x_2$, and $\beta=(10-x_1)(5-x_2)$. Having calculated the derivatives with respect to x_1 of functions $x_2=\alpha/x_1$, and $x_2=5-\beta/(10-x_1)$, and knowing that in the market equilibrium both indifference curves are tangent to the line of price ratio, i.e. $-\alpha/(x_1)^2=-p$, and $-\beta/(10-x_1)^2=-p$ lets us calculate α and β : $\alpha=p(x_1)^2$ and $\beta=p(10-x_1)^2$. When substituted to the definitions of indifference curves, it gives $p(x_1)^2=x_1x_2$ and $p(10-x_1)^2=(10-x_1)(5-x_2)$. After cancellations we get: $px_1=x_2$ and $p(10-x_1)=(5-x_2)$. The rest is as in the class. By the necessity of equating expenditures with revenues, we get $p_1(4-x_1)=p_2(x_2-4)$, i.e. $p_1/p_2=(4-x_2)/(x_1-4)$ (please remember that $x_1 \neq 4$). Hence the additional equation (needed in order to find the three unknowns) reads $p=(4-x_2)/(x_1-4)$. By substituting p into the previous equations we get: $x_1(4-x_2)=x_2(x_1-4)$ and $(10-x_1)(4-x_2)=(5-x_2)(x_1-4)$. These are two equations with two unknowns: $4x_1-2x_1x_2=-4x_2$ and $60-9x_1+2x_1x_2=14x_2$. By adding the two we get $x_2=6-x_1/2$ which can be substituted into the first one, yielding $(x_1)^2-10x_1+24=0$. The last quadratic equation has two roots: $x_1=6$ or $x_1=4$ implying $x_2=3$ or $x_2=4$, respectively.

The second solution has to be rejected (in order to calculate p properly; i.e. – as noted above – to avoid dividing into zero).

6.3 Demonstrate that in a pure exchange economy where the consumers have the allocations of the two goods (8,2) and (2,3), respectively, and their indifference curves are given by the following hyperboles – $(x_{2A})=\alpha/(x_{1A})$ and $(x_{2B})^{1/2}=\beta/(x_{1B})^{1/2}$, respectively ($\alpha,\beta>0$ – parameters) – the Walrasian equilibrium consists of the allocation (6,3,4,2) and the price ratio 1:2.

If you observe that by raising to the second power the definition of indifference curves for B you obtain the class example with constant β^2 instead of β , the calculations are the same as in the class, except that the constant β^2 should be put instead of β . Moreover, the calculations do not have to be carried out, since the Walras equilibrium has been calculated already: the allocation is (6,3,4,2), and the price ratio is 1:2. What needs to be done is to check that the slope coefficient of both indifference curves for the allocation (6,3,4,2) is the same and it is equal to -1:2.

The slope coefficient of the first indifference curve is $-\alpha/(x_1)^2=-18/36=-1/2$. The slope coefficient of the second indifference curve is $-\beta/(10-x_1)^2=-8/16=-1/2$. They are the same and equal to -1/2.

6.4 Please prove (graphically or algebraically) that in a pure exchange economy where consumers have the allocations of two goods (8,2) and (2,3), respectively, and their indifference curves are given by the following hyperboles: $x_{2A}=\alpha/x_{1A}$, and $x_{2B}=\beta/x_{1B}$, respectively ($\alpha,\beta>0$ – parameters), the allocation (7,3) and (3,2) is preferred by the first and offers the same well-being for the second, but it cannot be attained as a Walras equilibrium.

Graphical proof is perhaps easier, but it requires drawing very accurate pictures. Thus it is better to solve the problem algebraically. First we will demonstrate that the allocation (7,3) is preferred by the first player. According to the formula $\alpha=x_{1A}x_{2A}$, the initial allocation gave the utility of 16, and the second one – 21. Hence it is preferred. Now we will demonstrate that (3,2) gives (to the second consumer) the same utility as (2,3). But this is obvious, since $\beta=x_{1B}x_{2B}$ (in either case the product is 6). Hence the first part of the question is answered.

This allocation cannot be attained as a Walras equilibrium. In order to prove this, one has to recall that in Walras equilibria offers of both consumers bring them to the highest indifference curve they can afford given the prices. In other words, their indifference curves must be tangent to each other (if they are not tangent – if they intersect like in the picture on page 41 – both consumers are better off when they offer some additional exchange). Hence it is sufficient to check if the indifference curves are tangent or not. In (7,3) the slope coefficient of the tangent line to the first indifference curve is $-\alpha/(x_1)^2=-21/49=-3/7$. The slope coefficient of the tangent line to the second indifference curve is $-\beta/(10-x_1)^2=-6/9=-2/3$. This ends the proof.

6.5 Please prove the first corollary from the Walras Law.

The first corollary reads: "If the value of excess demand in $n-1$ markets is equal to 0, then also in the remaining n th market the value of excess demand is 0". Its proof is obvious. The Walras Law states that $\sum_j p_j z_j=0$. We can split this into two sums:

$$(p_1z_1+p_2z_2+\dots+p_{n-1}z_{j_{n-1}})+p_nz_n=0.$$

If the first is equal to zero, the second must be equal to zero too.

6.6 Please give an example of a good which is not desired (please see the fourth corollary from the Walras Law).

Our student library analyses annually what textbooks are read and what are not useful any more. The useless ones are displayed in the hall, and they can be taken for free. Apparently the demand for these textbooks is lower than the supply, despite the fact that their price is zero. Some of them are never picked up by anybody.

7 – Welfare economics

Welfare economics is a somewhat strange name for analyses of relationships between Pareto optima and Walras equilibria. In the previous lecture (QF-6) we observed that – at least in an Edgeworth box – a Walras equilibrium was achieved where the indifference curves were tangent to each other and thus economic agents enjoyed a Pareto optimum (please see its definition below). This finding was an example of the First Theorem of Welfare Economics. In general welfare economics theorems – the First and the Second – do not have to confine to Edgeworth boxes and to a pure exchange.

There are two fundamental theorems in so-called welfare economics: the first and the second. They establish a relationship between Pareto optima (PO) and Walras equilibria (WE). The first corresponds to the implication $WE \Rightarrow PO$, and the second: $PO \Rightarrow WE$. What they establish is almost an equivalence. Strictly speaking this is not an equivalence, because these theorems require different assumptions. As a rule, the first is easy, and it does not require almost any assumptions. The second is more difficult, and it requires some mathematical assumptions, like the convexity of certain functions. Welfare economics theorems are stated and proved in many areas of economics. In the beginning we will confine to the first one phrased in the language of a pure exchange economy.

But before proving it, let me make sure that everybody knows what we talk about. Most students are perhaps familiar with the definition of a Pareto optimum. Let me recall it briefly just in case somebody forgot it. A Pareto optimum is such an outcome that nobody can be made better off unless somebody else is made worse off. Assuming that there are 2 apples to be allocated to 2 people (both of them like apples), both can be given 1 apple. This is a Pareto optimum, since neither can be given 2 apples unless the other person is deprived of an apple (but this would make him or her worse off). The allocation can be different. One privileged person is given 2 apples while the other one is left without any. Surprisingly, this is a Pareto optimum too, since in order to make the deprived person better off, at least 1 apple should be taken away from the privileged person making him or her worse off.

Different allocations – some of them considered unfair – can satisfy the definition of a Pareto optimum. To see an allocation which violates it, let us consider 2 apples distributed in the following way. 1 apple goes to one person, the other person gets nothing, and 1 apple is left aside. This is not a Pareto optimum, because one person can be made better off by getting the apple left aside, and nobody will be made worse off. Intuitively, a Pareto optimum can be

understood as an allocation which uses up everything what is available (not necessarily in a fair way).

Our review of welfare economics starts with a theorem proved for the pure exchange economy case.

The first fundamental welfare economics theorem

In an Edgeworth box pure exchange economy, if X^* is a Walrasian equilibrium then X^* is a Pareto optimum.

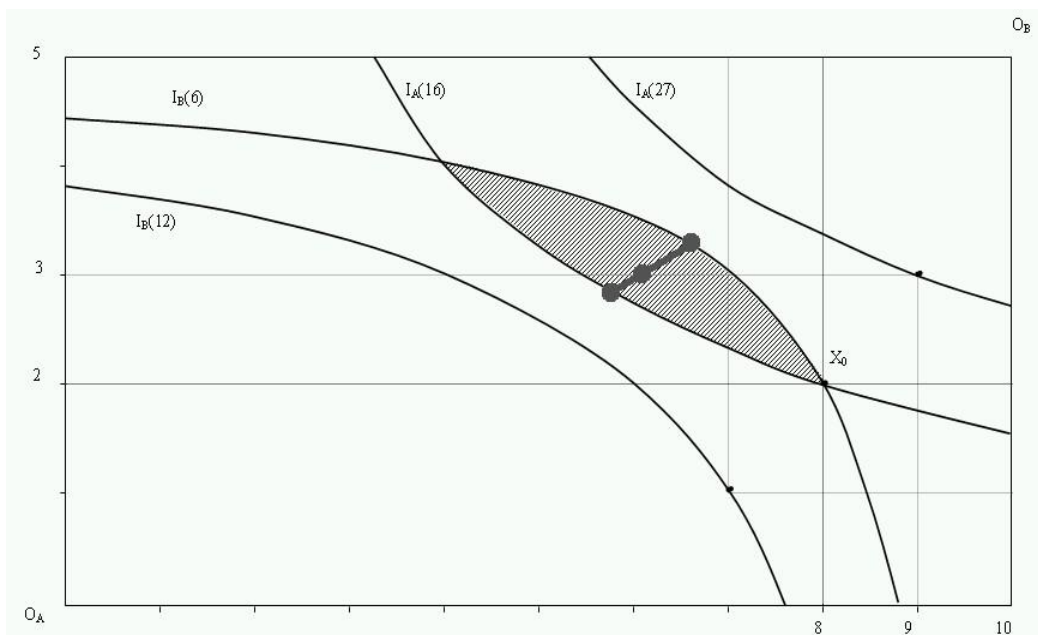
Proof:

It is sufficient to demonstrate that the allocation X^* cannot be improved in a sense that both consumers are made better off at the same time. If exchanges were to be carried out at given prices this would have been obvious given both inequalities (no improvements are possible) and the feasibility condition (each consumer can spend not more than he/she obtains from selling his/her endowments). But we have to take into account other exchange methods that do not necessarily apply the Walras price mechanism described in the previous lecture (QF-6).

Thus let us assume (in order to get a contradiction) that there is another feasible allocation X' which is at least as good as X^* for one consumer and strictly better for the other one. But then

$$p_1(x'_{1A} + x'_{1B}) + p_2(x'_{2A} + x'_{2B}) > p_1(x^*_{1A} + x^*_{1B}) + p_2(x^*_{2A} + x^*_{2B}),$$

which means that apparently X' must have been too expensive to be chosen in equilibrium (X^* was chosen instead). However, in a pure exchange economy: $x'_{1A} + x'_{1B} = x^*_{1A} + x^*_{1B} = \omega_1$, and $x'_{2A} + x'_{2B} = x^*_{2A} + x^*_{2B} = \omega_2$. Hence $p_1\omega_1 + p_2\omega_2 > p_1\omega_1 + p_2\omega_2$ which is a contradiction which ends the proof.



In order to complete the discussion of the pure exchange economy, an additional definition has to be introduced. A contract curve is the set of all allocations which are Pareto optima and

which are preferred by both players over their initial allocation. It is depicted as a thick line in the picture above (this is an Edgeworth box from the previous lecture, QF-6).

From now on, our review of welfare economics will be more general than a simple pure exchange.

There is asymmetry between the First and the Second welfare economics theorems. The First establishes implication $WE \Rightarrow PO$, and the Second – $PO \Rightarrow WE$. The First is easy to prove and it does not require many mathematical assumptions. The Second is more difficult to prove and it requires assumptions that were not necessary in the First one. As they require different assumptions, welfare economics theorems do not mean an equivalence ($\Leftrightarrow$) between the two concepts.

Nevertheless welfare economics theorems establish a close relationship between WE and PO. This is remarkable, since the two concepts seem to be rather far from each other. WE is a technical concept, requiring certain freedom (voluntariness of transactions) and the existence of market institutions. In contrast, PO is a universal notion that can be applied to any social organisation, looking at whether members of this organisation (consumers) can be made better off without hurting one another. And all of a sudden, it turns out that both ideas are close to each other. WE can be sought in a market economy only. Analysing PO makes sense in a non-market economy too. The predicament of people in ancient Egypt or in centrally planned economies can be analysed in terms of potential improvements, even though a full-fledged price mechanism is absent there.

The First welfare economics theorem requires two easy assumptions: (1) economic agents are price-takers; and (2) transactions are voluntary. The first assumption means that nobody can manipulate prices. A monopoly is not a price-taker, since by increasing its supply, the monopoly can lower the price, and by decreasing the supply, it can elevate the price (competitive firms do not have such a power; their individual supply is irrelevant for the market price). The second assumption is so obvious that many textbooks do not even mention it. Nevertheless it is important to be aware that the idea of WE is based on voluntariness of transactions: economic agents do not buy and sell unless they can increase their well-being.

The second welfare economics theorem is more complicated. In addition to what the first theorem assumes, it requires convexity of mathematical objects involved. Namely, it requires all firms to have convex production sets, and all consumers to have convex indifference curves. The theorems read:

First Welfare Economics Theorem

If $(\mathbf{p}^*, \mathbf{x}^*)$ is a Walras equilibrium in a market whose all participants are price-takers interacting with each other exclusively through voluntary transactions, then $\mathbf{x}^*$ defines a Pareto optimum.

Second Welfare Economics Theorem

Let us assume that

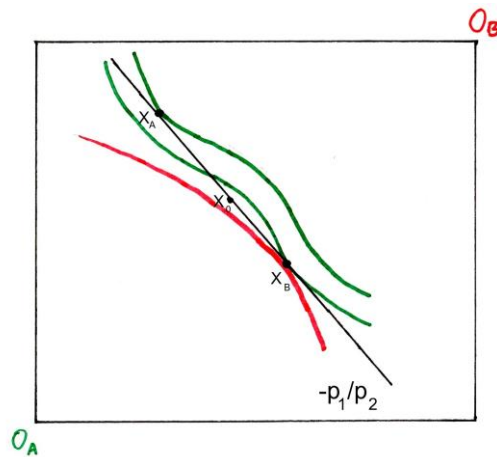
- all agents are price-takers,
- all firms have convex production sets,

- all consumers have convex indifference curves (surfaces).

Then for any Pareto optimum $\mathbf{x}^*$ there exists a price vector $\mathbf{p}^*$ and an initial allocation of endowments such that $\mathbf{x}^*$ can be attained as a Walras equilibrium $(\mathbf{p}^*, \mathbf{x}^*)$

As noted earlier, proofs of the theorems are not straightforward. The first one can be proved fairly easily (see above), at least for a pure exchange economy (for a more general case the proof is somewhat longer). The proof of the second theorem is more complicated.

An Edgeworth box can be used to graphically prove the Second Welfare Economics Theorem in a special case of a pure exchange economy with two goods and two consumers. Picture below demonstrates what may happen if the convexity assumption is violated.

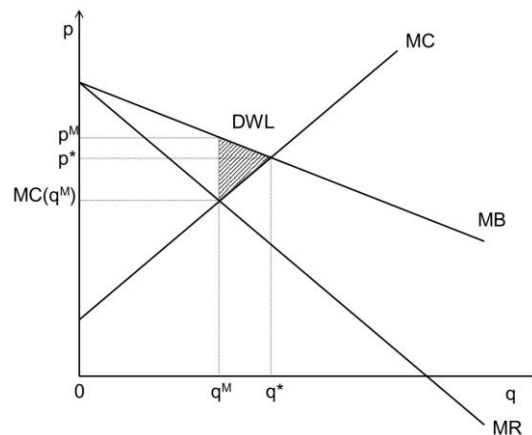


This is an Edgeworth box for two consumers: the green one (A), and the red one (B). The red one has convex indifference curves, while the green has indifference curves that are not convex. X_B is a tangency point of two indifference curves and therefore it is a Pareto optimum. The black line is a tangent to both indifference curves. Its slope coefficient ($-p_1/p_2$) gives the price ratio that can bring the consumers to this Pareto optimum. The two consumers have endowments X_0 . The point is located on this tangent line. But when the consumers face the price ratio p_1/p_2 (the only one that can bring them to this Pareto optimum), then the red one would like to sell some units of the first good (please note that B's horizontal axis runs from right to left), buy some units of the second good (please note that B's vertical axis is directed upside down), and enjoy the point X_B . At the same time, the green consumer would like to sell (not to buy!) some units of the first good, buy some units of the second good (not to sell!), and enjoy the point X_A . The consumers have offers that are not consistent (they would like to buy what the other one does not like to sell), and the Walras equilibrium cannot be attained. Please note that the same conclusion will be reached if the initial allocation was different than X_0 located on the same tangent line.

It is easy to see that Walras equilibrium in a monopolistic market is not a Pareto optimum (neither of the welfare economics theorems applies). In economic textbooks the following graph can be found.

Picture below explains why the equilibrium in a monopolistic market is not a Pareto optimum. It demonstrates what happens if buyers or sellers are not price-takers (they can manipulate prices). A monopolist can elevate the price from p^* to p^M by lowering the supply from q^* to

q^M , since the monopolist knows that the marginal revenue (MR) is lower than the marginal benefit (MB) for q^* ; the latter is equal to the price obtained, but the former is lower (since the price goes down if the supply goes up). By manipulating prices, the monopolist reduces the economic surplus by so-called *Deadweight Welfare Loss* (DWL). Students who are not very familiar with basic microeconomic concepts do not have to analyse the case in a detailed way. They simply should be aware of the fact that market equilibrium is not always a guarantee for the economic surplus maximisation which is a prerequisite for enjoying a Pareto optimum.



"Deadweight Welfare Loss" owes its name to the fact that what is a loss to somebody, can be a gain to somebody else; "deadweight" means that nobody gains. We do not like monopolies, since the elevated prices they charge make us worse off. However, monopolistic prices benefit somebody. Think of a local cooperative owned by poor farmers which enjoys a monopolistic power, but distributes its entire profit to the local community. Customers have to pay high prices, but the revenues go to the owner, i.e. to the local community. "Deadweight", however, means that the shaded part of the economic surplus, where the benefits originate from, is lost. In other words, potential beneficiaries have less to draw from compared to when the supply is higher (not lowered by the monopoly).

The following story illustrates another problem with monopolies. Let us suppose that a monopolist charges 10 € (instead of the marginal cost of 8 €). It has practised this for some time every day from 11:00 AM until 6:00 PM. One day, at 6:00 PM it announced that it would sell for 20 more minutes for the price of 9 €. Those who were willing to pay 10 €, bought it already, but there were some people who were not willing to pay 10 €, who are willing to pay 9 €; they did not buy before 6:00 PM for 10 €, yet they bought it after 6:00 PM and they enjoy some consumer surplus. Also the monopolist increased its (producer) surplus, since it made the profit of 1 € ($9 - 8 = 1$) per piece. What can be expected on the next day? Few customers will buy before 6 PM, since people will wait for the price be lowered again. That is why monopolistic markets are inefficient. The monopolist sets the price so as to maximise its profit, but the monopolistic price does not maximise the economic surplus (which can be increased by the trick mentioned above). Yet if the monopolist makes the trick, the monopolistic price cannot be maintained. In other words, a monopolistic equilibrium is not a Walras equilibrium (question 7.8 addresses the problem of differentiated prices we are exposed to).

The lack of price taking is not the only reason why Walras equilibria are not equivalent to Pareto optima. As indicated above, violating convexity assumptions affects the second welfare economics theorem too.

The two other important cases when welfare economics theorems cannot be applied are related to externalities and asymmetric information. The latter will be addressed in more detail in the next lecture (QF-8). Here we will look briefly at the externality problem. The lack of externalities and the lack of asymmetric information are often considered obvious and therefore they are not even mentioned in welfare economics theorems.

Externalities are linked to the voluntariness of transactions. The following example illustrates this.

Let us assume that a power plant located by a lake discharges wastewater and pollutes the lake. The lake contamination makes fishermen suffer from lower profits (the quantity or quality of fish goes down). This is a so-called negative externality (external cost). The polluter imposes some losses not because the price of fish goes down or the price of raw materials goes up, but imposes them on the fishermen directly.

Economists observe that an externality arises when there is no market for the factor responsible for the externality (for instance, when property rights are ill-defined). In our lake pollution example above the missing market means that the fishermen are not asked if the plant can discharge its sewage (the power plant does not have to buy any pollution permit). "Transactions" (in this case polluting the lake) are not voluntary. They simply take place without asking anybody for a permission.

Questions and answers to lecture 7

7.1 Do lower left and upper right corners of an Edgeworth box (0_A and 0_B) belong to the contract curve?

Yes, they do. The first corresponds to the allocation such that A does not have anything, and the second to the allocation such that B does not have anything. These are Pareto optima, since there is no transaction that makes both of them better off simultaneously.

7.2 Is it possible that a Pareto optimum cannot be achieved as a Walras equilibrium?

Yes. The second welfare economics theorem lists mathematical assumptions that need to be satisfied to let Pareto optima be attained as Walras equilibria. One condition is that indifference curves are convex. The lecture provides an example of what may go wrong when they are not. In addition not every Pareto optimum can be achieved as a Walras equilibrium given initial endowments. For instance 0_A and 0_B are Pareto optima (see 7.1 above). Yet if the initial allocation (endowment) is 0_A then 0_B cannot be achieved, and *vice versa*.

7.3 Is a Walras equilibrium a solution to some optimisation problem?

No. It does not have to be, even though for many years economists thought that an equilibrium must be an optimum of some sort. The definition of an equilibrium states that economic agents do not have incentives to change their decisions. The definition of Walras equilibrium states that consumers and firms do not have incentives to buy or sell different amounts. Consumers maximise their utility; firms maximise their profits. Whatever they do

cannot be interpreted as a maximisation of a single objective function. Nevertheless, for many generations economists believed that what is achieved as an equilibrium must be optimal in some sense (the First Theorem of Welfare Economics states that a Pareto optimum can be achieved in some cases indeed).

7.4 Is it possible to assume that economic agents are price takers in an Edgeworth box?

Yes, but it requires a special setup. Price-taking is often associated with a large number of participants, so that nobody's individual action (like selling or buying) can change the market price. Edgeworth box with two consumers endowed with two goods can be best described as a bilateral monopoly (the seller knows that he/she is the only seller, and the buyer knows that he/she is the only buyer). Both can act strategically and both can make offers which are aimed at manipulating prices.

In some economics textbooks, the price-taking assumption is reconciled with the small number of agents (two) by arguing that even though we see only two of them in the Edgeworth box, in fact there are many of them; what we see are just representatives of typical numerous consumers. This is not convincing to me.

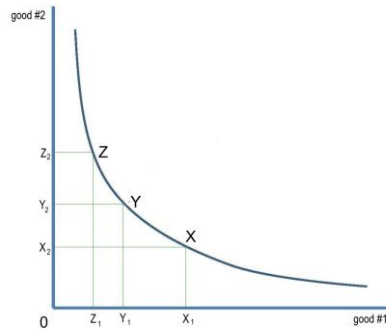
There is an explanation of how the two consumers can be considered price-takers (they cannot manipulate prices), but it requires the presence of the third agent called Walrasian Auctioneer (WA). The person is impartial, i.e. not interested in an attractive outcome for either of the consumers A and B. The WA announces a price ratio, say 1:1, and asks A and B to write their respective offers on pieces of paper so that they do not know what the other wrote. WA collects the offers and compares them. If A is willing to sell exactly what B is willing to buy, then WA asks them to make the exchange, as offered. Otherwise WA states that offers were not consistent (without stating who offered what), and announces a new price ratio, say, 3:1 and asks for written offers. If the offers turn out to be consistent, then A and B are allowed to make the exchange as offered. Otherwise WA announces a new price ratio, say, 2:1 and repeats the procedure once again. If WA is competent (knows that if the demand is higher than the supply, the price should go up), the procedure converges to WE. Please note that neither A nor B can manipulate prices; they have to take them as given.

7.5 What is the justification of convexity of indifference curves?

Convexity is justified empirically in many cases. For instance, production sets are found to be convex. If one technology allows to produce 200 bicycles, and another one allows to produce 20 motorcycles, then by taking 60% of the first one, and 40% of the second one, a firm producing 120 bicycles and 8 motorcycles can be contemplated. Convexity of indifference curves reflects the fact that consumers agree to decreasing the amount of one good only if this is compensated by a higher availability of the other good, and they do not want to have a very small amount of either good.

Picture below illustrates this. Bundles X, Y, and Z belong to the same indifference curve. X combines X_1 of the first good, and X_2 of the second good. The consumer is indifferent between this bundle and Y where Y_1 and Y_2 are combined. The amount of the first good decreased (from X_1 to Y_1), but the amount of the second good increased (from X_2 to Y_2). If the amount of the first good decreases further (from Y_1 to Z_1 ; that is much less than before), the consumer expects the amount of the second good to increase (from Y_2 to Z_2 ; that is more

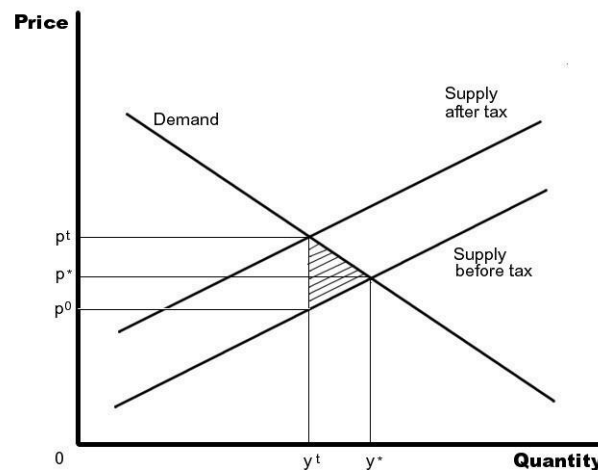
than before). This logic makes the indifferent curve convex. Non-convexity of indifference curves can be found sometimes, but this is not typical for consumer behaviour.



7.6 Can transactions be involuntary?

Yes, they can, but this is not allowed in the Walras market logic. For instance, some producers can be forced to sell something at prices dictated by buyers. During wars firms were forced to supply what was considered necessary for the success of military operations. In centrally planned economies farmers were obliged to sell their products to government agencies at prices set by the government. Real estate owners are not always free to take decisions regarding their property. If there are constraints of that sort, original Walras model cannot reflect what happens in the economy.

7.7 *Deadweight Welfare Loss* is expected as a result of taxation (not only as a result of monopolisation). Can you explain the mechanism?



The mechanism is similar to suffering DWL in a monopolistic market. Imposition of the tax increases the price from p^* to p^t . This is the price paid by the buyers. The sellers receive only p^0 . The difference $p^t - p^0$ is the unit tax. As a result, the quantity traded goes down from y^* to y^t . Economic surplus decreases. The shaded triangle illustrates this loss of economic surplus. It is not clear though, who benefits from the tax, and who loses. The tax revenue collected by the government is $(p^t - p^0)y^t$. This can be more than the shaded loss. It is up to the government to determine how the tax revenues are distributed. They can flow to the producers (e.g. in the

form of subsidies) or they can flow to consumers (e.g. in the form of scholarships for talented students, support for unmarried mothers, health care expenditures of senior citizens, or what have you). Nevertheless, irrespective of whether they flow to producers or consumers, the economic surplus is lower than before; the society loses (consumers gain less than producers lose or *vice versa*). Therefore this is a "Deadweight Welfare Loss".

7.8 The Walras model assumes that there is only one price of a given product – the same for everybody. How can this be reconciled with what we observe every day: prices charged in attractive locations are higher than elsewhere?

There are at least two ways to reconcile this logic with every-day experience. One makes use of the transportation cost concept, and another refers to the transaction cost concept.

The transportation cost was introduced to economics by Johann Heinrich von Thünen (1783-1850). He argued that the price of the product includes inevitable transportation cost (from the producer to the user). Therefore prices observed in different locations may reflect the fact that some effort is required to have products made available there. This explains why prices charged far from where most customers live but perhaps close to where the supply originates are lower.

The transaction cost was introduced to economics by Ronald Coase (1910-2013). It is defined as the cost of preparing and enforcing contracts. This sounds complicated, but often it boils down to whether buying something is safe or not. The same product if sold by an unknown vendor can be cheaper than sold by a credible supplier. Yet the transaction cost may be higher when we buy from a less credible seller.

An attractive location can be understood as the one which is reached easier or the one where the probability of failure (being left with an unsatisfactory purchase without a possibility to get the money back) is lower. Yet this does not exhaust situations where the same product may have different prices in different places. In financial economics this problem is solved by the theory of arbitrage (see lecture 13, QF-13). The arbitrage possibility means that somebody can buy a product at a lower price and sell it at a higher one. Neither transportation costs nor transaction costs are sufficient to explain the outcome. The same product – say, a soft drink of the standard quality – has very different prices in different places just because we are willing to pay more or less depending on circumstances. In order to reconcile this with the Walras logic, economists assume that these are different products. They are physically the same, but offered in different circumstances (traded in different markets). This allows them to be treated as different products (arbitrage is possible, but perhaps not attractive – because of the costs required).

Whatever can be said about different locations, can be repeated about different timing. Bars and restaurants may charge different prices at different times (for instance you have to pay for the lunch or beer less if you come in a certain moment – as observed by one of the students today). Firms charge various prices for the same product when it is justified by the demand elasticity. Economic theory addresses this problem by assuming that there are different markets (commodities).

8 – Asymmetric information

The definition of asymmetric information is obvious: the buyer has less information about the commodity than the seller or *vice versa*; acquiring information is possible, but costly. Consequences of asymmetric information can be serious, since it implies a market failure. There are two important cases:

- Equilibrium does not have to be a Pareto optimum;
- Equilibrium may not exist.

Typical examples of asymmetric information include:

- Used cars,
- Insurance policies,
- Employment contracts.

We will analyse these examples in order to see whether Walras equilibrium can be expected. Sometimes the buyer knows less than the seller, like in the second-hand car market. The seller of the car (who drove it for several years) knows it better than the buyer (who takes a test drive for several minutes). If a prospective buyer asks the seller "Is this car a good one?" the seller is likely to answer "Yes, of course!" irrespective of whether this is true or not. Everybody knows this (students confirmed this in the class), so buyers do not even ask questions of that sort. Instead they try to acquire as much technical information about the car as they can, but they will never know what the seller knows (but perhaps does not want to share).

The asymmetric information in the second-hand car market is harmful for economic efficiency as the following argument explains.

Let us assume that there are two kinds of cars: good ones (50%) and bad ones (50%). The owner of a good car is willing to sell it at the price of at least 8,000 € – this is what the car is worth to him (or her). A prospective buyer is willing to pay at most 10,000 € – this is what the car is worth to him (or her). They can agree to the price of 9,000 €. But irrespective of the price agreed, as a result of transaction, economic surplus increases by 2,000 €: in our example the seller receives 9,000 € for something he (or she) valued at 8,000 €, so he (or she) is left with 1,000 € more, and the buyer paid 9,000 € for something that is worth 10,000 €, so he (or she) enjoys the surplus of 1,000 €.

At the same time there are bad cars. The owner of a bad car is willing to sell it at the price of at least 3,000 € – this is what the car is worth to him (or her). A prospective buyer is willing to pay at most 5,000 € – this is what the car is worth to him (or her). They can agree to the price of 4,000 €. But irrespective of the price agreed, as a result of transaction, economic surplus increases by 2,000 €: in our example the seller receives 4,000 € for something he (or she) valued at 3,000 €, so he (or she) is left with 1,000 € more, and the buyer paid 4,000 € for something that is worth 5,000 €, so he (or she) enjoys the surplus of 1,000 €.

Yet in the second-hand market both good and bad cars are mixed. Buying a car in such circumstances is like taking a part in a lottery. With 50% probability you will get a good car, and with 50% probability you will get a bad one. By looking at the expected value of what you get, you are willing to pay at most 7,500 € for a car. Owners of bad cars (who were

willing to sell them for 3,000 €) will be happy to accept such a generous price. But owners of good cars will not agree to such a price, and they are likely to withdraw. This is what economists call "adverse selection" resulting from asymmetric information.

The market failure is caused by the fact that good cars should change owners as well. It is a social loss if a car stays with a person who values it at 8,000 € instead of being transferred to a person who values it at 10,000 €. There is a potential improvement of 2,000 €, but the market cannot realise this. The efficiency of the economy is impaired.

There are also examples when the seller knows less than the buyer. Insurance policy markets have to deal with such a problem. If you want to buy insurance against medical care expenditures, the seller may ask you whether any of your ancestors suffered from a certain disease. You may have such an information, but you do not have to disclose it. As a result, the company cannot estimate a fair price of your policy. It may be too high if the company is afraid that you may require an expensive treatment but you will not, or *vice versa*. It may be too low, if the company does not expect you to require an expensive treatment when in fact you will. Consequently insurance policy markets cannot function well.

Adverse selection is likely to result from asymmetric information. Its definition is simple: a better product is driven out from the market by a worse one; hidden information implies suboptimal demand or supply. Market failure is caused by the fact that when adverse selection makes certain transactions difficult, equilibrium will be achieved not in a Pareto optimum. In the case of used cars, good cars disappear (as explained above). In the case of insurance, low-risk groups disappear. People who are not likely to require expensive treatment will not buy a policy for the price which takes into account the average risk of treatment. As low-risk groups avoid the market, high-risk groups are overrepresented. This leads to higher prices of insurance policies, and adverse selection continues until low-risk groups are eliminated from the market completely.

Asymmetric information implies serious problems for economic efficiency. One way to proceed would be to give up efforts to overcome the problem: "it is a pity, but we cannot help it". Another way – a more constructive one – would be to look for a mechanism to improve the efficiency. Two such mechanisms in insurance market with asymmetric information on damages are frequently used:

- Efficiency can be improved by imposing obligatory insurance in order to bring in to the market low-risk groups,
- Low-risk groups can be also brought in by decentralized measures (i.e. without government intervention).

In either case the aim is to keep low-risk groups – otherwise disappearing (as a result of adverse selection) – in the insurance market. In the first case a government intervention is called for; in the second case, the government intervention is not necessary.

Civil liability insurance for car drivers is an example of the first approach. Without an obligation, safe drivers – who do not anticipate to cause an accident – would have an incentive not to buy insurance policy at the price which takes into account the presence of reckless drivers. As a result of adverse selection, reckless drivers would be even more

frequent buyers, and the price of the insurance policy would be even higher. Obligatory civil liability insurance keeps the low-risk group in the market.

If hidden information leads to suboptimal demand resulting from a loss imposed on somebody else (e.g. caused by the supply of a low quality commodity) then the government intervention can correct for the market failure, if it lowers the average level of this loss (*externality* in economic language). The problem of civil liability insurance fits this pattern. Safe drivers are victims of hidden information, because they cannot buy a policy at a reasonable price in the absence of the obligation.

A spontaneous (decentralised) approach can be explained by referring to health insurance plan offered to the University of Warsaw employees. A couple of years ago, I wanted to buy health insurance, but I learned that the price asked by insurance companies was much too high (unattractive for a person considered to be a member of the low-risk group). After some time I was informed that the University of Warsaw was approached by a well-known company who offered health insurance for a reasonable price. My first impression was that the firm miscalculated the price. I understood the mechanism when I learned that the condition for this attractive price was participation of the majority of the University of Warsaw employees.

University of Warsaw employs roughly 7,000 people, some of whom are healthy, and some may require costly medical treatment. Both high-risk and low-risk groups are represented. Perhaps in voluntary insurance schemes the former may be over-represented, but given the scale of the university, the adverse selection cannot be catastrophic.

Let us now come back to the second-hand car market with asymmetric information on the quality of a car. As explained before, adverse selection leads to a market failure. Yet efficiency can be improved spontaneously (e.g. by encouraging customers to buy cars offered with a guarantee). If buyers require a guarantee, and if guarantees become customary, then selling a car without a guarantee provides information about its quality; most likely it is a bad car.

In the beginning of this lecture acquiring information was considered expensive. Sometimes the cost of information is not that high. Once again, let me refer to my personal experience. In addition to guarantees, an inspection by a professional mechanic can reduce the asymmetry. Some time ago when I wanted to buy a beautiful Chevrolet station wagon for a very attractive price of \$400, an American colleague advised me to pay \$20 to a mechanic for an inspection. This proved to be the best investment in my life. The mechanic revealed that the engine required at least \$1,000 in order to fix defects that were likely to show up after driving 50-100 kilometres. My expenditure of \$20 saved me \$1,000. Thanks to this moderate expenditure, I acquired information that turned out to be very valuable.

There are many effective mechanisms to deal with problems that economists call "hidden information". But there is also "hidden action" that leads to adverse selection as well. The term is obvious, and it does not require an explanation. The most cumbersome aspect of it is the fact that economic agents may change their actions depending on circumstances that are difficult to observe and control. The idea of moral hazard has been defined in economics as the lack of *ex post* incentives to care for something that was *ex ante* assumed in the contract. Like hidden information, hidden action leads to suboptimal supply.

Let us refer to the car insurance market. Insurance against car theft can be bought. If the probability of a car being stolen is 2% per annum, then the owner of a car worth 10,000 € is willing to buy insurance policy for 200 €. If the car is not stolen then nothing happens. But if the car is stolen, the owner claims 10,000 €.

The likelihood of a car theft depends on whether the owner keeps keys in a secure place. It can be reduced if the car is equipped with alarm. It can be reduced if the driver always locks the door, never drives in unsafe districts, never leaves the car in an unguarded parking, etc. Insurance company probably informs the driver about these safety features, the driver promises – perhaps in good faith – to apply recommended measures. Nevertheless, once the car is insured he (or she) may change the behaviour. The 2% average probability of car theft is in fact lower or higher depending on how the driver behaves. This is the essence of "moral hazard". After signing a contract the behaviour can change.

Lawyers try to include in contract conditions precautions against "moral hazard". For instance, insurance against tourist accidents has two different prices: a lower one for non-skiers, and a higher one for skiers. A person who suffered an injury on a ski slope will not get a compensation if the insurance did not cover skiing activities. In this case it would be very difficult to hide one's activity (like it is difficult to pretend to be a non-smoker when buying health insurance). But there are many more difficult cases where the contract cannot predict all circumstances encouraging agents to undertake the "moral hazard".

As in the case of hidden information, hidden action may require government intervention sometimes. Yet solving the concern is often possible without such an intervention. The following statements conclude the problem.

- There is no market failure if actions can be perfectly controlled (e.g. smoking habits).
- If hidden action leads to suboptimal supply (as higher supply would give buyers an incentive for a moral hazard), then government intervention is typically un-purposeful, since the problem results from the cost of information, rather than from losses imposed on others (*externalities*).
- If behaviour cannot be observed, then efficiency in the insurance market requires that the insurance coverage is not complete.

The last statement can be explained by referring – once again – to the car insurance market. Some companies sell insurance policies against car-theft with not complete coverage. This means that if a € 10,000 car is stolen, the owner receives only 95% of its value, that is 9,500 €. The "missing" 500 € amount is considered an incentive against "moral hazard". The car owner who will lose 500 € if the car is stolen is expected to keep all the precautions included in the terms of the insurance.

A service added by some insurance companies – as mentioned by one of the students – to install a chip in a car to signal its location helps to find a stolen car. It does not provide an incentive for the car owner to behave carefully. I am afraid that – on the contrary – installing the chip may reduce the owner's motivation to undertake anti-theft measures.

The rest of the lecture is devoted to yet another aspect of asymmetric information, namely "signalling". Market failures are caused by the fact, that buyers and sellers have different information about the object of a transaction. In the second-hand car market this boils down to

the hypothetical question that a buyer may ask the seller: "Is the car a good one"? Let us assume that the seller answers: "Yes, of course". Theoretically this would solve the problem. But is such an answer a credible one? Signalling is about making information credible. The result is called a separating equilibrium.

The textbook model refers to the following story. A prospective employer would like to hire employees, and to pay them wages equal to their marginal productivity. It is known that these productivities are either a_2 or a_1 , with $a_2 > a_1$. The point is, however, that the employer cannot tell before the hiring who belongs to what group. The only knowledge that is available is about the entire population which consists of b people with the high productivity, and $1-b$ with the low productivity. The members of either group belong to it for ever; no education, no training can change it. The problem of the employer can be formalised in the following way (high *versus* low productivity is identified with being "smart" or "dull").

- Marginal productivity of workers (unobserved): dull a_1 , smart a_2 ; $a_2 > a_1$.
- Relative share in the population (observed): smart b , dull $(1-b)$.
- If the employer cannot tell who is smart and who is dull (but $\partial Q/\partial L_1 = a_1$ and $\partial Q/\partial L_2 = a_2$, where Q – production, L_1 – the employment of the dull, L_2 – the employment of the smart), then the wage offered should be the same for all and equal to $w = (1-b)a_1 + ba_2$.

The same wage offered to the smart and the dull prevents the employer from making losses (by paying – on average – more than an employee produces), but it does not solve the labour market problem. If the smart were to be paid the same as the dull ones, then adverse selection should be expected: the smart ones would disappear from the labour market. They can claim: "we are smart, not dull, and we deserve a higher wage", but this information will not be considered credible.

Let us assume, however, that some sort of formal education is possible. Everybody can go through it, but it is more expensive for the dull ones than for the smart ones. The higher cost per unit of education c_1 for the dull than c_2 for the smart can be interpreted in terms of out-of-pocket expenses and opportunities lost. In order to learn, the dull have to pay for extra instruction. They need to read the textbook three times, whereas the smart ones know everything having read the textbook only once. Consequently the smart ones can spend their time on earning money (e.g. by giving lessons) rather than spending money (e.g. by taking lessons). The smart ones go through the education cheaply. But also the dull ones can undertake the education. Even though it is expensive for them, they expect that they will be considered the smart ones and offered a higher wage a_2 rather than a_1 .

The employer can request a certificate of acquiring education at certain level. Let us denote this level e^* (the number of units of education, say, years of schooling). If it is set adequately, it can serve as a "signalling device": those who present it are likely to be smart; those who do not present it are likely to be dull. By separating the smart from the dull, the employer can offer different wages to both groups. The adverse selection process is avoided. The mechanism of choosing e^* can be formalised in the following way.

- The unit cost of acquiring education is c_1 for the dull and c_2 for the smart, and $c_2 < c_1$.
- Let the level of education e^* be selected so as $(a_2 - a_1)/c_1 < e^* < (a_2 - a_1)/c_2$.
- Then the dull choose $e_1 = 0$, and the smart ones $e_2 = e^*$,

since for the smart ones: the gain = $a_2 - a_1 > c_2 e^*$ = the cost,
while for the dull ones: the gain = $a_2 - a_1 < c_1 e^*$ = the cost.

An education certificate corresponding to e^* signals to the employer the quality of its prospective employee. It is possible that a dull one will bring a diploma, and a smart one will fail to present a diploma. Nevertheless the likelihood of such a misclassification is rather small (the dull ones are unlikely to spend on the certificate more than they can gain in the labour market, and the smart ones are unlikely to fail getting a certificate if they know that it plays a positive role in the labour market).

The last comment is on the adjective "separating". If an adjective is added to a noun, the noun usually changes its meaning badly (think of "chair" and "electric chair"). Welfare economics theorems let us think about "equilibria" positively. Signalling introduces the concept of a "separating equilibrium" which suggests that the latter is less positive than the former. The following efficiency analyses explain why economists prefer to speak about "separating equilibrium" rather than simply equilibrium implied by signalling.

Asymmetric information (and the resulting adverse selection) causes a market failure: economic efficiency is compromised. We can contrast two correcting mechanisms – guarantees and signalling. Selling a second-hand car with a guarantee does not imply unnecessary cost. If the car breaks down, it has to be repaired. Selling the car with or without a guarantee differs in who pays for the repair. It does not change the fact whether a repair is to be paid. Thus the guarantee does not imply unnecessary costs. On the contrary, signalling assumes that education is irrelevant for economic efficiency (a_1 and a_2 do not depend on whether somebody undertook the education or not). From this point of view, it is an unnecessary cost. It is born only to fight asymmetric information. That is why economists call it "separating equilibrium" rather than simply "equilibrium".

Questions and answers to lecture 8

8.1 How do insurance companies try to reduce the asymmetric information before selling a health insurance policy?

Insurance companies face asymmetric information. They know less than a potential buyer. The risk they have to cope with is that medical treatment required by the person insured is more expensive than what they expected. The first question they ask is about hereditary diseases and smoking habits. People can cheat, so the company typically requires professional blood and urine tests, and a visit to a doctor who works for the company. In addition, they exclude injuries caused by certain risky activities like skiing, parachuting, etc. If an insured person tries to cheat and wants the company to pay for a procedure that resulted from a forbidden activity, the company can check the cause of the injury and refuse to pay when they find that it was exempt from the insurance. They cannot reduce the asymmetry completely, but they try to keep it below what they can cope with.

8.2 How do employers try to reduce the asymmetric information resulting from hiring an employee?

In many cases hiring a new employee involves a complex procedure. The person should pass a number of tests in order to demonstrate his (or her) aptitude to play the roles the employer expects. This is a relatively easy part of the procedure. The more difficult part is to agree on a contract that reduces moral hazard. In addition, employers prefer to offer short-term contracts to be extended only when the employee is found to have the expected productivity. After several months the employer has a much better knowledge of the employee's aptitude than at the time of recruitment.

8.3 Please discuss the moral hazard problems faced by students signing education contracts.

An education contract lists rights and obligations of the university and students. The list can never be specific enough to predict all scenarios. A typical moral hazard is when the university promises to provide effective lectures, and in fact they turn out to be boring or outdated. It is impractical to characterise the lectures up to the last detail, so students sign contracts without a guarantee that they will receive what they expect and what the contract promises. It would be impractical to offer students a guarantee that lectures will result in high test scores, attractive employment, and so on. In other words, exposure to some moral hazard is inevitable.

Nevertheless students try to make sure that the university does not take the opportunity of moral hazard. They cannot guarantee that the university gets rid of lazy or incompetent lecturers immediately, but they prefer to deal with universities with good reputation. It is believed that in order to save reputation good universities have internal procedures of hiring, evaluation and promotion such that bad lecturers are not employed at all. In some drastic cases lecturers are fired when students prove that classes are run inadequately, but usually this simply does not happen. Moreover students trust that good universities offer curricula that provide them with skills appreciated in the labour market. On top of that they trust that the diploma of a good university will "signal" their quality, but this is a separate issue.

8.4 Auctions can be used in order to reduce asymmetric information. If the owner of an object wants to sell it then – instead of asking a specific price, he (or she) can organise an auction, and to sell the object to a buyer who offered the highest price. There are a number of auction types, and we will focus on so-called sealed bid auction (potential buyers send their offers in sealed envelopes, so that nobody knows at the time of making offers who bade what). The highest bidder wins the auction. But there are first-price auctions and second-price auctions. In the former the winner pays the own price offered; in the latter the winner pays the price offered by the first loser. The second-price auction is also called Vickrey auction since it was suggested by William Vickrey (1914-1996; he was awarded the Nobel prize 3 days before his death). Can you guess why economists prefer the Vickrey auction rather than the first-price auction?

On average both types of auctions provide the owner with the same revenues. This paradoxical result can be proved formally. Let us confine to two potential buyers. The object can be purchased by either of two buyers which value the object v_1 and v_2 , respectively. Each valuation is kept secret. But both the owner and the buyers know the distribution that these values originate from (as realisations of some random variable θ). The potential buyers state their bids s_1 and s_2 which are not necessarily equal to v_1 and v_2 , respectively. The object goes to the first bidder, if $s_1 > s_2$, or to the second bidder if $s_1 < s_2$.

How much does the winner pay (and how much does the owner receive)? Let us consider the first-price auction first. The winner pays s_i , and gets the net benefit $v_i - s_i$. If $s_i > v_i$ then the net benefit would be negative and the winner would have a motivation to be the loser rather than the winner (in other words, he does not have a motivation to overstate the valuation). If $s_i < v_i$ then the winner risks losing the auction and hence $s_i = v_i$ makes the preferred bid. On average, the payment is $E\theta$. Now let us consider the second-price auction. The motivation for the winner not to overstate the bid is like before. A likely motivation to understate in order to pay less disappears, because the winner pays the bid of the loser (rather than his/her own).

The reason economists prefer the Vickrey auctions is that they provide stronger incentives not to underestimate bids, because the winner does not have to pay his (or her) bid, but rather the loser's bid. Thus the motivation to underestimate one's own bid ("If I win the auction I will have to pay what I declared") disappears. One can argue that second-price auctions provide sellers with even higher revenues.

8.5 Can an unsuccessful Vickrey auction (see 8.4) result in adopting the first-price auction next time?

It should not. An average Dutch is taller than an average Greek. But there are tall Greeks, as well as there are short Dutch. The height of an individual from a population has some statistical distribution. The height of a specific person is a realisation of the variable with this distribution. Meeting a tall Greek and a short Dutch does not imply that an average Greek is taller than an average Dutch.

Auctions are statistical events. It may be the case that the winner bade high and the loser bade disappointingly low. The owner of the object for sale can be upset by the fact that he (or she) gets the second price instead of the first one. Nevertheless, rather than taking a decision based on one realisation, the owner should make sure that next auctions reduce the probability of collusion of auction participants, or any other circumstances resulting in bids far from the average.

9 – Nash equilibrium

Games are used in order to analyse human behaviour found in situations when nobody can fully control the outcome; what happens results from independent decisions of several entities. The following definition captures this idea.

Non-zero sum two-person game is a representation of a decision situation with a table of pairs of numbers (P_{ij}, D_{ij}) . Index $i=1, \dots, m$, where m is the number of strategies (decision variants) for the first player, and $j=1, \dots, n$, where n is the number of strategies (decision variants) for the second player. The numbers P_{ij} are payments to the first player, and D_{ij} – to the second player, if the first chose the i th strategy, and the second – the j th one. Payments can be interpreted as utilities. The table including pairs of numbers (P_{ij}, D_{ij}) is called payoff matrix.

The terminology has to be explained. The term 'two-person game' does not need any explanation. The definition can be easily generalized into n -person games. The notation becomes more complex, because simultaneous decisions of three or more players cannot be represented by a single matrix. The 'non-zero sum' expression allows for 'zero sum' games as

well. 'Non-zero sum' refers to the fact that $P_{ij}+D_{ij}$ is not necessarily zero. If incidentally $P_{ij}+D_{ij}=0$ (i.e. $P_{ij}=-D_{ij}$) the entire framework works, even though certain additional facts can be observed. The next definition introduces the concept of domination.

The strategy i_0 of the first player is called (strictly) dominant, if for any strategy i of the first player, and any strategy j of the second player, $P_{i_0j} > P_{ij}$; likewise, strategy i_0 is (strictly) dominated, if there exists strategy i such that for any strategy j , $P_{i_0j} < P_{ij}$ (analogously for strategies of the second player).

The following example explains these concepts.

		Second	
		L	R
First	U	(1,3)	(4,1)
	D	(0,2)	(3,4)

U is (strictly) dominant for the First. Hence D is not likely to be played. Therefore the game reduces to:

		Second	
		L	R
First	U	(1,3)	(4,1)

Now it turns out that R is (strictly) dominated for the Second player, and therefore it is not likely to be played. Thus the game is iteratively reduced to

		Second
		L
First	U	(1,3)

In other words, the solution is (U,L) the First player will get 1, and the Second will get 3. This is what will happen if the players behave rationally. Please note that the same result would be obtained if (strictly) dominated strategies of the Second player were eliminated first. If the domination is somewhat weaker (the next definition will explain this concept formally), the result may depend on the sequence.

The strategy i_0 of the first player is called weakly dominant, if for any strategy i of the first player, and any strategy j of the second player, $P_{i_0j} \geq P_{ij}$ with strict inequality for at least one strategy of the second player; likewise, strategy i_0 is weakly dominated, if there exists strategy i such that for any strategy j , $P_{i_0j} \leq P_{ij}$ with strict inequality for at least one strategy of the second player (analogously for strategies of the second player).

Let us look at the following example.

		Second	
		L	R
First	U	(5,1)	(4,0)
	M	(6,0)	(3,1)
	D	(6,4)	(4,4)

Strategies U and M of the first player are weakly dominated by D. If U is eliminated, the game reduces to:

		Second	
		L	R
First	M	(6,0)	(3,1)
	D	(6,4)	(4,4)

Now the Second player can eliminate L as weakly dominated by R. Then the game reduces to:

		Second
		R
First	M	(3,1)
	D	(4,4)

Finally the First player can eliminate M, and the game reduces to:

		Second
		R
First	D	(4,4)

However strategies can be eliminated in a different order. If M is eliminated first, the game reduces to:

		Second	
		L	R
First	U	(5,1)	(4,0)
	D	(6,4)	(4,4)

Now R becomes weakly dominated, and the game reduces to:

		Second
		L
First	U	(5,1)
	D	(6,4)

The First player may eliminate U now, and the game reduces to:

		Second
		L
First	D	(6,4)

Please note that this is a different outcome than in the previous elimination sequence. Weak dominance does not justify a reasonable iterative elimination procedure, since the choice of a specific sequence is arbitrary. In the first case we started with U, and in the second one we started with M.

We will define the most important concept in game theory, namely what we call now Nash Equilibrium – called also Nash strategy. John Nash (1928-2015; Nobel prize in 1994) invented a new type of equilibrium in 1950 (he did not call it "Nash Equilibrium"). A traditional equilibrium was understood as a situation when people do not have a motivation to move away. Nash changed this understanding slightly by assuming that nobody has a motivation to move away unilaterally. This slight change turned out to be extremely important, and forced economists to look at their discipline differently. The difference between moving away and moving away unilaterally is essential. Moving away unilaterally assumes, that other decision makers will keep their choices (and will not move). It precludes moving away jointly, even though this is what economists understood earlier when they argued that an equilibrium is synonymous with an optimum. Nash equilibrium does not have to be an optimum, as we will see soon.

Nash strategy (equilibrium) means any pair of strategies (i_0, j_0) such that $P_{i_0 j_0} = \max_i \{P_{ij_0}\}$ and $D_{i_0 j_0} = \max_j \{D_{i_0 j}\}$.

There is an easy theorem which links the concept of domination and Nash equilibrium. Please note, however, that the implication leads in one direction only (i.e. there may be Nash equilibria consisting of strategies that are not dominant).

Theorem (corollary of definitions)

If players have (either strictly or weakly) dominant strategies, then their pair is a Nash equilibrium.

Proof:

If there is i_0 such that $P_{i_0 j} \geq P_{ij}$ for any j , and there is j_0 such that $D_{ij_0} \geq D_{ij}$ for any i , then (i_0, j_0) is a Nash strategy. Please also note that if both inequalities are strict then the equilibrium is unique.

An alternative definition of Nash equilibrium reads: if players are in a Nash equilibrium, then – if they wish to maximize their payoffs – neither has a motivation to unilaterally change his or her strategy

This alternative definition of Nash equilibrium is sometimes adopted as the main one. Indeed it is very convenient to apply, as you will see in a moment. Nevertheless only one definition can be adopted formally, and anything else needs to be proved to be equivalent. That is why the criterion above is a theorem (even though its proof is a trivial one).

There exist games with no Nash equilibrium as the following example demonstrates:

		D	
		Y	N
P	Y	(4,2)	(-2,3)
	N	(2,1)	(-1,0)

Please also note that strategies making a Nash equilibrium (printed in bold) do not have to be dominant ones (nor unique):

		D	
		Y	N
P	Y	(4,2)	(-2,2)
	N	(2,1)	(-3,0)

Now we will look at the best known game called "Prisoner's Dilemma". The story behind is that two criminals are caught at a minor offence – like stealing a purse from an old lady in a supermarket. They were caught by the local security, everything is recorded, so they cannot deny. You go to jail for 1 month for such an offence. However this is not the only crime they committed. Sometime earlier they committed a heavier crime – something you go to jail (when convicted) for 1 year (12 months). But they were not caught then and nobody saw them, so unless one of them confesses, they cannot be convicted. They promised each other not to confess about what they did.

The police interrogates them in two separate rooms. The police do not have any doubts about their minor offence (everything is recorded and certain), but whenever criminals are arrested, they are interrogated about all previous unsolved mysteries. When the police asked them "did you do it", they could deny without any hesitation many times. But when they are asked about what they actually did, how will they react? They did it, but should they confess? They promised each other to deny. Yet would you trust the word of honour given by a criminal? In the class today you observed that you would not trust such a partner. Each of the thieves has a dilemma. "Indeed I promised my partner to deny, and he promised me the same. But shall I trust him? What if he breaks the promise and confesses?"

The numbers you see in the payoff matrix below refer to their predicament. "C" stands for "Confess", and "D" stands for "Deny". If both of them confess, they both go to jail for 12 months. If both of them deny then they go to jail for the minor offence they were caught at, and the major crime stays unsolved. If the first confesses but the second denies, then the first is freed and the second goes to jail for 18 months. The asymmetry is caused by the fact, that the police rewards the first criminal for his collaboration; both blames are forgiven (the major crime and the minor offence). The second is punished not only for the crime that he did (and the police can convict him now) but also for the fact that he lied to the police (by denying).

Example (Prisoner's Dilemma)

		Second	
		C	D
First	C	(-12,-12)	(0,-18)
	D	(-18,0)	(-1,-1)

In order to identify a Nash equilibrium, payoffs given in every cell need to be analysed. Let us start with the (D,D) cell. Both have a motivation to unilaterally move away. If the first player switches from D to C, then he will be given 0 rather than -1. If the second player switches from D to C, then he will be given 0 rather than -1. Hence this is not a Nash equilibrium. When (C,D) is checked, then the first player does not have a motivation to unilaterally switch (he would be given -1 rather than 0), but the second player by switching from D to C will be given -12 rather than -18. Likewise (D,C) is not a Nash equilibrium. Even though the second player does not have a motivation to unilaterally switch (he would be given -1 rather than 0),

the first player by switching from D to C will be given -12 rather than -18. The only outcome where neither of the players has a motivation to unilaterally move away (from C to D) is (C,C).

Nash equilibrium is thus in (C,C) which is the very worst outcome for both players. Nash equilibrium does not necessarily 'optimize' the global outcome (which – in this case – would be (D,D) giving the payoffs (-1,-1)). The 'Prisoner's Dilemma' demonstrates that Nash equilibrium is not an optimum, but in fact it can be the very worst outcome: 24 months in jail, *versus* 18 or 2 in other combinations.

That is why economists were shocked when John Nash introduced this strange equilibrium concept. Nevertheless the awareness grew gradually that this is a kind of equilibrium that was analysed by them earlier as well. For instance Antoine Augustin Cournot (1801–1877) studied oligopolistic competition and calculated what we now call 'Cournot equilibrium'. If oligopolists want to maximise their joint profit by charging the monopolistic price, they should calculate the joint supply corresponding to that price, and to allocate the total supply to participating firms. But having done this, every oligopolist would have a motivation to unilaterally cheat (by increasing its supply somewhat, spoiling the price a little bit, and enjoying a higher profit at the expense of others). If all of them do so, all will lose. Cournot calculated a solution such that nobody has a motivation to unilaterally move away (to cheat). This is a Nash equilibrium. It also demonstrates the difference between moving away jointly, and moving away unilaterally (that is assuming that others will not move).

Many social arrangements turn out to be Nash equilibria that are not optima. For instance, Windows operating system is used by almost all students, but it has a number of shortcomings (and besides, it is expensive). Some time ago I contested it using an alternative one. But my students complained that they could not read my electronic letters. And *vice versa* – I had problems with what my students sent me. Finally I gave up, I switched to Windows and I do not have a motivation to unilaterally move away. Likewise, other people do not have a motivation to unilaterally move away. If everybody switched to something else, all could have gained, but this requires coordination that elementary game theory does not study. Windows is a Nash equilibrium – in a sense that nobody has a motivation to unilaterally move away – even though perhaps it is not optimal.

Nash equilibrium has a formal mathematical definition, but it can be interpreted in an informal way too. In particular, it is interpreted as a 'stable social convention'. It is a social convention in a sense that people stick to it irrespective of whether it is optimal; they simply do not have incentives to unilaterally move away. It is also stable in a sense that it can last for many years.

Oligopolistic models analysed in microeconomics are popular examples of how people's behaviour can be explained by the game theory. An oligopoly is when the supply comes mainly from a small number of firms each of which has some impact on the market price. A duopoly is the simplest case of an oligopoly: the supply comes from two firms only. In an oligopolistic market, the price rule derived for a perfectly competitive market ($p = MC(y^*) = AC(y^*)$) may not hold. Prices depend on how oligopolists compete.

We will apply the following notation:

- y_i – supply from the i th rival
- y – aggregate supply
- p_i – price charged by the i th rival
- p – price determined by the market

Key features of oligopolistic competition relate to two questions: (1) do rivals make decisions simultaneously, or is there one of them (known as the leader) who takes a decision first while others (known as followers) make decisions afterwards; and (2) do rivals compete with prices or quantities.

Accordingly we have four models of oligopolistic competition:

1. Cournot: rivals make quantity decisions simultaneously
2. Bertrand: rivals make price decisions simultaneously
3. Quantity leadership (Stackelberg): one makes a quantity decision first
4. Price leadership: one makes a price decision first

The general idea of analysing oligopolistic markets is based on calculating key economic variables:

- Determine a demand that the market will reveal: $D(p)=y$
- Let $p_i y_i - TC(y_i)$ be the profit enjoyed by the i th firm (TC – total cost)
- Let p_i^* and y_i^* be profit maximizing decisions of the i th firm
- Firm i makes its decisions which maximize its profit, taking into account expected decisions of its rivals
- We look for an equilibrium in a sense that all firms make decisions exactly as they were expected by their rivals

In addition (for simplicity) we assume:

- The number of firms is two (i.e. $i=1$ or $i=2$)
- The demand curve is known and linear (i.e. $p=a-by$)
- Both firms have identical cost functions and $MC_1=MC_2=AC_1=AC_2=c=\text{const}$

By looking at alternative combinations of types of oligopolistic competition we can develop models 1-4 listed above. We start with the Cournot model.

Cournot duopoly model:

- The price is implied by the total supply: $p=a-b(y_1+y_2)$
- The first rival makes a quantity decision that maximizes its profit expecting that the second rival does the same
- Thus they solve two simultaneous maximization problems:
 - $\max_{y_1} \{(a-b(y_1+y_2))y_1 - cy_1\}$
 - $\max_{y_2} \{(a-b(y_1+y_2))y_2 - cy_2\}$

FOC for these problems are:

- $a - 2by_1 - by_2 - c = 0$
- $a - 2by_2 - by_1 - c = 0$

Solving these equations yields:

- $y_1 = y_2 = (a-c)/3b$
- $y = 2(a-c)/3b$
- $p = (a+2c)/3$

Please note that $p = (a+2c)/3$ is greater than competitive price under 'typical conditions', i.e. greater than c . To see this, note that $a > c$ (otherwise in the competitive market the demand and supply would not intersect). Thus $(a+2c)/3 > (c+2c)/3 = c$. Hence the Cournot price is higher than the competitive one.

Bertrand duopoly model

- Each rival makes its price decision p_i
- The one whose price is higher loses all the buyers
- The one whose price is lower wins all the buyers
- The lower price becomes the market price p
- The winner enjoys profit $(p-c)(a-p)/b$
- If $p_1 = p_2$ then $p = p_1 = p_2$, $y = (a-p)/b$, and $y_1 = y_2 = (a-p)/(2b)$

In equilibrium $p=c$, and profits vanish. The proof is simple. No price $p < c$ can be sustained since it implies losses. No price $p > c$ can be sustained since at least one firm has a motivation to offer a price $\in (c, p)$ – i.e. higher than c , and lower than p – in order to undercut the rival's price, get all the clients and increase the profit.

Instead of studying more complicated models of price or quantity leadership, we will ask a more practical question. When do duopolists compete *à la* Bertrand, and when do they compete *à la* Cournot? In the Bertrand model they compete by prices while in the Cournot model – by quantities. The textbook example of the former is the software market, while that of the latter: Boeing *versus* Airbus.

Let us start with software. Some time ago, computer programmes were sold on CDs. Burning a CD cost 1 eurocent or so, which was negligible. Now they are sold in the internet; you have to pay something and then you can download it to your computer. The producer had to upload it on the seller's server only once, and it does not matter, whether afterwards it is downloaded by users 1 thousand times or 1 million times. The cost does not change for the seller. One crucial decision to be taken by the seller is the price, while the quantity can be arbitrary; it will be determined by the market.

Boeing with its main production facility in Seattle is the most important American airplane supplier, and Airbus with its main production facility in Toulouse is the most important European airplane supplier. The two firms control the global market of wide body passenger jets. There are some other producers, but their market shares are much smaller. The global demand for large airplanes is very limited, maybe 500 per year. If Boeing and Airbus are to produce, say, 300 each per year, the market will not absorb such a supply at a satisfactory price. But the quantity decisions are irreversible. If the global demand turns out to be higher than expected, say, 700 per year, neither Boeing nor Airbus can increase their production rapidly, because investing in appropriate capacity takes time. If they built facilities to produce

less, they cannot produce more. Can they produce less than the capacity permits? Not really. The reason to decrease production was that the market did not accept the price the producer asked for. If the price is going to be higher (and it must be higher, because the cost of investing in capacity has been already spent – so-called sunk cost), the demand will be even lower. In other words, the quantity cannot be lowered without incurring financial losses.

The Bertrand and Cournot models differ in the way the producers compete with each other. In the former they set prices. Whoever sets a lower price wins all the clients. In the latter they decide on quantities supplied, and the market determines the price that both of them have to face. They can be surprised by this price, but it is too late to change their quantity decisions. In order to use either of the models, economists have to determine whether the duopolists compete by prices or by quantities. If quantities can be adjusted easily, Bertrand model is perhaps more adequate. Otherwise, the Cournot model is more suitable. There are also mixed models that require more sophisticated mathematical treatment.

Questions and answers to lecture 9

9.1 Can a three-person game be described by payoff tables?

Yes, but not by a single payoff table. An example is fairly complicated. Two players choose rows and columns, respectively, as in two-person games. However, they do not know which payoff matrix is applied. The third player chooses the payoff matrix. The game reads:

	L	R		L	R		L	R
U	0,1,3	0,0,0	U	2,2,2	0,0,0	U	0,1,0	0,0,0
D	1,1,1	1,0,0	D	2,2,0	2,2,2	D	1,1,0	1,0,3
	A			B			C	

First chooses U or D, Second chooses L or R, Third chooses the payoff matrix A, B, or C. This game has the unique Nash equilibrium (D,L,A) with payoffs 1 to each of the players. To see that (D,L,A) is a Nash equilibrium indeed, let us check that nobody has a motivation to move away unilaterally. If the First was to switch from D to U, the payoff would drop from 1 to 0. If the Second was to switch from L to R, the payoff would drop from 1 to 0. If the Third was to switch from A to B or C, the payoff would drop from 1 to 0 too. It is easy to check that there are no other Nash equilibria: if you look at the other 11 cells, you will always find that somebody has a motivation to move away unilaterally.

9.2 Can a zero sum game be considered a non-zero sum game?

		Second		
		1	...	n
First	1	$(x_{11}, -x_{11})$	...	$(x_{1n}, -x_{1n})$
	.	.	...	.
	.	.	...	.
	m	$(x_{m1}, -x_{m1})$	...	$(x_{mn}, -x_{mn})$

Yes. The payoff matrix has the above form; payments received by the first player and by the second player have the same absolute value, but opposite signs.

9.3 Can a game have several Nash equilibria?

Yes. Please note that both solutions obtained by eliminating weakly dominated strategies in the class example – (6,4) and (4,4) – are Nash equilibria. If (D,L) is chosen, then the First will get 6 or 5 if he (or she) switches to M or U. The Second will get 4 if he (or she) switches to R. Similarly, if (D,R) is chosen, then the First will get 3 or 4 if he (or she) switches to M or U. The Second will get 4 if he (or she) switches to L. Hence both (6,4) and (4,4) are Nash equilibria.

		Second	
		L	R
First	U	(5,1)	(4,0)
	M	(6,0)	(3,1)
	D	(6,4)	(4,4)

9.4 In prisoner's dilemma, why do not they take the joint decision to deny? If they deny, then the police cannot convict them.

They cannot agree on a decision, because they are being interrogated in two separate rooms. The First player does not know the decision taken by the Second one and *vice versa*. Professional interrogators do not reveal such an information. Even if they know that one confessed or denied they would not tell this to the other one. Prisoners promised each other to deny, but they do not know if this is 100% sure (please note that people who promised are criminals after all).

9.5 Martin Shubik (an excellent economist and one of the founders of game theory) said that "you need to know John Nash in order to understand the Nash equilibrium". Please comment on this opinion.

This is not a nice opinion. Nash equilibrium is about taking decisions for your own sake. You check whether something can improve your situation, and you do not ask what would be the result for anybody else. In other words, Nash equilibrium is about egoistic decision making. You do not try to achieve a situation which can be considered an optimum for a wider group (the society at large, or even a small group consisting of yourself and the other player). According to Martin Shubik, John Nash took decisions with his own wellbeing in mind only. Have you seen the movie 'Beautiful mind' (nobody answered this question positively in the last class)?

9.6 Why is the Cournot equilibrium considered a Nash equilibrium?

In the Cournot model, sellers decide about the supply they are going to offer. They would like to maximise their profits. They understand that the price will depend on the total supply available in the market. They understand that the more they offer, the lower the price. Thus – in order to maximise the profit – they would like to sell more, but not too much in order to enjoy a decent price. They cannot control what others will supply, but they assume that the others have similar objectives. Everybody makes certain assumptions regarding the supply offered by the others and adjusts their own supply in order to maximise the profit (please keep

in mind that the price depends on the supply). If the decisions actually taken are consistent with what oligopolists assumed, they do not have a motivation to change these decisions. This is what Cournot supposed. But this is exactly what the Nash equilibrium states: nobody has a motivation to unilaterally move away. If the oligopolists colluded, they could have taken the advantage of the monopolistic price. But in the Cournot model they do not collude and they take decisions independently.

9.7 Please discuss whether the right hand side traffic can be interpreted as a Nash equilibrium.

This is a social convention. In some countries cars drive on the left and in some countries on the right. Both solutions are good, but once something has been agreed upon people should obey the rules they adopted. Actually some people say that the left hand side traffic – prevalent in Europe before Napoleon – is a more natural one (given the fact that most of us are right-handed, if you rode your horse on the left you could have your sword kept in the right hand ready to hit an enemy who approached you riding on the left too). Nevertheless, let us assume that we have the right hand side traffic as a rule, and in principle everybody drives on the right. Does anybody have an incentive to switch to the left? The price to be paid for such a diversion would be very high. Therefore, even though driving on the left is equally good (some people say that it is even better), we do not have a motivation to change unilaterally. That is why driving on the right can be interpreted as a Nash equilibrium.

9.8 Can the QWERTY keyboard be considered 'a stable social convention'?

Yes. In many countries we use the QWERTY keyboard (the name comes from the first letter keys in the upper left of your keyboard). Not everybody applies the same standard. For instance in French speaking countries they use the QWERTZ keyboard. So when I asked a colleague in her office in Brussels to use her computer, when I tried to type my name, instead of Zylicz I saw Yzlicy on the computer screen. If you are used to a QWERTY keyboard, then you lose a lot of time when you type on another one.

Likewise, some of the visitors who come to my office have to spend extra effort if they wish to use my computer. QWERTY keyboard was adopted some time ago in order to allow easy typewriting, and soon it became a standard. In the middle of the 20th century studies were carried out in order to optimise keys' positions. It turned out that QWERTY keyboard was not a good solution (at least not good for English language users). There were attempts to introduce these better keyboards, but they were abandoned. The reason behind this was that the cost to be paid when switching to a better standard would be excessive. Please imagine a situation that the standard is QWERTY, and you want to unilaterally abandon it (even if there are thousands of users of the new keyboard, there are millions of users of the old standard). If everybody left the old standard, and moved to the new one, the situation would be different.

QWERTY keyboard is a social convention. It proved to be stable, because attempts to move to something different (better) failed. They did, because the cost of a unilateral – that is not universal – change would be excessive. This is what we check when we look for Nash equilibria. QWERTY keyboard should be considered a sort of a Nash equilibrium.

10 – Topics in game theory

The example quoted on page 68 of the previous lecture (QF-9) demonstrates that there may be games that do not have Nash equilibria. In order to avoid such problems, game theorists defined the concept of a so-called mixed strategy. Strategies chosen by players resemble bundles chosen by consumers. Some bundles are preferred to others if they satisfy consumer's needs better. Likewise, some strategies are preferred to others if they provide players with higher payoffs. In the fifth lecture (QF-5) we observed that preferences with respect to lotteries take into account the fact that certain outcomes happen with certain probabilities. Also in games we can think of a probability that a player applies a given strategy. The definitions below formalise this idea.

From now on, strategies defined previously are called *pure*. A game can be defined where pure strategies are selected by players randomly with certain probabilities $\pi=(\pi_1,\dots,\pi_m)$ and $\delta=(\delta_1,\dots,\delta_n)$, respectively (for the first and the second player), where $\pi_1,\dots,\pi_m\geq 0$, $\pi_1+\dots+\pi_m=1$ and $\delta_1,\dots,\delta_n\geq 0$, $\delta_1+\dots+\delta_n=1$. The pair (π,δ) is called a mixed strategy selection. If the players select mixed strategies, then the payoffs are understood as expected payoffs from their pure strategies. In other words, the payoff for the first is $\sum_{ij}\pi_i\delta_jP_{ij}$, and for the second it is $\sum_{ij}\pi_i\delta_jD_{ij}$.

Once again let us compare the concept of mixed strategies to lotteries. Lotteries with probabilities equal to 1 were called 'degenerate' ones. Formally they are lotteries, but their outcomes are certain rather than random. A 'traditional' game (with pure strategies) can be interpreted as a mixed-strategy game where probabilities are either 0 or 1 (they are 'degenerate'). In other words, everything we learnt about 'traditional' games applies to 'degenerate' games with mixed strategies directly. For non-degenerate games we need to introduce new definitions. In particular the Nash equilibrium can be defined. Please check that it will coincide with earlier definitions for 'degenerate' games. A pair of mixed strategies (π^0,δ^0) is a Nash equilibrium, if

- $\sum_{ij}\pi_i^0\delta_j^0P_{ij}=\max_{\pi}\{\sum_{ij}\pi_i\delta_j^0P_{ij}\}$, and
- $\sum_{ij}\pi_i^0\delta_j^0D_{ij}=\max_{\delta}\{\sum_{ij}\pi_i^0\delta_jD_{ij}\}$.

It has the same economic interpretation like the Nash equilibrium in 'traditional' games: it is a pair of strategies that nobody has a motivation to unilaterally move away. "Moving away" in this case means changing probabilities unilaterally, that is assuming that the other player sticks to the same probabilities as before. Unlike games with pure strategies, games with mixed strategies have always Nash equilibria. The proof can be derived from the Brouwer's fixed-point theorem, but it is difficult, and it will not be provided here. Instead we will check that the game from page 68 has a Nash equilibrium in mixed strategies. The game has the following payoff matrix:

		M	
		Y	N
F	Y	(4,2)	(-2,3)
	N	(2,1)	(-1,0)

If p is the probability of choosing Y by the player F, and q is the probability of choosing Y by the player M, then $(1/2, 1/2; 1/3, 2/3)$ is the Nash equilibrium in mixed strategies (when $p=1/2$, and $q=1/3$ the players do not have incentives to unilaterally change these probabilities).

You can check that if the player F switches from probabilities (1/2,1/2) while M sticks to (1/3,2/3), or if M switches from (1/3,2/3) while F sticks to (1/2,1/2) their average payoffs will not increase. In what follows we will see how these numbers can be calculated.

The expected payoff for F is

$$E(P)=4pq-2p(1-q)+2(1-p)q-(1-p)(1-q)=4pq-2p+2pq+2q-2pq-1+q+p-pq=-p+3q+3pq+1.$$

The expected payoff for M is

$$E(D)=2pq+3p(1-q)+(1-p)q=2pq+3p-3pq+q-pq=3p+q-2pq.$$

F would like to maximise $E(P)$, and M would like to maximise $E(D)$, therefore we need to solve system of equations

$$\partial E(P)/\partial p=0$$

$$\partial E(D)/\partial q=0$$

that is

$$-1+3q=0; q=1/3$$

$$1-2p=0; p=1/2.$$

This method has to be used carefully. Please note that vanishing derivatives can indicate a minimum, a maximum or an inflexion point of a function. Higher derivatives need to be calculated in order to make sure that we found what we looked for. In this particular case the second derivatives vanish as well, so determining what we found is difficult.

In some game theory textbooks F stands for a female, M for a male, Y stands for "yes, I am loyal to my spouse", and N for "no, I am not loyal to my spouse". If payoffs resemble benefits from being loyal or not loyal, the probabilities calculated as Nash equilibrium explain why men and women may be characterised by different loyalty patterns.

In real life there are more complicated decision situations than those analysed in the previous lecture. "Entry deterrence" game provides an example of a struggle between two firms that want to serve the same market. Originally the market was served by one firm – a monopolist – called Incumbent (I). Another firm (F) – called Entrant – contemplates entering the market. It has two strategies: "Enter" (E) and "Do Not Enter" (D). If it chooses D, the payoffs are (0,4) (its payoff is 0, and the payoff of the Incumbent is 4 as before). If it chooses E, the two firms will play a duopolistic game which boils down either to attacking (A) e.g. by a price war or to cooperating (C), e.g. *à la* Cournot or Stackelberg (or by sharing the market if they make a monopolistic collusion). If both attack then the payoffs are (-3,-1), if they cooperate then the payoffs are (3,1). If I attacks and F tries to cooperate the payoff is (-2,-1), and if I tries to cooperate and F attacks it is (1,-2). The story is reflected by the following payoff matrix:

		Incumbent (I)	
		C if F enters	A if F enters
Entering Firm (F)	D&C	0,4	0,4
	D&A	0,4	0,4
	E&C	3,1	-2,-1
	E&A	1,-2	-3,-1

This game has three Nash equilibria (D&C,A), (D&A,A), and (E&C,C) with payoffs (0,4), (0,4), and (3,1), respectively (printed in bold). In order to check that (0,4) is a Nash equilibrium it is sufficient to observe that if F changes its decision (either D&C, or D&A) while I sticks to its decision (A), the payoff for F will be -2, or -1; at the same time if I moves

from A to C unilaterally, then its payoff will not change. Hence neither of the firms has an incentive to move away unilaterally. Also (3,1) is a Nash equilibrium. To see this, please note that if F moves away unilaterally then its payoff goes down from 3 to 1 or 0. If I moves away unilaterally, then its payoff goes down from 1 to -1.

But in fact the game can be played in two stages. In the first stage I does not have to take any decision. F takes a decision whether to enter. If F took a decision not to enter, then there is no second stage; things remain as they were. If F took a decision to enter, then the second stage needs to be played. In this stage both firms have to decide whether to fight or not: whether to attack (A) the rival or to cooperate (C).

Formally these two decisions could have been taken in the first stage already. The entrant, while taking the decision whether to enter or not, could have stated, that – if the second stage takes place – it will attack or cooperate. But this statement was not really necessary if the decision was D (do not enter). Please note that the payoffs to be received by F would have been the same (0). Likewise the incumbent could have stated, that – if the second stage takes place – it will attack or cooperate. What could be the purpose of such a statement? If the statement is "A" ("I will attack"), it is a threat. But is such a threat credible? If you hear "do not do it, because if you do it, I will kill you", are you afraid? You may consider this threat not credible. Somebody who makes such threats is not obliged to fulfil them afterwards.

This is what may happen in the second stage. If the second stage is considered alone, the payoff matrix will be:

		Incumbent (I)	
		C	A
Entering Firm (F)	C	3,1	-2,-1
	A	1,-2	-3,-1

The only Nash equilibrium in this game is (C,C). We see that there is a difference between (3,1) and (0,4). The former shows up in the second stage, while the strategy A ("attack") of the incumbent does not show up in (C,C). The potential entrant may suspect that when confronted with its presence in the market, the incumbent will cooperate despite threats. Hence such an entry deterrence strategy of the incumbent is not effective.

To be credible, a threat must include a decision that will be taken irrespectively of what happened in the previous stage. This is called a pre-commitment. Examples of pre-commitment include: (in military planning) automatic destruction of one's own infrastructure in the case of attack, and (in economics) lack of possibility of changing a business decision before a certain date.

Let me elaborate on pre-commitments starting with the military example.

In Switzerland everything is undermined. There are explosives under every tunnel, every highway, and every bridge. This should deter from attacking the country. Any aggressor knows that if attacked, the Swiss will damage all the country's infrastructure. They will hurt themselves, but the aggressor will be hit as well. Thus attacking Switzerland does not make sense.

Is this threat credible? This is questionable. It is doubtful whether the Swiss – when attacked – will decide to hurt themselves. They would like to prevent the attack, of course, but if it happened, then it would be better to cooperate with the aggressor rather than to take desperate steps. Hence the threat can be considered not credible. In order to be credible, the threat should be based on a pre-commitment. If the action to explode everything is to be taken automatically rather than following a decision, the threat must be considered credible.

Now the economic example. Let us assume that there are negotiations between a seller and a buyer. The seller offers something for the price of 80€ per piece, and the potential buyer says that 60€ is the maximum acceptable price. The seller warns that unless the buyer pays 80€ today, there will be no opportunity to buy over the next 30 days. Is this threat credible (assuming that the buyer needs to buy it before)? It depends on whether it is backed by a pre-commitment. The potential buyer may consider the threat to be a negotiation trick: "the seller insists on 80€, but if I do not yield to this pressure, soon the price will go down somewhat".

Yet the seller can demonstrate a decision of the company's supervisory board which authorises the sale for the price not lower than 80€. If the price is to be lowered, then a new decision of the supervisory board needs to be taken. But according to the company's by-laws, the supervisory board meetings have to be announced at least 30 days ahead. Therefore, indeed, there is no possibility of lowering the price over the next 30 days. The threat is credible.

Questions and answers to lecture 10

10.1 Why can pure strategies be interpreted as a special case of mixed strategies?

Because pure strategies can be considered mixed with a degenerate probability distribution. Namely a pure strategy i can be considered mixed with probabilities $(0, \dots, 0, 1, 0, \dots, 0)$, where strategies $1, \dots, i-1$, and $i+1, \dots, m$ are never applied (that is they are applied with the probability 0), and the i th strategy is applied for sure (that is with the probability 1).

10.2 Please try to calculate the Nash equilibrium in mixed strategies for the following game

		M	
		C	B
F	C	(3,1)	(0,0)
	B	(0,0)	(1,3)

This game is called "the battle of sexes" (but it has nothing to do with the game analysed in the class earlier). The story is that two friends – a girl (F), and a boy (M) – would like to go out. The girl prefers a concert (C), and the boy prefers a baseball game (B). The girl enjoys the baseball game somewhat, and the boy enjoys the concert, but this is not what they prefer. Yet the worst outcome is when the girl insists on a concert, and the boy insists on the ball. Then they do not go together anywhere. The girl would like the boy to accompany her for the musical performance, and the boy wants the girl to accompany him for the sport event. It is easy to check that there are two Nash equilibria in pure strategies: (C,C), and (B,B). You can suspect that if they agree sometimes to what they do not prefer, both can be satisfied to some

extent. If they choose randomly between C and B with 50% probabilities they will get 1 each (on average). If they choose either of the Nash equilibria all the time (that is they avoid (C,B), and (B,C)) they will get 2 each (on average). What will happen if they switch from C to B randomly, with different probabilities? Let us try to calculate the Nash equilibrium in mixed strategies.

By applying the method explained in the class, we can calculate that $E(F) = 3pq + (1-p)(1-q) = 3pq + 1 - p - q + pq = 4pq - p - q + 1$, and $E(M) = pq + 3(1-p)(1-q) = pq + 3 - 3p - 3q + pq = 4pq - 3p - 3q + 3$. By taking partial derivatives and equating them to zero, we get the system of equations

$$4q - 1 = 0, \text{ and } 4p - 3 = 0 \text{ which yields } p = 3/4, \text{ and } q = 1/4.$$

In other words, the girl should randomly choose C in 75% cases, and the boy should randomly choose C in 25% of cases. The solution has an intuitive appeal, because – after all – the girl prefers the concert, and the boy prefers the baseball. But if they stick to such mixed strategies they will achieve even less (3/4 each on average) than by switching from C to B with 50% probabilities.

If there is a Nash equilibrium in pure strategies then calculating it in mixed strategies does not make sense (even without a Nash equilibrium in pure strategies a caution is necessary, as explained earlier). Let us recall that the game has two Nash equilibria in pure strategies. If the players stick to these equilibria they can get the payoff of 4 jointly. One Nash equilibrium (C,C) favours the girl, and the other one (B,B) favours the boy. An obvious way to introduce a sort of symmetry is to toss a coin. If it is Heads, then they choose (C,C), if it is Tails, then they choose (B,B).

10.3 What kind of a 'stable social convention' does the game on page 76 allude to?

The Nash equilibrium in mixed strategies was calculated as (1/2, 1/2; 1/3, 2/3). It means that women are loyal in 50% cases, and men in 33% cases only. If a woman is found not to be loyal, then many people are outraged. They are less outraged if a man is found not to be loyal. The social convention which supports such opinions has been stable over the last centuries (if not millennia).

10.4 What will happen in the 'entry deterrence' game, if the oligopolistic competition in its second stage is *à la* Bertrand?

In the Bertrand model oligopolists compete by prices (not by quantities, like in the Cournot model). In the Bertrand model the only equilibrium price is the same as in the competitive markets: if firms have identical linear cost curves, the price is equal to the marginal cost. In the long run their profits must vanish. This will change payoffs in the second stage and perhaps will motivate the entrant not to enter.

10.5 Can rationality be reconciled with pre-commitments?

It depends on how the rationality is understood. If it is understood narrowly, as taking decisions justified by what is most attractive in a given predicament, pre-commitments can be harmful. Pre-commitments do not allow making choices based on comparing the payoffs in a given stage; game theorists say that decisions are 'renegotiation proof'. Choices are made at earlier stages. Yet they may influence choices to be made by other players (the player who made a pre-commitment cannot choose any more). Pre-commitments can be rational when understood in a wider perspective.

11 – Industrial Organisation

Industrial Organisation (IO) analyses the strategic behaviour of firms, regulatory policy, and market competition. To some extent it overlaps with game theory, but it is more general. While game theory looks at incentives economic agents may have to take certain decisions, IO studies their behaviour in a broader context. In particular, IO demonstrates that Nash equilibrium does not explain economic behaviour in many circumstances. Before we elaborate on the Discrete Hotelling Model, let us analyse Network markets – as an example of an 'atypical' behaviour. The following four examples will be referred to:

- Compatibility and product standards
- External effects of consumption
- *Switching costs* and loyalty programmes
- Production scale effects

Young people are probably not aware of the fact that recording amateur videos and viewing movies – now enjoyed on mobile phones mainly – required a fairly sophisticated analogue technology in the 1970s and 1980s. There were two major incompatible formatting standards: VHS and Beta. Beta offered somewhat better quality than VHS, but VHS offered a possibility of longer recordings. VHS was developed by Panasonic, and Beta by Sony. Both firms lobbied governments to set their technologies as obligatory standards. The Japanese government – approached by them at the beginning – refused to intervene and the firms started to fight in the market. VHS eventually won the war, and emerged as the dominant home video format throughout the tape media period (i.e. until the 1990s). Only one format could survive, because consumers preferred to stick to what allowed them to record and to play whatever was available in the majority of cases (if the other format was used by a minority of consumers it was doomed to elimination; most people did not want to buy two types of equipment). Product standards imply that the "winner takes all".

Not all people who use Google navigation systems are aware of the fact that whenever they use the application, they benefit other users by letting the system know their location and speed. At the same time, they themselves benefit from information provided by other users (who let the system know their location and speed). In other words, the more people use the application, the more accurate estimations of traffic conditions, as observed in the class today. This is the prime example of "external effects of consumption": every consumer of a good increases the utility of using the good by others.

Consumers of mobile phone services benefit from the so-called "number portability". Several years ago they started to benefit from keeping the old number when they switch from one company to another. Earlier mobile phone companies claimed that – for technical reasons – when a customer switched to another company, he/she should abandon the old number. For me this would not be a big problem – just to inform my family and friends about the new number. Yet for business customers this would imply a significant cost (like printing new flyers and advertisements). Mobile phone companies knew that this would keep the old clients, even if there were cheaper options available somewhere else. As a result, prices they charged were excessive, because they knew that their customers are reluctant to leave even

when they learn that other companies charge less. Now – with the "number portability" – customers are less "loyal" and prices of mobile phone services went down significantly.

There are also a number of other methods of keeping the clients even if cheaper options are available elsewhere. Many companies run "point" programmes. They offer something (typically of a very low value) if buyers are loyal. The advantage of such "point" programmes is not spectacular, but people derive satisfaction from earning something "for free". This casts doubts on the rationality of decisions, but this is how people choose (and hence economists should acknowledge).

Production scale effects are probably the best understandable methods of winning customers. Chain hotels provide a good example. Because of the scale of operations, chain hotels can offer attractive prices sometimes. People choose their offers, despite the fact, that non-chain hotels often offer higher quality services. Nevertheless – because of their smaller scale – they cannot offer equally attractive prices.

In all examples above, "network markets" provide incentives for firms and consumers to make decisions that would be difficult to explain by referring to standard rules of economic analysis. In particular we will see in the "Discrete Hotelling Model" that Nash equilibrium is inadequate to predict their behaviour.

Discrete Hotelling model

1. Firms a and b produce a discernible product. There are n_a consumers who prefer a, and n_b who prefer b. Production costs are zero.
2. Consumer demand functions are unitary, and the disutility from consuming the non-preferred product is $\delta > 0$.
3. The utility of a type a consumer is:
 - $U_a = -p_a$ when buying from the supplier a
 - $U_a = -p_b - \delta$ when buying from the supplier b
4. The utility of a type b consumer is:
 - $U_b = -p_b$ when buying from the supplier b
 - $U_b = -p_a - \delta$ when buying from the supplier a
5. Therefore the numbers q_a and q_b of consumers who buy from a and b are, respectively:
 - $q_a = 0$, if $p_a > p_b + \delta$,
 - $q_a = n_a$, if $p_b - \delta \leq p_a \leq p_b + \delta$,
 - $q_a = n_a + n_b$, if $p_a < p_b - \delta$,
 - $q_b = 0$, if $p_b > p_a + \delta$,
 - $q_b = n_b$, if $p_a - \delta \leq p_b \leq p_a + \delta$,
 - $q_b = n_a + n_b$, if $p_b < p_a - \delta$

The best way to understand the "Discrete Hotelling Model" is to think of "a" as Coca Cola, and "b" as Pepsi Cola. Both of them are soft drinks and they are close substitutes. Nevertheless there are people who are Coca Cola fans, and people who are Pepsi Cola fans. They can taste the difference, and they prefer one over the other. However, if the price of their preferred brand is much higher than the price of the substitute, they will switch to it.

The assumption about the zero production cost is not crucial (it is adopted for arithmetic simplicity). It can be relaxed by adding a fixed term to the price (see question 11.5). Likewise

the assumptions 3 and 4 are not crucial. They can be relaxed by adding a fixed term to the utilities defined in 3 and 4, so that they become positive. The assumption about the unitary demand is for arithmetic simplicity. It says that irrespective of price and income circumstances, the demand for a and b (jointly) is the same.

The six cases listed in 5 can be justified by the common-sense calculations. Let us start with the first bullet. The price of a is so high, that even the Coca Cola fans are better off if they drink Pepsi Cola. Their utility is $-p_b - \delta$ which is more than $-p_a$. Thus nobody drinks Coca Cola: $q_a = 0$, and $q_b = n_a + n_b$ (please note that the inequality in the sixth bullet is the same as in the first one). In the second (and the fifth) bullet everybody drinks what they prefer: the prices of their not preferred substitutes are not attractive enough. Thus $q_a = n_a$ and $q_b = n_b$. The third (and the fourth) bullet replicates the sixth (and the first) bullet, except that the roles of Coca Cola and Pepsi Cola are reversed. Nobody drinks Pepsi Cola, since its price is excessive even for its fans. Hence $q_a = n_a + n_b$ and $q_b = 0$.

It turns out that that the model does not fit the usual game theoretic framework as the following theorem explains.

Theorem

There is no Nash equilibrium in the discrete Hotelling model.

Proof:

On the contrary, let us assume that (p_a, p_b) is a Nash equilibrium. Then one of the following three conditions holds:

1. $|p_a - p_b| > \delta$,
2. $|p_a - p_b| < \delta$,
3. $|p_a - p_b| = \delta$.

In each of these cases, one of the suppliers has a motivation to increase its profit by changing the price somewhat. Thus – contrary to our supposition – this cannot be a Nash equilibrium. It can be assumed that $p_a \geq p_b$ (if $p_a \leq p$ then their roles could be reversed). We will check what happens in all three cases. If 1 takes place, then $p_a - p_b > \delta$, that is $p_a > \delta + p_b$. This means that the price of Coca Cola is excessive even for the Coca Cola fans. The producer of Pepsi Cola can increase its price by a little bit (say, by $\frac{1}{2}$ of the number $p_a - p_b - \delta$). The price of the Coca Cola will be still excessive, and the profit of Pepsi Cola can increase. Let us look at the second case now. If 2 takes place, then $p_a - p_b < \delta$, that is $p_a < \delta + p_b$. By the assumption, $p_b \leq p_a$. If δ is subtracted from the left hand side, the inequality will read: $p_b - \delta < p_a$. This means that the second bullet of 5 takes place. Fans of Coca Cola drink Coca Cola, and fans of Pepsi Cola drink Pepsi Cola. The Pepsi Cola is cheaper than Coca Cola by less than δ . Thus the price of Pepsi Cola can be increased by a little bit (so that it is still attractive for its fans; it is still cheaper). It will not lose its fans, and it will increase its profit. If 3 takes place ($p_a = \delta + p_b$) the second bullet of 5 takes place too. Fans of Coca Cola drink Coca Cola, and fans of Pepsi Cola drink Pepsi Cola. The Pepsi Cola is cheaper than Coca Cola by δ . If Pepsi Cola decreases its price by a little bit, it is still acceptable for Pepsi Cola fans certainly, but all the Coca Cola fans start to drink Pepsi Cola, since the difference in prices started to be higher than δ . The number of clients will increase, and the profit of Pepsi Cola will increase as well.

The third case above illustrates what is meant by "price undercutting". This term means that by lowering its price one can increase its profit by winning (some of) the rival's clients. The formal definition below looks complicated, but analysing carefully what happens with the rivals' profits when prices change explains the idea.

Definition (*Undercut-proof prices*)

Supplier a undercuts b if $p_a < p_b - \delta$ (i.e. deprives it of potential buyers).

Prices (p_a^U, p_b^U) are undercut-proof, if:

1. p_a^U is the maximum price that for given p_b^U and q_b^U satisfies:

$$\Pi_b^U = p_b^U q_b^U \geq (p_b^U - \delta)(n_a + n_b)$$
2. p_b^U is the maximum price that for given p_a^U and q_a^U satisfies:

$$\Pi_a^U = p_a^U q_a^U \geq (p_a^U - \delta)(n_a + n_b)$$

By calculating the number of clients if prices are given by the following formulae, one can prove the following theorem.

Theorem (*Undercut-proof equilibrium*)

The following pair makes an undercut-proof equilibrium:

- $p_a = \delta(n_a + n_b)(n_a + 2n_b) / ((n_a)^2 + n_a n_b + (n_b)^2)$,
- $p_b = \delta(n_a + n_b)(n_b + 2n_a) / ((n_a)^2 + n_a n_b + (n_b)^2)$

We will avoid making calculations in the general case. Instead, we will check that the pair of prices identified in a special case when the number of Coca Cola and Pepsi Cola fans are equal, are undercut-proof indeed.

Corollary

If $n_a = n_b$, then $p_a^U = p_b^U = 2\delta$.

The formula $p_a^U = p_b^U = 2\delta$ is implied by the theorem. Let us check what happens if the prices are set according to it. Let $n_a = n_b = x$. The profits made by a and b are the same and they are equal to $2\delta x$. If a increases its price even slightly – say, by ε – then b can decrease its price by $\delta - \varepsilon/2$ undercutting the rival. The new prices will be $p_a = 2\delta + \varepsilon$ and $p_b = \delta + \varepsilon/2$. This is the case corresponding to the first bullet of item 5 in the definition of the Discrete Hotelling Model on page 2: $q_a = 0$ and $q_b = n_a + n_b = 2x$. In other words, if a increases its price even slightly, b can undercut it. If a decreases its price slightly (i.e. less than δ) then no clients of b will switch to a, and the only consequence will be its profit decrease. If a decreases its price significantly (i.e. by more than δ) then all the clients of b will switch, but the profit will go down (it will be less than $2\delta x$). If a decreases its price exactly by δ (from 2δ to δ), then the number of clients will double, but the profit will not increase. Thus a does not have an incentive to change the price. By symmetry, b does not have an incentive to change its price either.

Industrial Organisation is an approach to analyse economic questions without making assumptions adopted in Game Theory. In particular, it is possible to look for equilibria that allow for simultaneous changes undertaken by economic agents, not only for unilateral ones (that is assuming that others will not change anything). The Discrete Hotelling model introduces the concept of Undercut-Proof Price Equilibrium which analyses explicitly what an agent does, anticipating that his or her change may motivate the rival to do something different than so far.

Questions and answers to lecture 11

11.1 Please argue that the co-existence of two incompatible standards implies problems for users.

Co-existence of two incompatible standards is inconvenient. It requires the users to invest in (perhaps expensive) equipment to use any of the two standards. For instance, it is possible to run a railway system using two different gauge widths, like in Europe: on the Polish-Ukrainian border trains have to switch from 1.435 m to 1.520 m. Yet it takes time, and it requires certain hardware. If switching from one standard to another one is not very frequent, then its inconvenience is not very acute. But if households were to invest in two different video standards (VHS and Beta) in order to play what their family or friends bring, it would be expensive. They would rather expect that everybody sticks to a single standard. They do not care what standard it is, but it should be a single one. The saying is "the winner takes all"; it means that the "loosing" standard disappears quickly, since no new buyers are likely to show up (and perhaps many of the old users abandon it).

11.2 What disco are you likely to choose (if you think of the pre-COVID-19 time): the one which is not frequented by many users or the one where people crowd?

Of course there may be different attitudes, but many young people choose places where numerous friends can be met. Thus, the club that is frequented by more users provides each of their users with higher utility. This is the essence of what is called "external effects of consumption" in "network markets" analyses.

11.3 Are loyalty programmes effective?

Surprisingly, they are. There was an empirical research on the effectiveness of marketing campaigns. The famous result is that when faced with price differentials of 2% (one store sold a boom-box for 200 €, and another one sold it for 196 €) people were less excited than when faced with price differentials of 50% (one store sold a pen for 1 €, and another one sold it for 0.50 €); in this second case they were willing to undertake a larger effort to enjoy the lower price. Loyalty programmes are based on this attitude. If you tank fuel sold by the same company for a long time, then you can get a gift of 2 € or so. It does not make sense, because in order to get this gift you may forego an opportunity to pay less for the fuel elsewhere, but this is how we behave. Loyalty programmes rely on our emotions. They do not refer to any traits of rationality.

11.4 Think of the fast food market. If you know how to serve fast food, would you start your own independent business, or would you prefer to buy a franchise license to sell meals under the brand of McDonald's or some other large firm?

It depends on your ambition level, among other things. If you wish to operate locally then opening an independent business is an attractive option probably. Even if the production scale is low (which implies high costs), you are likely to manage to have an adequate number of clients. Franchise is a way to lower production costs by increasing the scale of operations. If you are interested in new technologies, entering new markets, and reaching new clients, then a franchise licence is an adequate choice probably.

11.5 In a more realistic version of the discrete Hotelling model, an assumption of zero production costs can be relaxed. How this translates into undercut-proof price policy?

Let us stick to the Coca Cola / Pepsi Cola example. Their production is probably similar, so it can be assumed that the unit cost of production is the same (say, c). It can also be assumed that "disutility" from drinking the non-preferred soft drink is the same for both types of fans, that is δ . The corollary derived by the end of the class "if $n_a = n_b$, then $p_a^U = p_b^U = 2\delta$ " can be rephrased as "if $n_a = n_b$, then $p_a^U = p_b^U = c + 2\delta$ ". In other words, companies should not charge their marginal cost of production c . If they charge 2δ more (the more consumers are idiosyncratic with respect to tasting the "wrong" drink, the higher the mark-up should be) they do not risk to expose their prices to undercutting. For instance, if the production cost is 0.30 €, and the "disutility" parameter is 0.10 €, then both drinks can be sold at the price of 0.50 € without a possibility of price undercutting.

If Coca Cola and Pepsi Cola are sold at, say, 0.55 €, then Pepsi Cola can offer the price, say, 0.44 €. Then it will attract not only its own fans, but also Coca Cola drinkers (who will suffer the disutility of 0.10 €, but – at the same time – they will save 0.11 € on the price). By undercutting the Coca Cola price, Pepsi Cola will make an extra profit of 0.14 € per every unit sold. By not undercutting the price – that is by charging the previous price – Pepsi Cola would make a profit of 0.25 € per piece, but the number of pieces would be only a half of the total (n_a instead of $n_a + n_b$) and thus the total profit would be lower; undercutting is profitable.

Now let us assume that both companies sell at, say, 0.45 €. Then if Pepsi Cola increases its price to 0.46, it will increase its profit. But Coca Cola can increase its profit as well by increasing its price to the same level. They can continue until they hit the level of 0.50 €. If one exceeds this level then the competitor can undercut the rival's price (like in the previous paragraph) and claim all the market.

Thus the only equilibrium price is 0.50 €. This simply demonstrates that the theorem from the lecture can be extended to the case when $c > 0$. If the price charged is higher than $c + 2\delta$, then it can be undercut. If it is lower, then rivals are not in equilibrium, because they will be better off if they increase it.

11.6 Why is the undercut-proof price equilibrium not a Nash equilibrium?

The undercut-proof equilibrium does not provide incentives to change the price for either of the rivals. In the Nash equilibrium definition, changes are unilateral. In the undercut-proof equilibrium motives to change prices are different. It is not sufficient to have an incentive to

change the price assuming that the rival does not react. The rival who changes the price anticipates that the new price will provide an opportunity for the other rival to react. Hence in the undercut-proof price equilibrium changes can be bilateral. There may be no incentives for a bilateral price change, but if one rival disregards the possibility of another rival to respond, then he or she may have an incentive to change the price unilaterally. As a result, the undercut-proof price equilibrium is not a Nash equilibrium.

12 – Portfolio selection

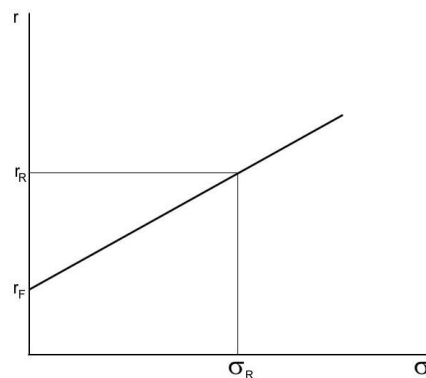
This lecture is devoted to what Quantitative Finance students are probably most interested in. It deals with the problem of a consumer who contemplates what to do with the money not spent on current consumption. Of course, everybody would like to get as much as possible. But markets do not allow to enjoy gains without a risk. There is a trade-off between the rate of return and the risk of losing money. In mathematical language it means that the rates of return on alternative investment projects are considered random variables. In models that we are going to analyse, it is assumed that means and variances of these random variables are known. It is also assumed that – as a rule – if one wants to achieve a higher rate of return, one needs to face a higher risk (measured e.g. by the variance of the rate).

A practical problem that needs to be solved can be summarised in the following way:

- How to take into account the rate of return and risk simultaneously? While
- Risk is to be measured with standard deviation (variance) of the rate of return of an asset

The most popular model addressing this problem is called Mean-variance model (see picture below). It answers the question about two assets with different mean rates of return (r), and variances of the respective rate of return (σ^2). One is risk-free, and the other is risky. The model assumptions read:

- Risk-free asset: $r_F > 0$, and $(\sigma_F)^2 = 0$ ($\sigma_F = 0$)
- Risky asset: $r_R > r_F$, and $(\sigma_R)^2 > 0$ ($\sigma_R > 0$).



The optimization problem reads:

- For a given level of expected rate of return find a portfolio with the least variance; or
- For a given level of variance (riskiness) find a portfolio with the highest expected rate of return.

The solution depends on an arbitrary adoption of the preferred rate of return (and minimising the variance), or an arbitrary adoption of the acceptable variance (and maximising the rate of return). Let α be a portion of the money spent on the risk-free asset ($1-\alpha$ is the portion of money spent on the risky asset). If $\alpha=1$ then only the risk-free asset is in the portfolio. If $\alpha=0$ then only the risky asset is in the portfolio. In the first case $r=r_F$ and $\sigma^2=(\sigma_F)^2=0$. In the second case $r=r_R$ and $\sigma^2=(\sigma_R)^2$. In general:

- Expected rate of return on both assets combined is: $r = \alpha r_F + (1-\alpha)r_R$, and
- Variance of rate of return is (for uncorrelated assets): $\sigma^2 = \alpha^2(\sigma_F)^2 + (1-\alpha)^2(\sigma_R)^2$.

The last equality can be simplified to $\sigma=(1-\alpha)\sigma_R$ (since $(\sigma_F)^2=0$, and $\sigma=(\sigma^2)^{1/2}$). It can be concluded that any solution (with rate of return between r_F and r_R , and risk measured by standard deviation not higher than σ_R) can be achieved by taking a convex combination of the risk-free and the risky asset (see picture above). The model can be extended by allowing $\alpha<0$. Negative α 's correspond to the fact that the risk-free asset can be sold rather than bought (money can be borrowed). Then any expected rate of return $r>r_R$ can be achieved (at the cost of variance increased over $(\sigma_R)^2$). This paradoxical finding is based on the fact that the rate of return is calculated with respect to the money owned, not with respect to the sum of money owned and borrowed.

This can be illustrated by the following example. An investor has 100 €. The risk-free asset gives the rate of return of 2%. The risky asset gives the expected rate of return of $r_R=10\%$, but the standard deviation of this rate is $\sigma_R=5$. How the investor should split his (or her) money? In other words, what α should be chosen? If $\alpha=1$, then $r=2\%$, and $\sigma=0$ (all the money is invested in the risk-free asset). If $\alpha=0$, then $r=10\%$, and $\sigma=5$ (all the money is invested in the risky asset). By taking $\alpha=0.5$ (i.e. investing 50 € in the risk-free asset and 50 € in the risky one), $r=6\%$, and $\sigma=2.5$. If the investor accepts the risk of, say, $\sigma=3$ but not more, then he (or she) should take $\alpha=0.4$ (i.e. invest 40 € in the risk-free asset and 60 € in the risky one). The expected rate of return will be 6.8% (maximum of what can be achieved by keeping the standard deviation at the level of 3 or below). If the investor wants to achieve the expected rate of return, say, of 8%, then $\alpha=0.25$ (25 € should be invested in the risk-free asset, and 75 € in the risky asset). Standard deviation of this rate will be 3.75 (the minimum of what can be achieved if the expected rate of return is to be 8%).

Now let us assume that money could be borrowed (α is negative, and the risk free asset can be used not as an investment opportunity, but rather as a 'disinvestment' opportunity). If the investor borrows, say 200 €, then he (or she) has 300 € to be invested in the risky asset. The expected rate of return will be 10%, i.e. 30 €. Out of this money the investor has to pay interest 4 € (2% of 200 €) which leaves 26 € as the 'net' gain. If this gain is related to 300 € invested, it will correspond to 8.67%. If the rate of return is related to the original amount of money owned by the investor, that is 100 €, it makes 26%. When calculating the expected rate of return when borrowing is possible, one needs to clarify whether the rate is related to the amount invested or to the amount originally owned.

What an investment strategy should be adopted by a consumer who can borrow money (at the rate r_F)? The choice follows the same rules as in the basic mean-variance model (minimise the variance for a given rate of return; or maximise the expected rate of return for a given variance). Please note that when borrowing is allowed (α can be negative) then the rate of return can be made arbitrarily high: $r = \alpha r_F + (1-\alpha)r_R$. But the riskiness will be high too: $\sigma = (1-\alpha)\sigma_R$ (please note that $1-\alpha$ is a positive number now). If 20 is the maximum standard deviation acceptable for the investor, then α can be calculated from the last equation: $20 = (1-\alpha)5$. The solution is $\alpha = -3$. This means that investor should borrow 300 € from the risk-free asset, and invest 400 € (300 € borrowed and 100 € originally owned) in the risky asset. The standard deviation will be 20 (as assumed) The expected return will be 40 €. Out of this 6 € should be paid for the amount borrowed. Hence the investor is left with 34 €. The rate of return is 8.5% when related to the amount invested or 34% when related to the amount originally owned.

This basic mean-variance model can be extended further by introducing two risky assets. Let there be two risky assets A_i with respective rates of return and variances: r_i , and $(\sigma_i)^2$ for $i=1,2$. They are combined with weights α_1 and α_2 , where $\alpha_2=1-\alpha_1$. The rate of return of this combination is given by the formula:

$$r(A_1, A_2) = \alpha_1 r_1 + \alpha_2 r_2,$$

and the variance is given by the formula:

$$\text{var}(A_1, A_2) = (\alpha_1)^2(\sigma_1)^2 + 2\alpha_1\alpha_2\text{cov}(A_1, A_2) + (\alpha_2)^2(\sigma_2)^2.$$

Hence for $\alpha_1, \alpha_2 > 0$ we have:

- $\text{var}(A_1, A_2) = (\alpha_1)^2(\sigma_1)^2 + (\alpha_2)^2(\sigma_2)^2$, if $\text{cov}(A_1, A_2) = 0$,
- $\text{var}(A_1, A_2) > (\alpha_1)^2(\sigma_1)^2 + (\alpha_2)^2(\sigma_2)^2$, if $\text{cov}(A_1, A_2) > 0$,
- $\text{var}(A_1, A_2) < (\alpha_1)^2(\sigma_1)^2 + (\alpha_2)^2(\sigma_2)^2$, if $\text{cov}(A_1, A_2) < 0$.

Students who remember that covariance is a measure of correlation of two random variables see that variance (riskiness) can be lowered by investing in assets that are negatively correlated. The following example explains the idea.

A consumer contemplates investing into two businesses:

- (1) umbrella; or
- (2) ice-cream.

But profitability depends on the weather which can be:

- (A) sunny; or
- (B) rainy

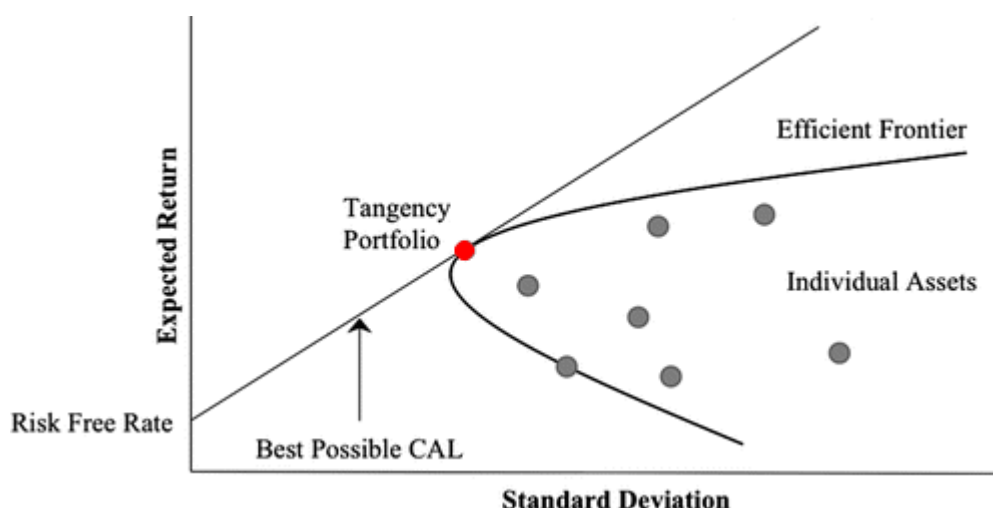
The weather is a random variable with 50%-50% probabilities (impossible to predict at the time of investment). Expected rate of return on investment is given in the following table:

	Rainy	Sunny
Umbrella	6%	2%
Ice-cream	2%	6%

If the consumer invests 50% of money into the umbrella and 50% into the ice-cream business, then – irrespective of the weather – the expected rate of return will be 4%, and variance will be 0. These numbers result from the formulae given above, but they can also be calculated directly from the table.

Mean-variance model for two assets can be generalized for any number of assets. In what follows we will introduce the Markowitz model. Harry Markovitz, born in 1927, got the Nobel prize in 1990. The model analyses investment opportunities of a person who faces many options, each of which is characterised by two numbers: r and σ (that is by the expected rate of return and standard deviation). When plotted in the r - σ plane, formulae for the expected rate of return and the variance of the combination of assets define a parabola (the "efficiency frontier"). This is the famous "Markovitz bullet", where:

- "tangency portfolio" represents the best combination of assets (given an opportunity to invest in a risk-free asset); and
- CAL is the Capital Allocation Line (see picture on page 87).



The model

- Explains how to optimally choose a combination of risky assets to invest (any combination is characterised by an expected rate of return and variance according to formulae from page 89).
- If a risk-free asset is available, and money can be borrowed, CAL indicates the highest rate of return possible given the variance of a combination of assets.
- The rate of return (and variance) indicated by the red dot can be achieved by investing all the money owned in the "tangency portfolio" (an optimum combination of assets).
- A higher rate of return (at the cost of a higher variance) can be achieved by borrowing money (investing a "negative" amount in the risk-free asset) and buying more units of the "tangency portfolio".

The "tangency portfolio" (the red dot in the picture above) does not have to be a specific asset; it can be a combination of other assets. The parabolic shape of the "Efficient frontier" results from the formulae on page 88 since the expected rate of return is a linear function of alphas, and the expected variance is a quadratic function of alphas. The Markovitz model illustrates how a rational investor chooses his or her portfolio. The "red dot" is a hypothetical

risky asset. It is assumed that the investor has enough money to invest in it. If the investor has less money, then he or she can afford buying a fraction of it only. The entire asset can be bought by borrowing money from the risk-free asset and investing in the risky one (as picture on page 1 illustrates). Any point along the CAL can be achieved as a combination of the risk-free asset and the "red dot" when borrowing is allowed. If borrowing is not allowed then only the segment between the vertical axis and the red dot can be achieved.

Questions and answers to lecture 12

12.1 How can the risk-free asset be interpreted?

The risk-free asset can be interpreted as a safe deposit in a bank. No banks are perfectly safe, but in many countries bank deposits enjoy state guarantees up to a certain limit. Interest rates offered by banks are close to zero ($r_F \approx 0$). If a bank is insolvent then a consumer can claim his (or her) deposit (without interest) from the state. The state can turn out to be insolvent either, but the probability of such an outcome is rather low. If one wants to take this into account then the assumption $\sigma_F = 0$ should be abandoned. Nevertheless, for practical purposes, bank deposits can be interpreted as risk-free assets.

12.2 Please prove that when money cannot be borrowed then r satisfies $r_F \leq r \leq r_R$

The inequalities are obvious when one looks at picture on page 87. They can be proved algebraically too. If money cannot be borrowed, then $\alpha \in [0, 1]$. The formula for the expected rate of return reads: $r = \alpha r_F + (1 - \alpha) r_R$. Observing that $r_F \leq r_R$, when r_F is substituted for r_R , one gets the first inequality, and when r_R is substituted for r_F , one gets the second inequality.

12.3 Is the assumption about borrowing at the risk-free rate realistic?

To answer this question one needs to ask if banks offer credits at the rate paid to those who deposit. If transaction costs are disregarded then they should. Yet transaction costs cannot be neglected. The banks serve as intermediaries between those who deposit and those who borrow, but they need to subtract from the money that flows through them in order to exist. Moreover, they have to take into account the risk of losing the money paid to some borrowers. Taking these circumstances into consideration, a more realistic assumption would be to allow borrowing at the rate somewhat higher than r_F .

12.4 Which method of calculating the expected rate of return do you consider more adequate: relative to what you own or relative to what you invest?

Examples calculated on pages 88-89 demonstrate that the difference between the two approaches can be quite significant (irrespectively of whether r_R and r_F differ largely or not). If related to the amount of money owned this rate of return can be made arbitrarily high. At the same time, the answer to 12.5 confirms (what some people intuitively feel) that the possibility of borrowing money does not allow to enjoy a return on investment higher than r_R (even when no money is borrowed); if the amount of money borrowed is very high, this return will be close to $r_R - r_F$. In my opinion the first method is more informative. The expected rate of return can be made arbitrarily high indeed, but its riskiness is reflected in the increased

variance. People who borrow a lot of money and invest in risky assets must take into account this variance.

12.5 Please prove that when money can be borrowed, then r understood as the expected rate of return on the amount invested satisfies: $\lim_{\alpha \rightarrow -\infty} r = r_R - r_F$.

If money can be borrowed then α can be negative. The amount invested is equal to $1-\alpha$ (1 corresponds to what was owned, and $-\alpha$ to what was borrowed). Thus the rate of return on investment reads $((1-\alpha)r_R - (-\alpha)r_F)/(1-\alpha) = r_R - (-\alpha)r_F/(1-\alpha)$. Please note that $-\alpha/(1-\alpha) \rightarrow 1$ when $\alpha \rightarrow -\infty$ which completes the proof.

12.6 Your client seeks your advice regarding investment in two projects. The first is to invest in photovoltaics, and the second is to invest in electric cars. What will you advice?

If the projects were among routine activities, then the mean-variance model could be applied. Photovoltaics and electric cars are fairly new, and it is difficult to predict expected rates of return and their variances. If there is a reliable data base, calculations are possible. But emerging technologies – especially those that are related to energy consumption – depend crucially on political decisions rather than market considerations. Professional advice should be based on studying political mechanisms likely to drive the future demand on renewable energy and environment-friendly transport.

12.7 In picture on page 90, the "tangency portfolio" (red dot) is characterised by lower riskiness than any other asset portrayed. Is it a rule?

No. It is not a rule; there may be assets with lower standard deviation. There are two features of the "Markovitz bullet" that characterise the "tangency portfolio". First, the rate of return on "tangency portfolio" is higher than on the risk-free asset (Capital Allocation Line, CAL, has a positive slope). Second, the "efficient frontier" is a part of the parabola running from the "tangency portfolio" to the right. There might be some points in this parabola which offer a lower standard deviation than the "tangency portfolio". But they lay below CAL (there is a possibility of enjoying the same standard deviation with a higher expected rate of return, if an appropriate convex combination of the risk-free asset and "tangency portfolio" is chosen).

13 – Introduction to financial economics

There are a number of formal assumptions adopted in financial economics. Many of them are so obvious that they do not need to be listed. The most important (and a controversial) one is the no-arbitrage assumption. It can be seen as equivalent to the one-price law or to an assumption that there is no possibility of enjoying profits without risk.

The *no-arbitrage* assumption can be ridiculed. Obviously everybody knows that there are arbitrage opportunities. Very often people observe that the same thing in one place may have a different price than somewhere else. They are happy to satisfy their consumption needs cheaper than otherwise. If they take commercial advantage of that, buy at a low price and resell at a higher one, they are arbitrageurs. The *no-arbitrage* assumption can be reconciled with the assumption that there are *arbitrageurs*, i.e. economic agents who try to make money

on arbitrage (and they succeed occasionally). In other words, the reason that – as a rule – there is no arbitrage, is because there are arbitrageurs who try to take advantage of asymmetric information, buy at a lower price, and sell at a higher one immediately. This cannot continue forever. Hence whenever there is such an opportunity, it disappears quickly.

Key concepts of financial economics are that of a "derivative" and of an "option". Their formal definitions are quoted below.

Definition

A derivative instrument ("derivative") is a contract with a payoff which depends on an (*ex ante*) unknown realization of a random variable.

- The random variable does not have to be a purely economic phenomenon (e.g. weather).
- A typical random variable is economic (e.g. exchange rate, performance index, asset price, etc.). In the Black-Scholes model it is the price of an underlying stock.

Definition

- An option is a contract which gives its party the right (but not the obligation) to buy or sell a commodity at a predetermined price. A "call option" ("call") gives the right to buy, and a "put option" ("put") – the right to sell.
- In a European option, the right is to be executed at a given date (typically on the third Friday of a given month). In an American option, the right is to be executed any time prior to the given date. In a Bermudan option, the right is subject to specific conditions.
- Options can be seen as derivatives, since the execution of the right depends on the market price of the commodity (considered a random variable).

The concept of a risk-free asset was introduced in the previous lecture (QF-12). Here it is called "money".

Definition

The risk-free asset is called money. Its value increases over time at the pace of the discount rate. If B_t is the value of money at time t then its value at time τ is $B_\tau = e^{r(\tau-t)}B_t$, where r is a discount rate. Please note that if $\tau=t$, then – according to the formula – $B_\tau=B_t$. If $r>0$, then $B_\tau>B_t$ for $\tau>t$, and $B_\tau<B_t$ for $\tau<t$. This is not risk-free profiting since – by definition of the discount rate (see lecture 4, QF-4) – economic agents are indifferent between B_τ at time τ and B_t at time t .

Financial specialists speak a professional jargon which makes use of two adjectives: "short" and "long". Meaning of these words is different than in the everyday language.

Definition

Short selling – selling an asset that is not owned by the seller. This can be interpreted as borrowing from its owner, and enjoying the right to use the asset. If the owner asks for the asset, then it needs to be bought in the market and given back to the owner. By short selling, economic agents expect that the market price of the asset will decrease and thus they will

make a profit (having returned the asset to where they borrowed from, but keeping the difference between the original higher and final lower price).

Definition

Long position – buying an asset whose value is expected to rise. Investors in 'long position' expect to make a profit by selling the asset when its price is higher than the original one. With respect to options:

- Long position on call options – investor expects the price of the underlying asset to increase;
- Long position on put options – investor expects the price of the underlying asset to decrease.

Financial specialists use also specific expressions in order to distinguish between what is written in a contract, and what the market determines.

Definitions

- Strike price (exercise price) – the price indicated in an option (for call options: the holder has the right to buy the underlying asset at this price; for put options: the holder has the right to sell at this price).
- Spot price – the price established by the market.
- The call option is Out-of-The-Money (OTM) if the spot price of the underlying asset is lower than the strike price (it is better not to execute the option, as observed in the class). It is In-The-Money (ITM) if the spot price of the underlying asset is higher than the strike price (it is profitable to execute the option, likewise, as observed in the class).
- For put options it is the other way around.

Now we have to introduce notation necessary to define the famous Black-Scholes model. It establishes a relationship between stocks (capital assets) and values of their derivatives. The stock can be understood as a firm, and derivatives can be understood as contracts which use the stock in order to accomplish certain transactions.

- S – the (spot) price of the underlying stock,
- $V(S,t)$ – the price of its derivative,
- $C(S,t)$ – the price of a European call option,
- $P(S,t)$ – the price of a European put option,
- K – the strike price of the option,
- σ – the standard deviation of the stock's returns,
- r – rate of return on the risk-free asset,
- T – expiry date ("maturity"; established in the option).

Fischer Black (1938-1995) and Myron Scholes (born in 1941) developed the most famous equation in financial mathematics in 1973 (see below). It was based on sophisticated modelling of how to invest in order to reconcile the desire to make profits with avoiding excessive risk. Later on it was refined with the assistance of Robert Merton (born in 1944). In some textbooks the model is called BSM (referring to three names: Black, Scholes and

Merton). Scholes and Merton got the Nobel prize in 1997. Soon afterwards LTCM (*Long-Term Capital Management*) – with Scholes and Merton among its key advisors – made a 4.6 billion dollar loss and collapsed spectacularly. Nevertheless the collapse was caused not by the inadequacy of the model, but by a wrong estimation of risks involved.

Black-Scholes equation

$$\partial V / \partial t + (1/2)\sigma^2 S^2 \partial^2 V / \partial S^2 + rS \partial V / \partial S = rV$$

This is a partial differential equation established as a result of analysing how to optimally manage investment risk (accomplished through "hedging", i.e. combining options and shares in the underlying stock). Please note that the equation does not involve C and P explicitly. As derivatives, they need to be substituted for V. If one does this, the following relationships can be derived (following tedious calculations). They are called Black-Scholes formula.

Black-Scholes formula

- For European call options:

$$C(S_t, t) = N(d_1)S_t - N(d_2)Ke^{-r(T-t)},$$

- For European put options:

$$P(S_t, t) = N(-d_2)Ke^{-r(T-t)} - N(-d_1)S_t$$

where:

$$d_1 = 1/(\sigma(T-t)^{1/2})(\ln(S_t/K) + (r + \sigma^2/2)(T-t))$$

$$d_2 = d_1 - \sigma(T-t)^{1/2}$$

N is the normal distribution function, i.e.

$$N(x) = 1/(2\pi)^{1/2} \int_{-\infty}^x (\exp(-z^2/2))dz$$

This is what students can read in financial economics textbooks, and this is what caused the bankruptcy of LTCM. The equation has a firm theoretical justification, but the formula depends on specific probabilistic assumptions about risks involved. In what you see above, it is assumed that random variables are distributed normally. Sometimes they are, but in order to make reliable predictions (i.e. to avoid a bankruptcy), a deeper inquiry into the nature of risks involved is recommended.

Black-Scholes formula seems to be very complicated at the first glance. In fact, it can be interpreted in fairly simple terms, if one rewrites it in the following way:

$$C(S_t, t) = (\text{Risk factor}) \times S_t - (\text{Risk factor}) \times \text{PV}(K),$$

$$P(S_t, t) = (\text{Risk factor}) \times \text{PV}(K) - (\text{Risk factor}) \times S_t.$$

Call and put options allow the holder to buy and sell, respectively, something that has the price S_t now for the price of K in the future. $\text{PV}(K)$ is the present value of this future price. Disregarding risk factors, the price of the call option is the difference between the current price and the future price of the asset. The price of the put option is the difference between the

future price and the current price. Risk factors modify these differences in order to take into account the fact that future prices are uncertain.

The Black-Scholes formula is derived from the Black-Scholes equation based on the assumption of the normal distribution of random effects and the lack of factors disturbing financial transactions. Its extensions (not quoted in this lecture) take into account non-normal distributions, and relaxing "classical" assumptions (such as zero transaction costs, no taxes, etc.). For many politicians and their constituencies, financial transactions – especially those involving derivatives – are considered a promising taxation target. There are at least two reasons for that. First, financial transactions are undertaken usually by individuals who are better off. Common wisdom suggests that taxing them does not affect the poor. (Whether this is really true, an appropriate general equilibrium model – see QF-6 – has to be analysed.) Second, several acute economic crises were triggered by financial speculations. It is expected that when financial transactions are more difficult to carry out, people do not undertake them recklessly, and inevitable mistakes do not accumulate quickly.

The Black-Scholes formula establishes European call and put option unique prices based on:

- S_t – the current price of the underlying stock, and
- K – the strike price of the option.

A careful scrutiny of the formulae on page 95 lets discover that these prices depend on S_t rather than S_T . Thus they are not based on S_T (i.e. the stock price at the time of maturity of the option), but rather on its current price, S_t . This is probably the most interesting conclusion from the Black-Scholes model. Once again, it should be emphasised that practical applications of the model are risky because of the rigorous probabilistic assumptions required.

Questions and answers to lecture 13

13.1 Is the no-arbitrage assumption realistic?

The no-arbitrage assumption can be ridiculed by referring to the following joke. Two men walk along a street. The non-economist spots a 10-euro bill by the kerb and he wants to pick it up. The economist tells him: "don't bother, don't pick it up; it's fake". The non-economist asks "how do you know that it's fake?" The economist responds: "if it was real somebody else would have picked it up already".

According to economists, arbitrage is not possible, because if it were possible, somebody else would have taken advantage of it already. The arbitrage possibility means that if there are two different prices of the same good, you can buy it for the lower price, and resell it immediately for the higher price. Let us analyse possible circumstances when two different prices of the same good co-exist. One possibility is that the lower price is enjoyed somewhere at a distant location. You go to this distant location, buy and resell it in a more attractive location. If you take into account the travel cost, the price difference shrinks and perhaps disappears. Another possibility is that the same good has two different prices in places very close to each other (think of a chocolate bar in a grocery store and in a nearby restaurant; the restaurant price is higher because customers pay not only for the product, but for the service as well). You buy a

good at a low price and immediately – almost without any travel – resell it at a high price. In order to take advantage of this you have to bear some risk (namely that you will end up with the good that was not bought by anybody). Yet another possibility is that you buy at a wholesale (low) price and resell at a retail (high) price. Here again you act as an entrepreneur and you bear some (perhaps acceptable) risk that you will end up with an inventory of what you expected to sell. Finally there is a possibility that the selling takes place in the next period, and – if discounting is taken into account – the higher price is simply the same as the buying one after discounting.

It is possible to take advantage of buying cheaply and selling expensively – this is the essence of entrepreneurship. Nevertheless, it implies some risk. The no-arbitrage assumption says that whenever one takes such advantage, one bears some risk. Of course, it is possible to find a real 10-euro bill on the street. Yet such opportunities are not frequent and economics is about what happens routinely.

13.2 Please give an example of a derivative contract with a non-economic random variable.

Flood insurance. If there is no flood, you get nothing. If there is a flood, you get the compensation that was agreed upon in the contract.

13.3 Please give an example of a derivative contract with an economic random variable.

A typical derivative contract is to sell something on June 30, 2021 for the price of, say, 10 euros if the company X enjoys sale growth between 2019 and 2020, and for the price of, say, 12 euros if X's sales stagnated or shrank between 2019 and 2020. Sales of the company is considered a random variable.

13.4 The European "call" option must have a lower price (*caeteris paribus*) than the American "call" option. Is this a true statement?

If the market takes into account all relevant circumstances, European options should be cheaper than American ones with the same parameters. The reason is that a holder of the American option can take advantage of it on the very last day (like in the European option). But the American option offers a greater flexibility; it can also be realised on some earlier date if this is found attractive. In particular, if the spot price of the "call" is higher than the strike price and it is expected to fall, the holder may choose to execute the right to buy (for the strike price) and then sell immediately (for the spot price). The profit (the difference between the spot price and the strike price) is expected to be higher now than at the time of maturity.

13.5 According to the definition of money (risk-free asset, see page 93), its value at time τ is $B_\tau = e^{r(\tau-t)}B_t$, where r is a discount rate, and B_t is its value at time t . What is the present value at t of what has the value B_τ at τ ?

The present value is B_t . To calculate this, one needs to apply the 'continuous time' version of the formula $X_T = X_0(1+r)^T$ which reads $X_T = X_0e^{rT}$ (page 25 of the fourth lecture, QF-4). In this case $T = \tau - t$. Hence the present value of B_τ is $e^{r(t-\tau)}B_\tau = B_t$.

13.6 The strike price of a put option is 10 €. Its spot price at the time of maturity turned out to be 12 €. Is it ITM?

No. The 10 € strike price of the put option means that the holder has the right to (but not an obligation) sell something for 10 €. If the spot price is 12 €, then the holder of the option prefers to sell it for this price rather than for the price guaranteed by the option. Hence the option is OTM.

13.7 Please prove that if the price of an option does not depend on the price of the underlying stock, then its price grows at the rate of interest (the rate of growth of the value of money, i.e. the risk-free asset).

This is rather unusual, but we can assume formally, that $\partial V/\partial S=0$. Also the second derivative vanishes then, $\partial^2 V/\partial S^2=0$. Under these circumstances the Black-Scholes equation reads: $\partial V/\partial t=rV$. This is a simple differential equation for V as a function of time. Its solution is $V(t)=V(0)e^{rt}$, i.e. exactly according to the formula for money (see page 83). Usually the price of an option does depend on S , and the equation cannot be solved so easily.

13.8 Let us assume that a European call option has the maturity date 2022, the strike price is 100 €, the price of its underlying asset in 2020 is 120 €, standard deviation of the stock's returns is 4, and the risk-free interest rate is 2%. What is the price of this option in 2020 according to the Black-Scholes formula?

According to the Black-Scholes formula, $C(S_t, t) = N(d_1)S_t - N(d_2)Ke^{-r(T-t)}$, where:

- $d_1 = 1/(\sigma(T-t)^{1/2})(\ln(S_t/K) + (r + \sigma^2/2)(T-t))$,
- $d_2 = d_1 - \sigma(T-t)^{1/2}$, and
- N is the normal distribution function, i.e. $N(x) = 1/(2\pi)^{1/2} \int_{-\infty}^x e^{-z^2/2} dz$.

Please note that $T-t=2$, $r=0.02$, and $\sigma=0.04$. We can calculate that $d_1=3.9558$, and $d_2=3.8992$. $N(d_1)=0.99996$, and $N(d_2)=0.99995$. Hence, according to the Black-Scholes formula, $C(120, 2020)=23.92$ €. The price of the underlying stock is 120 € now. The option gives an opportunity to buy it in two years from now for the price of 100 €. If the price of the asset does not change, this is like a profit of 20 € to be enjoyed in two years. In fact, the holder of the option can enjoy an even higher profit if the price of the stock grows above 120 €. But the price can also decline. Taking into account its volatility (reflected by σ), and the possibility of investing in the risk-free asset (reflected by r), the formula gives the price which "hedges" investors against excessive risk. If random variables follow non-normal distributions (as it happens often) then numbers calculated by the Black-Scholes formula above are misleading.

13.9 Let us assume that a European put option has the maturity date 2022, the strike price is 120 €, the price of its underlying asset in 2020 is 100 €, standard deviation of the stock's returns is 4, and the risk-free interest rate is 2%. What is the price of this option in 2020 according to the Black-Scholes formula?

According to the Black-Scholes formula, $P(S_t, t) = N(-d_2)Ke^{-r(T-t)} - N(-d_1)S_t$, where:

- $d_1 = 1/(\sigma(T-t)^{1/2})(\ln(S_t/K) + (r + \sigma^2/2)(T-t))$,
- $d_2 = d_1 - \sigma(T-t)^{1/2}$, and
- N is the normal distribution function, i.e. $N(x) = 1/(2\pi)^{1/2} \int_{-\infty}^x e^{-z^2/2} dz$.

We can use some numbers from 13.8. They let us calculate $d_1=-2.4859$, and $d_2=-2.5427$. $N(-d_1)=0.99354$, and $N(-d_2)=0.9945$. Hence $P(100, 2020) = N(-d_2)120e^{-r(T-t)} - N(-d_1)100=15.31$.

14 – Econophysics

Econophysics is an approach to microeconomic problems by assuming that economic objects behave like molecules (studied by physics), and therefore the results of their behaviour can be modelled by equations analysed in physics. It is a fast growing area of inquiry, but economists raise doubts with respect to some of its findings. Also our brief class discussion demonstrated that a field where this approach can be applied is limited. Nevertheless econophysics provides useful insights into many problems of financial economics.

In particular, it is assumed that certain economic phenomena can be modelled by the well-known in mathematics *Hit Diffusion Equation*:

$$\partial u / \partial t = \alpha (\partial^2 u / \partial x^2 + \partial^2 u / \partial y^2 + \partial^2 u / \partial z^2).$$

Sometimes the equation is written in a more concise form:

$$\partial u / \partial t = \alpha \nabla^2 u,$$

where

∇ is the so-called Laplace operator. In a 3-dimensional space: $\nabla^2 u = \partial^2 u / \partial x^2 + \partial^2 u / \partial y^2 + \partial^2 u / \partial z^2$

The equation models temperature changes $u(x,y,z,t)$ of a spatial object if there is no external source of heat. The temperature varies in time t , and in all spatial dimensions x , y , and z . If there is an external source of heat, an additional term is added at the right-hand-side of the equation.

The model can be simplified by assuming that there is only one spatial dimension:

$$\partial u / \partial t = \alpha \partial^2 u / \partial x^2.$$

In this case, the equation predicts the distribution of temperature in a one-dimensional object (a rod) which is perfectly insulated (there is no exchange of heat with its environment).

Financial economic models apply this approach by assuming that assets are priced in processes similar to those which determine distribution of the temperature in physical objects. To some extent this is justified by the fact that people who buy and sell communicate with each other, observe prices established in the market, and everything takes time. The spatial dimension has to be interpreted in a reasonable way. The most immediate idea is to think of a geographical distance from where some transactions take place. But this is not consistent with how information is transported. It is difficult to indicate a convincing interpretation of economic processes in terms of movement of molecules. The *no-arbitrage* assumption is said to correspond to the First Law of Thermodynamics. In Thermodynamics it is determined that the available energy is constant; it may be transformed from one type to another one, but it is neither created nor annihilated. If coal is burnt in a power plant, a part of the chemical energy contained in the fuel is transformed into electricity, a part can be used by a central heating system, a part is used to run the plant, and a part is dissipated in the form of wasted thermal energy; but nothing vanishes. Likewise in economic systems. Value can be created only in

production processes where something can be added to what existed before as production inputs only in exchange for a risk. The *no arbitrage* assumption states that it is not possible to enjoy risk-free profits.

Robert Brown (1773-1858) was a Scottish botanist who first realised that particles move in random directions (he used a microscope to observe pollens in a droplet). By Brownian motion we understand random movement of molecules in fluids. This concept – explained by physicists convincingly – is used in economics in order to model the movement of economic variables that are considered random. There are two versions of Brownian motion:

- Random walk; and
- Wiener process (continuous version of a random walk).

In the random walk we observe changes in subsequent discrete time moments. It is more convenient sometimes to analyse these changes in continuous time. This is how a Wiener process – a mathematical representation of a Brownian motion – is interpreted. Its formal mathematical definition is stated below:

- $W_0 = 0$ almost surely (with probability 1 $W_0=0$),
- W has independent increments: for every $t>0$, the future increments $W_{t+u}-W_t$, $u\geq 0$, are independent of the past values W_s , $s<t$,
- W has Gaussian increments: $W_{t+u}-W_t$ is normally distributed with mean 0 and variance u , i.e. $W_{t+u}-W_t \sim N(0,u)$; in other words, the variance of an increment grows with its length, but its mean is zero,
- W has continuous paths almost surely (with probability 1 W_t is continuous in t).

The language of this definition requires explanations. First of all, the expression 'almost surely' has to be explained. The words 'almost surely' are used in mathematics when a lay person would say 'for sure'. 'Almost surely' is a term used when the measure of the set where something does not happen is zero. For example, the length of a segment $[1,4]$ is three (in mathematical language, its measure is 3). A single point, like $\{2\}$ has zero length (its measure is zero). In this case expression 'almost surely' means that something happens for $x\in[1,4]$, and the set of points where it does not happen consists of isolated points (like $\{2\}$). The first bullet of this definition says that $W_0=0$, perhaps except for some rare circumstances. The last bullet says that the path of W is continuous, perhaps except for some isolated time moments.

The second bullet says that $W_{t+u}-W_t$, are independent of the past values (W_s). It means that what will happen after s does not depend on what happened before s . For instance if W grew before s , it does not mean anything for what will happen after s : it can grow, or it can decline – the probabilities do not change.

The third bullet interpretation is the most complicated. The normal distribution is fairly obvious. Many real-life processes follow the distribution discovered by Carl Friedrich Gauss (1777-1855): everything can happen, but moving away from the average by x , decreases its likelihood x^2 times. The first part of the bullet says that everything can happen, but large changes are unlikely. The second part of the bullet is more difficult to grasp. It says that if the

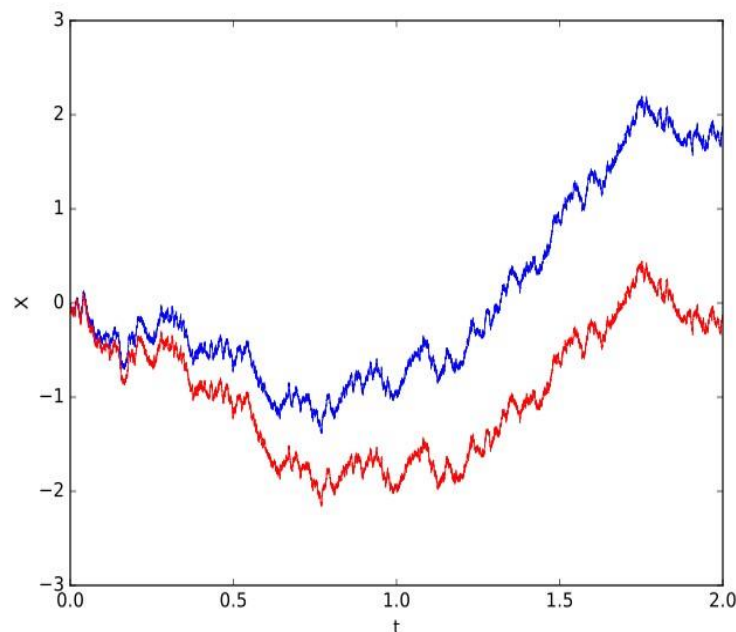
time span analysed doubles, then standard deviation doubles as well. This is easy to comprehend. The more difficult part is to see what happens with the standard deviation when the time elapsed gets closer to zero. According to the formula from the third bullet, it gets closer to zero too. The proportion between the standard deviation and time, $\sqrt{t}$, is constant and equal to 1. Mathematicians say that infinitesimal standard deviation is 1. The infinitesimal standard deviation of a Wiener process is 1.

The next definition introduces the concept of a stochastic process:

$$X_t = \mu t + \sigma W_t,$$

where:

- W – Wiener process,
- μ – drift (long-term trend)
- σ^2 – infinitesimal variance (terminology of stochastic processes; interpreted as "volatility" – the standard deviation is proportional to time t elapsed since the start: $\sqrt{t}\sigma$)



- Red line – without drift ($\mu=0$),
- Blue line – with drift ($\mu>0$).

Picture above illustrates a stochastic process based on a Wiener process. Please note that in any moment t , W_t goes either up or down, and it is impossible to predict the direction. Between 0 and 1, it went down more often, but between 1 and 2 it was the other way around; movements were random. Standard deviation (σ) determines the height of an individual "pixel": the larger the σ the higher the "pixel". The blue line reflects a growing trend (corresponding to $\mu>0$). A stochastic process with a declining trend ($\mu<0$) would be illustrated by a line below the red one.

A similar equation (called Geometric Brownian Motion) was used in order to establish the Black-Scholes model. Mathematicians consider it Stochastic Differential Equation:

$$dS_t = \mu S_t dt + \sigma S_t dW_t$$

whose solution is $S_t = S_0 \exp((\mu - \sigma^2/2)t + \sigma W_t)$.

Wiener processes in financial economics provide a theoretical foundation for analyses of derivatives:

- W can be interpreted as the underlying economic variable (e.g. the return on an asset); and
- S can be interpreted as an index characterising a derivative instrument (e.g. a European option).

It is not always clear whether economic phenomena can be researched as random events. By developing so-called fundamental analysis, economists try to explain economic variables by what happens in the economy; if the return on an asset grows, it may result from technology changes, introducing new regulations, finding new clients, and so on. At the same time, other analysts try to explain these phenomena by studying historical patterns of the trends; this is called technical analysis. Irrespective of how investors prepare their decisions, they can always be surprised with the results. Hence returns on an asset have to be considered random to a large extent.

Econophysical methods proved to fit empirical data fairly well. They provide inspiration for a number of researchers to quantify and explain what used to be considered impossible to analyse systematically. Nevertheless many economists observe that more fundamental work is needed to justify using equations developed to explain movements of molecules in order to reflect what happens in an economy.

Questions and answers to lecture 14

14.1 When does heat diffusion equation ($\partial u / \partial t = \alpha \partial^2 u / \partial x^2$) predict the linear shape of the u function?

The linear shape of the u function implies its second derivative vanishing: $\alpha \partial^2 u / \partial x^2 = 0$. By the heat equation, it means that $\partial u / \partial t = 0$. This, in turn, means that the function u does not depend on time. $u(x, t) = \beta x + \gamma$ is an example of such a function: it does not depend on t, and it is linear in x. Intuitively this means that if the distribution of temperature along an object does not change in time (it is stable), it must be linear (not necessarily constant).

14.2 Please provide an example of a random walk.

Perhaps the simplest example of a random walk is registering the outcomes of tossing a coin. In a single step it can be either Heads or Tails. Irrespective of how many Heads or Tails happened in the past, the current probability of Heads or Tails is 50%. Heads can be represented by going up, and Tails can be represented by going down. Pictures below illustrate possible outcomes of 17 such steps. The first diagram corresponds to 17 Heads in a row. The second represents 17 Tails in a row. The third one represents a regular series Heads-Tails-Heads-Tails and so on. The last one looks 'random' indeed; outcomes are not regular –

sometimes it is Heads, sometimes it is Tails. Please note that the probability of any of these paths is the same: $1/2^{17}=1/131072=0.00000762394531$ (a very small number). Even though the last one looks as the most 'random' one, in fact all of them are equally unlikely.



14.3 Is W defined by the following procedure a Wiener process? $W_0=0$, and $W_{t+u}-W_t$ is normally distributed with mean W_t and variance u , i.e. $W_{t+u}-W_t \sim N(W_t, u)$.

No. Increments are not independent of the past values. A positive W_t makes it more likely to go up, and a negative W_t makes it more likely to go down.

14.4 Please provide economic interpretation of the Geometric Brownian Motion equation ($dS_t = \mu S_t dt + \sigma S_t dW_t$).

S_t can be interpreted as the value of a stock and W_t as a realisation of some random variable. The value of the stock goes up or down randomly. How much it goes up or down depends on σ . If σ is close to zero, then the value of the stock is stable; otherwise it adds to, or subtracts from, S_t proportionally to σ . On top of that there is a long-term trend μ . Irrespective of random ups and downs, the value of the stock tends to increase (if $\mu > 0$), or decreases (if $\mu < 0$) as time elapses.

Please note that a different interpretation can be offered as well (like in the lecture). For instance, W can be interpreted as an "underlying economic variable" (e.g. the return on an asset); it depends on a number of factors difficult to predict – thus it is a random variable. At the same time, S can be interpreted as an index characterising a derivative instrument (e.g. a European option); it is also random, but it can be explained by what happens with the asset.

14.5 Can the price of a stock be seen as a random variable?

To some extent, yes. Stock prices reflect the ability of the company to bring profits. Its profitability depends on a number of quantifiable factors such as production capacity, management solutions, market demand, sector competitiveness, government regulations and so on. But there are also factors difficult to predict, such as technical innovations, changing consumer preferences, politics, and so on. Analysts who try to build econometric models calculating stock prices have to admit, that – no matter how many explanatory variables they take into account – some variability cannot be explained. It is this unexplained part of the stock price variability that must be seen as a random variable.

14.6 Can econophysics explain economic phenomena?

Perhaps some of them can be explained by models developed by physicists. If economic agents behave like molecules (they move randomly in various directions), the supply or demand can be modelled using the heat diffusion equation. Yet some decisions cannot be seen as random; there are unpredictable, but not necessarily random. They are taken by individuals who try to be smarter than their neighbours.

Molecules do not have free will. Albert Einstein did not want to acknowledge the fact that some physical quantities take probabilistic values. The mainstream physics now accepts the theory that they do. Human activities are even more complicated: in addition to their probabilistic aspect, they can be unpredictable because of the free will. In economics it is reflected by the consumer sovereignty assumption. Evolution of preferences can be predicted to some extent only.

15 – Kuhn-Tucker theorem

The Kuhn-Tucker theorem (published in 1951 by Harold W. Kuhn, 1925-2014, and Albert W. Tucker, 1905-1995), is perhaps the most widely used mathematical theorem in economics. It extends the 18th and 19th century findings on constrained maxima to account for non-strict inequalities (the original analyses dealt either with strict inequalities or equalities). It does not have important applications in physics, and it does not belong to the obligatory material of lectures in mathematical analysis. At the same time it is too complicated to be proved in economic textbooks (the mathematical appendix to Mas-Colell A., Whinston M. D., Green J., *Microeconomic Theory*, Oxford University Press 1995 is an exception). As a result, most economics students are informed that many findings are based on the Kuhn-Tucker theorem, but they do not have an opportunity to see it being proved.

The theorem is associated usually with two names only: Kuhn and Tucker. It turns out that the necessary conditions for the problem had been stated and proved by William Karush (1917-1997) in his master's thesis in 1939 (University of Chicago). Thus the theorem is sometimes referred to as KKT conditions. This lecture provides a proof of the theorem, but it requires numerous preparations.

Definition (conditional maximum or minimum)

A function $f: \mathcal{R}^n \rightarrow \mathcal{R}$ has a conditional maximum (minimum) subject to $g_1(\mathbf{x})=b_1, \dots, g_m(\mathbf{x})=b_m$ for $\mathbf{x}^* \in \mathcal{R}^n$, if for any $\mathbf{x} \in \mathcal{R}^n$ satisfying $g_1(\mathbf{x})=b_1, \dots, g_m(\mathbf{x})=b_m$ there is: $f(\mathbf{x}) \leq (\geq) f(\mathbf{x}^*)$. The problem of finding a conditional maximum (minimum) is coded as:

$$\text{Max(Min)}_{\mathbf{x}} \{ f(\mathbf{x}) : g_1(\mathbf{x})=b_1, \dots, g_m(\mathbf{x})=b_m \}.$$

The Lagrange function was introduced in the first lecture (QF-1). Nevertheless it will be redefined here, making sure that the notation is consistent with the current one.

Definition (Lagrange function)

For the conditional maximum problem the Lagrange function $L: \mathcal{R}^{n+m} \rightarrow \mathcal{R}$ is defined as:

$$L(\mathbf{x}, \lambda_1, \dots, \lambda_m) = f(\mathbf{x}) + \lambda_1(b_1 - g_1(\mathbf{x})) + \dots + \lambda_m(b_m - g_m(\mathbf{x}))$$

Variables $\lambda_1, \dots, \lambda_m$ are called "Lagrange multipliers".

The following theorem is universally taught and hence known widely, but it is a necessary part of our argument.

Theorem (Lagrange)

Let us assume that functions $f, g_1, \dots, g_m$ are continuously differentiable and the matrix of derivatives of $g_1, \dots, g_m$ (the Jacobi matrix) has the rank m . Then:

1. If $\mathbf{x}^*$ is the solution of the conditional maximum problem, then there exist $\lambda_1^*, \dots, \lambda_m^*$, such that $(\mathbf{x}^*, \lambda_1^*, \dots, \lambda_m^*)$ solve system of equations $\partial L / \partial \mathbf{x} = \mathbf{0}, \partial L / \partial \lambda = \mathbf{0}$.
2. If functions $f, g_1, \dots, g_m$ are twice continuously differentiable, and matrix of second derivatives $\partial^2 L / \partial \mathbf{x}^2$ (the Hesse matrix of L) is negatively definite, then the point $\mathbf{x}^*$ from the solution $(\mathbf{x}^*, \lambda_1^*, \dots, \lambda_m^*)$ of the system $\partial L / \partial \mathbf{x} = \mathbf{0}, \partial L / \partial \lambda = \mathbf{0}$ solves the original problem.

Please note that equations $\partial L(\mathbf{x}^*, \lambda_1^*, \dots, \lambda_m^*) / \partial \mathbf{x} = \mathbf{0}$ are equivalent to $\partial f(\mathbf{x}^*) / \partial \mathbf{x} = \lambda_1^* \partial g_1(\mathbf{x}^*) / \partial \mathbf{x} + \dots + \lambda_m^* \partial g_m(\mathbf{x}^*) / \partial \mathbf{x}$, where the equality applies to each of the coordinates, and $\partial g_j(\mathbf{x}^*) / \partial \mathbf{x}$ are vectors of partial derivatives with respect to $x_1, \dots, x_n$. Equations $\partial L(\mathbf{x}^*, \lambda_1^*, \dots, \lambda_m^*) / \partial \lambda = \mathbf{0}$ are equivalent to $b_j - g_j(\mathbf{x}^*) = 0$ (for $j=1, \dots, m$). Please also note that in the Lagrange theorem (2) for a conditional minimum, the Hesse matrix should be positively definite.

Lagrange multipliers have an easy interpretation. Taking into account that $L(\mathbf{x}^*, \lambda_1^*, \dots, \lambda_m^*) = f(\mathbf{x}^*)$, then $F(\cdot) = f(\mathbf{x}^*)$ understood as a function of right hand sides of constraints $(b_1, \dots, b_m)$ yields as derivatives: $\lambda_j^* = \partial F / \partial b_j$.

Up to this moment, everything could be found in the 19th century textbooks. Now we aim to extend this theory towards problems encountered in economics.

Definition (mathematical programming; optimization)

$\text{Max}_{\mathbf{x}} \{f(\mathbf{x}): g_1(\mathbf{x}) \leq b_1, \dots, g_m(\mathbf{x}) \leq b_m, \mathbf{x} \geq \mathbf{0}\}$ (likewise for the minimum). Compared to the 'classical' conditional maximum, in this definition inequalities substitute for equalities, and variables must be non-negative.

It is easy to observe that in the optimization problem defined above an equality $\mathbf{g}(\mathbf{x}) = \mathbf{b}$ can be represented by a pair of inequalities: $\mathbf{g}(\mathbf{x}) \leq \mathbf{b}$ and $-\mathbf{g}(\mathbf{x}) \leq -\mathbf{b}$. In an optimization problem, if a variable x_i may be a negative one, then it can be substituted with a pair of non-negative variables x_i^- and x_i^+ , such that $x_i = x_i^+ - x_i^-$. It lets us state (as a corollary) that a conditional maximum (minimum) problem is a special case of an optimization problem as defined above. The following lemma has an easy proof.

Lemma

If f is twice continuously differentiable, then a solution $\mathbf{x}^*$ to the optimization problem $\text{Max}_{\mathbf{x}} \{f(\mathbf{x}): \mathbf{x} \geq \mathbf{0}\}$ must satisfy the following $2n+1$ conditions:

$$\partial f(\mathbf{x}^*) / \partial \mathbf{x} \leq \mathbf{0}, \mathbf{x}^* \geq \mathbf{0} \text{ and } \partial f(\mathbf{x}^*) / \partial \mathbf{x} \cdot \mathbf{x}^* = 0.$$

The last condition ("complementarity") can be written as:

$$\partial f(\mathbf{x}^*) / \partial x_1 \cdot x_1^* + \dots + \partial f(\mathbf{x}^*) / \partial x_n \cdot x_n^* = 0$$

and – given the other $2n$ conditions – is equivalent to the following series of implications:

- if $x_i^* > 0$, then $\partial f(\mathbf{x}^*)/\partial x_i = 0$, and
- if $\partial f(\mathbf{x}^*)/\partial x_i < 0$, then $x_i^* = 0$.

Before we proceed further, let us observe that inequalities in the optimization problem ($\mathbf{g}(\mathbf{x}) \leq \mathbf{b}$) can be substituted with equalities, as in the conditional maximum problem, once m auxiliary variables $s_i \geq 0$ are introduced and inserted into the constraints: $\mathbf{g}(\mathbf{x}) + \mathbf{s} = \mathbf{b}$. The Jacobi matrix of this new constraints systems has the rank m , since it includes a unitary matrix of that rank). Therefore an appropriate assumption from the Lagrange theorem is satisfied.

Definition (differential Kuhn-Tucker conditions for an optimization problem)

- $\partial f(\mathbf{x})/\partial \mathbf{x} \leq \lambda_1 \partial g_1(\mathbf{x})/\partial \mathbf{x} + \dots + \lambda_m \partial g_m(\mathbf{x})/\partial \mathbf{x}$,
- $(\partial f(\mathbf{x})/\partial \mathbf{x} - \lambda_1 \partial g_1(\mathbf{x})/\partial \mathbf{x} + \dots - \lambda_m \partial g_m(\mathbf{x})/\partial \mathbf{x}) \cdot \mathbf{x} = 0$ ("complementarity" conditions),
- $\mathbf{x} \geq \mathbf{0}$,
- $b_1 - g_1(\mathbf{x}) \geq 0, \dots, b_m - g_m(\mathbf{x}) \geq 0$,
- $\lambda_1(b_1 - g_1(\mathbf{x})) + \dots + \lambda_m(b_m - g_m(\mathbf{x})) = 0$ ("complementarity" conditions),
- $\lambda_1, \dots, \lambda_m \geq 0$.

Two corollaries are implied by this definition.

Corollary 1

The differential Kuhn-Tucker conditions as necessary for a solution to an optimization problem can be derived from the Lagrange theorem using the lemma and substituting auxiliary non-negative variables $s_1 = b_1 - g_1(\mathbf{x})$, ..., $s_m = b_m - g_m(\mathbf{x})$. (The problem to be solved can be then written as $\text{Max}_{\mathbf{x}, \mathbf{s}} \{f(\mathbf{x}): \mathbf{g}(\mathbf{x}) + \mathbf{s} = \mathbf{b}, \mathbf{x} \geq \mathbf{0}, \mathbf{s} \geq \mathbf{0}\}$.)

Corollary 2

If the definition of the Lagrange function is generalized for optimization problems, then differential Kuhn-Tucker conditions can be written as:

- $\partial L/\partial \mathbf{x} \leq \mathbf{0}$,
- $(\partial L/\partial \mathbf{x}) \cdot \mathbf{x} = 0$ ("complementarity"),
- $\mathbf{x} \geq \mathbf{0}$,
- $\partial L/\partial \lambda_1 \geq 0, \dots, \partial L/\partial \lambda_m \geq 0$,
- $\lambda_1 \partial L/\partial \lambda_1 + \dots + \lambda_m \partial L/\partial \lambda_m = 0$ ("complementarity"),
- $\lambda_1 \geq 0, \dots, \lambda_m \geq 0$.

Definition (saddle point)

A function $h: \mathfrak{R}^{n+m} \rightarrow \mathfrak{R}$ has a non-negative saddle point $(\mathbf{y}^*, \mathbf{z}^*)$, if for every $\mathbf{y} \in \mathfrak{R}^n$, $\mathbf{y} \geq \mathbf{0}$ and every $\mathbf{z} \in \mathfrak{R}^m$, $\mathbf{z} \geq \mathbf{0}$ there is: $h(\mathbf{y}, \mathbf{z}^*) \leq h(\mathbf{y}^*, \mathbf{z}^*) \leq h(\mathbf{y}^*, \mathbf{z})$.

The rest of the lecture is devoted to proving the theorem.

Theorem (Kuhn-Tucker)

Solutions to $\text{Max}_{\mathbf{x}}\{f(\mathbf{x}): g_1(\mathbf{x}) \leq b_1, \dots, g_m(\mathbf{x}) \leq b_m, \mathbf{x} \geq \mathbf{0}\}$ can be characterized in the following way:

1. If $(\mathbf{x}^*, \lambda_1^*, \dots, \lambda_m^*)$ is a non-negative saddle point of the Lagrange function, then $\mathbf{x}^*$ is a solution to the original optimization problem;
2. If f is a concave function, and $g_1, \dots, g_m$ are convex functions, and if there is $\mathbf{x}^0 \geq \mathbf{0}$ such that $g_1(\mathbf{x}^0) < b_1, \dots, g_m(\mathbf{x}^0) < b_m$ (so-called Slater "constraint qualification"), then for every solution $\mathbf{x}^*$ to the original optimization problem, there exist $\lambda_1^*, \dots, \lambda_m^*$ such that $(\mathbf{x}^*, \lambda_1^*, \dots, \lambda_m^*)$ is a non-negative saddle point of the Lagrange function.

Proof:

- (1) Let $(\mathbf{x}^*, \lambda_1^*, \dots, \lambda_m^*)$ be a non-negative saddle point of the Lagrange function. Hence for all $\mathbf{x} \geq \mathbf{0}$:

$$f(\mathbf{x}) + \lambda_1^*(b_1 - g_1(\mathbf{x})) + \dots + \lambda_m^*(b_m - g_m(\mathbf{x})) \leq f(\mathbf{x}^*) + \lambda_1^*(b_1 - g_1(\mathbf{x}^*)) + \dots + \lambda_m^*(b_m - g_m(\mathbf{x}^*))$$

and for all $\lambda_1^*, \dots, \lambda_m^* \geq 0$:

$$f(\mathbf{x}^*) + \lambda_1^*(b_1 - g_1(\mathbf{x}^*)) + \dots + \lambda_m^*(b_m - g_m(\mathbf{x}^*)) \leq f(\mathbf{x}^*) + \lambda_1(b_1 - g_1(\mathbf{x}^*)) + \dots + \lambda_m(b_m - g_m(\mathbf{x}^*))$$

The second inequality can be written as

$$(\lambda_1 - \lambda_1^*)(b_1 - g_1(\mathbf{x}^*)) + \dots + (\lambda_m - \lambda_m^*)(b_m - g_m(\mathbf{x}^*)) \geq 0 \text{ for } \lambda_1, \dots, \lambda_m \geq 0.$$

Since $\lambda_1, \dots, \lambda_m$ can be arbitrarily large, the following inequalities should hold:

$$g_1(\mathbf{x}^*) \leq b_1, \dots, g_m(\mathbf{x}^*) \leq b_m.$$

On the other hand, by substituting $\lambda_1 = \dots = \lambda_m = 0$ we will get $\lambda_1^*(b_1 - g_1(\mathbf{x}^*)) + \dots + \lambda_m^*(b_m - g_m(\mathbf{x}^*)) \leq 0$. However, since every number in the parentheses is non-negative, the inequality is simply an equality:

$$\lambda_1^*(b_1 - g_1(\mathbf{x}^*)) + \dots + \lambda_m^*(b_m - g_m(\mathbf{x}^*)) = 0.$$

By substituting this equality to the first inequality of the saddle point definition we get:

$$f(\mathbf{x}^*) \geq f(\mathbf{x}) + \lambda_1^*(b_1 - g_1(\mathbf{x})) + \dots + \lambda_m^*(b_m - g_m(\mathbf{x}))$$

for all $\mathbf{x} \geq \mathbf{0}$. If additionally $\mathbf{x}$ satisfies the constraints, then $f(\mathbf{x}^*) \geq f(\mathbf{x})$ (since $\lambda_1^*, \dots, \lambda_m^* \geq 0$), which completes this part of the proof.

- (2) Let $\mathbf{x}^*$ be a solution to the optimization problem, so $\mathbf{x}^* \geq \mathbf{0}$, $\mathbf{g}(\mathbf{x}^*) \leq \mathbf{b}$ and $f(\mathbf{x}^*) \geq f(\mathbf{x})$ for all $\mathbf{x}$ satisfying $\mathbf{x} \geq \mathbf{0}$ and $\mathbf{g}(\mathbf{x}) \leq \mathbf{b}$ ($\mathbf{g} = [g_1, \dots, g_m]'$, $\mathbf{b} = [b_1, \dots, b_m]'$). Let us define two sets in $\mathfrak{R}^{m+1}$: $A = \{(a_0, \mathbf{a}): \exists \mathbf{x} \geq \mathbf{0} [a_0 \leq f(\mathbf{x}) \text{ and } \mathbf{a} \leq \mathbf{b} - \mathbf{g}(\mathbf{x})]\}$ and $C = \{(c_0, \mathbf{c}): c_0 > f(\mathbf{x}^*) \text{ and } \mathbf{c} > \mathbf{0}\}$, where $a_0, c_0 \in \mathfrak{R}$, $\mathbf{a}, \mathbf{c} \in \mathfrak{R}^m$ (column vectors). The set C is convex as an interior of a convex set. From the concavity of f and convexity of $\mathbf{g}$ we know that the set A is also convex (this is fairly easy to verify). As $\mathbf{x}^*$ is a solution, then the sets are disjoint.

Thus by the separation theorem there exists a non-zero vector $(y_0, \mathbf{y}) \in \mathfrak{R}^{m+1}$ ($\mathbf{y}$ is a row vector) such that for all $(a_0, \mathbf{a}) \in A$ and all $(c_0, \mathbf{c}) \in C$ we have: $y_0 a_0 \leq y_0 c_0$ and $\mathbf{y} \cdot \mathbf{a} \leq \mathbf{y} \cdot \mathbf{c}$. From the definition of sets A and C $(y_0, \mathbf{y}) \geq \mathbf{0}$ (otherwise, by substituting negative $\mathbf{a}$ and a_0 and positive $\mathbf{c}$ and c_0 we get a contradiction with the separation theorem). The point $(f(\mathbf{x}^*), \mathbf{0})$ belongs to the border of the set C , so all (non-strict) inequalities that hold for C , hold for this point too. Therefore: $y_0 a_0 \leq y_0 f(\mathbf{x}^*)$ and $\mathbf{y} \cdot \mathbf{a} \leq \mathbf{y} \cdot \mathbf{0}$ (i.e. $\mathbf{y} \cdot \mathbf{a} \leq 0$).

Summing up these two inequalities yields: $y_0 a_0 + \mathbf{y} \cdot \mathbf{a} \leq y_0 f(\mathbf{x}^*)$. In particular, by substituting $a_0 = f(\mathbf{x})$ and $\mathbf{a} = \mathbf{b} - \mathbf{g}(\mathbf{x})$, we get $y_0 f(\mathbf{x}) + \mathbf{y} \cdot (\mathbf{b} - \mathbf{g}(\mathbf{x})) \leq y_0 f(\mathbf{x}^*)$. Now we will prove that $y_0 > 0$. From the constraint qualification (there exists $\mathbf{x}^0 \geq \mathbf{0}$, for which $\mathbf{b} - \mathbf{g}(\mathbf{x}^0) > \mathbf{0}$ for all coordinates) and

from the non-negativity of the vector $\mathbf{y} \neq \mathbf{0}$ we get $\mathbf{y} \cdot (\mathbf{b} - \mathbf{g}(\mathbf{x}^0)) > 0$. If $y_0 = 0$, then the inequality $y_0 f(\mathbf{x}) + \mathbf{y} \cdot (\mathbf{b} - \mathbf{g}(\mathbf{x})) \leq y_0 f(\mathbf{x}^*)$ could not have been satisfied.

Knowing that $y_0 > 0$, then both sides of this inequality can be divided by y_0 . Then one gets $f(\mathbf{x}) + \mathbf{y}^* \cdot (\mathbf{b} - \mathbf{g}(\mathbf{x})) \leq f(\mathbf{x}^*)$, where $\mathbf{y}^* = \mathbf{y}/y_0$. By substituting $\mathbf{x} = \mathbf{x}^*$ one gets $\mathbf{y}^* \cdot (\mathbf{b} - \mathbf{g}(\mathbf{x}^*)) \leq 0$. Given the fact that all elements of this sum are non-negative (as products of non-negative numbers), this in fact implies that $\mathbf{y}^* \cdot (\mathbf{b} - \mathbf{g}(\mathbf{x}^*)) = 0$. The pair $(\mathbf{x}^*, \mathbf{y}^*)$ is the non-negative saddle point of the Lagrange function (one needs to substitute $\lambda_1 = y_1, \dots, \lambda_m = y_m$), since it is easy to check that $L(\mathbf{x}, \mathbf{y}^*) \leq L(\mathbf{x}^*, \mathbf{y}^*) \leq L(\mathbf{x}^*, \mathbf{y})$, i.e. $f(\mathbf{x}) + \mathbf{y}^* \cdot (\mathbf{b} - \mathbf{g}(\mathbf{x})) \leq f(\mathbf{x}^*) + \mathbf{y}^* \cdot (\mathbf{b} - \mathbf{g}(\mathbf{x}^*)) \leq f(\mathbf{x}^*) + \mathbf{y} \cdot (\mathbf{b} - \mathbf{g}(\mathbf{x}^*))$.

Note

In the Kuhn-Tucker theorem the concavity/convexity assumptions for f and $\mathbf{g}$, as well as constraint qualifications for $\mathbf{g}$ can be weakened by assuming, respectively, quasi-concavity, quasi-convexity and constraint qualifications for those constraints only whose left-hand side is non-linear (this modified constraint qualification is called Martos condition).

Questions and answers to lecture 15

15.1 The classical constrained maximisation problem reads:

$$\text{Max}_{\mathbf{x}} \{f(\mathbf{x}) : g_1(\mathbf{x}) = b_1, \dots, g_m(\mathbf{x}) = b_m\},$$

while the mathematical programming (optimisation) is:

$$\text{Max}_{\mathbf{x}} \{f(\mathbf{x}) : g_1(\mathbf{x}) \leq b_1, \dots, g_m(\mathbf{x}) \leq b_m, \mathbf{x} \geq \mathbf{0}\}.$$

Please provide an example of an economic issue which cannot be modelled as the classical maximisation problem.

Consumer optimal choice may provide such an example. Let us assume that a consumer's utility function reads $u(x_1, x_2) = x_1^2 x_2^3$, the prices of goods one and two are: $p_1 = 1, p_2 = 2$, the consumer has the amount of money $m = 30$, and there is an additional requirement (caused, say, by weight or storage capacity): $x_1 \leq 20$. The problem can be written as

$$\text{Max}_{x_1, x_2} \{x_1^2 x_2^3 : x_1 + 2x_2 \leq 30, x_1 \leq 20, x_1, x_2 \geq 0\}.$$

The solution is $x_1^* = 12, x_2^* = 9$. The first inequality ($x_1 + 2x_2 \leq 30$) could be substituted with equality ($x_1 + 2x_2 = 30$), but the second one ($x_1 \leq 20$) could not (it is satisfied as a strict inequality, $x_1^* < 20$).

15.2 The definition of a quasi-convex function $f: \mathfrak{R}^n \rightarrow \mathfrak{R}$ is: for every $z \in \mathfrak{R}$, the set $\{\mathbf{x} \in \mathfrak{R}^n : f(\mathbf{x}) \leq z\}$ is convex. Please prove that every convex function is quasi-convex.

One of the definitions of convexity of functions says that for every $z \in \mathfrak{R}$, the set $\{\mathbf{x} \in \mathfrak{R}^n : f(\mathbf{x}) \leq z\}$ is convex which could complete the proof. Nevertheless let us apply the convexity definition that most students are more likely to be familiar with: a function $f: \mathfrak{R}^n \rightarrow \mathfrak{R}$ is convex if and only if for any $\lambda \in [0, 1]$, $f(\lambda a + (1-\lambda)b) \leq \lambda f(a) + (1-\lambda)f(b)$. Let z be the higher of the two numbers: $f(a)$ and $f(b)$. In order to prove that the set $C = \{\mathbf{x} \in \mathfrak{R}^n : f(\mathbf{x}) \leq z\}$ is convex, we need to prove that when $a \in C$ and $b \in C$, then $\lambda a + (1-\lambda)b \in C$. If $a, b \in C$, then $f(a) \leq z$, and $f(b) \leq z$. If f is a convex function then $f(\lambda a + (1-\lambda)b) \leq \lambda f(a) + (1-\lambda)f(b) \leq \lambda z + (1-\lambda)z = z$ which completes the proof.

15.3 Give an example of a quasi-convex function which is not convex.

The function illustrated in the following (blue) graph has the desired properties. It is not convex (because the segment linking $(x_1, f(x_1))$ and $(x_2, f(x_2))$ is below the graph), but the set $C = \{x \in \mathbb{R} : f(x) \leq z\}$ is convex.

