

Decision Theory

IE 5545: Decision Analysis
Prof. Krishnamurthy Iyer*

May 2, 2026

In these notes, we will develop a framework to represent and analyze static decision problems (with or without uncertainty) faced by a single *rational* decision maker. This framework will serve as the building block to analyze more complex decision problems that extend over many periods and/or involve multiple agents.

1 Structure of the decision problem

The structure of a decision problem can be characterized by four components:

1. the set of actions (or *strategies*¹) available to the decision maker;
2. the set of outcomes that can result from the decision;
3. how each action determines the realized outcome (or “rules”); and
4. the decision maker’s preference over different outcomes.

Let $\mathcal{S} = \{s_1, s_2, \dots, s_n\}$ denote the set of all available actions to the decision maker. Similarly, let $\mathcal{O} = \{o_1, o_2, \dots, o_m\}$ denote the set of all outcomes that can arise from the decision. Here, we have implicitly made the assumption that both the sets $\mathcal{S}$ and $\mathcal{O}$ are finite. The framework carries over naturally to settings where there are infinitely many actions or outcomes (for example, in the case where a firm is deciding a price for a product, the set $\mathcal{S}$ can be modeled as an interval of real numbers.)

Next, the rules underlying a decision problem capture how each action influences the realized outcome. First, consider a decision problem that does not involve any uncertainty or randomness. Then, each action $s \in \mathcal{S}$ leads to a definite deterministic outcome in $\mathcal{O}$. Denote the outcome realized after choosing action $s \in \mathcal{S}$ as $x(s)$. Note that $x(\cdot)$ is a function that takes as input an action $s \in \mathcal{S}$ and gives as output an outcome $x(s) \in \mathcal{O}$. For the simple setting of a static decision problem with no uncertainty, the function $x(\cdot)$ represents the “rules” of the problem.

In a decision problem with uncertainty, the outcome resulting from an action may not be deterministic. Instead, the outcome may be stochastically determined. To capture this, we need to consider distributions over outcomes.

A *distribution* over outcomes is given by a tuple $(p_1, p_2, \dots, p_m)$, where $p_i \geq 0$ for each i , and $\sum_{i=1}^m p_i = 1$. Under this distribution, outcome o_1 is realized with probability p_1 , outcome o_2 is realized with probability p_2 , and so on. We let $\Delta(\mathcal{O})$ denote the set of all distributions over $\mathcal{O}$:

$$\Delta(\mathcal{O}) = \left\{ (p_1, p_2, \dots, p_m) : p_i \geq 0, \text{ for all } i, \text{ and } \sum_{i=1}^m p_i = 1 \right\}.$$

In the case of a decision problem with uncertainty, we can now specify the “rules” by specifying a function $\bar{p}(\cdot)$ that, for each input $s \in \mathcal{S}$, gives as output a distribution $\bar{p}(s) \in \Delta(\mathcal{O})$ over outcomes.

*©2026 Krishnamurthy Iyer

¹For the time being, since we are focusing on static settings, we will use the terms “actions” and “strategies” interchangeably. Later, while studying dynamic settings involving multiple stages, we will be more careful and formally differentiate between the two terms.

Summarizing, depending on the decision problem, the rules are specified either by a function $x : \mathcal{S} \rightarrow \mathcal{O}$, or a function $\bar{p} : \mathcal{S} \rightarrow \Delta(\mathcal{O})$. (Note that we can always choose the latter representation even in the case of a deterministic decision problem, by defining $\bar{p}(s)$ to be the distribution where the outcome $x(s)$ gets probability 1 and all other outcomes get probability 0.)

Having described formally the set of actions, the set of outcomes and the rules underlying the decision problem, we are left with formalizing the notion of a decision maker's preferences. We focus on this in the rest of these notes.

1.1 Preferences

The decision maker has a preference over which outcome is realized as a consequence of her decision. This preference can be represented as a relation between two outcomes that specifies which outcome between the two the decision maker prefers.

Formally, the decision maker's preference over the set of outcomes $\mathcal{O}$ can be represented by an "order relation" $\preceq$ over the set of outcomes $\mathcal{O}$. For any two outcomes $a, b \in \mathcal{O}$, we let the relation $a \preceq b$ mean that the decision maker prefers outcome b at least as much as outcome a .

Note that this "weak preference" relation can be used to describe other kinds of preferences. For example, if $a \preceq b$, but not $b \preceq a$, then the decision maker prefers b at least as much as a , but not vice versa. In other words, the decision maker strictly prefers b over a , which we denote by $a \prec b$. Similarly, if $a \preceq b$ and $b \preceq a$, then the decision maker prefers the two outcomes equally, which we denote by $a \sim b$. (Note that this does not mean that outcomes a and b are the same, just that the decision maker is indifferent between the two outcomes.)

What properties must the relation $\preceq$ satisfy if it is to represent a "rational" decision maker's preferences? This is an important question, since currently, we have imposed no restriction on $\preceq$.

One condition that is natural to require is that a rational decision maker must always be able to compare any two outcomes. Without this assumption, one might end up in a situation where the decision maker cannot compare two outcomes, and hence cannot choose between the two. Since the decision maker *must* eventually decide, this gives us the first axiom:

Axiom 1 (Completeness). *For any $a, b \in \mathcal{O}$, we have either $a \preceq b$, or $b \preceq a$, or both.*

A preference relation that satisfies the above axiom is said to be *complete*. (Another term is "total".) Note that this condition does not mean that the decision maker has to strictly prefer either a or b . We allow the decision maker to be indifferent between the two outcomes. We only require that the decision maker should be able to compare any two outcomes.

A second natural condition to impose is that a rational decision maker's preference should not involve cycles. For example, it must not be the case that a rational decision maker strictly prefers apples over bananas, strictly prefers bananas over carrots, and strictly prefers carrots over apples. (If such a scenario holds, the decision maker's preferences can be exploited by others — how?)

To avoid such cycles, we impose the following axiom on preferences:

Axiom 2 (Transitivity). *For any $a, b, c \in \mathcal{O}$, if $a \preceq b$ and $b \preceq c$, then we must have $a \preceq c$.*

Thus, a rational decision maker's preference is specified by stating an order relation $\preceq$ on $\mathcal{O}$ that is *complete* and *transitive*.

We next look at preferences that can be represented by a *utility function*:

Definition 1 (Utility representation). *Given any preference $\preceq$, we say a utility function $u : \mathcal{O} \rightarrow \mathbb{R}$ represents $\preceq$, if we have*

$$a \preceq b, \quad \text{if and only if,} \quad u(a) \leq u(b) \quad \text{for all } a, b \in \mathcal{O}.$$

It is quite straightforward to see that if a preference $\preceq$ can be represented by a utility function $u : \mathcal{O} \rightarrow \mathbb{R}$, then it must be complete and transitive. It turns out that the converse is also true. Namely, if a preference $\preceq$ is complete and transitive, then there exists a utility function $u : \mathcal{O} \rightarrow \mathbb{R}$ that represents it. (Showing this is left as an exercise).

We note here that a preference $\preceq$ can be represented by many different utility function. As an example, suppose $\mathcal{O} = \{o_1, o_2, o_3\}$, and the preference $\preceq$ is given by $o_1 \prec o_2 \prec o_3$. Then, one utility function u that represents this preference is given by $u(o_1) = 1, u(o_2) = 2, u(o_3) = 3$. However, the same preference can be represented by the utility function v where $v(o_1) = 2, v(o_2) = 3, v(o_3) = 5$. The reason here is that the preference $\preceq$ only states which outcome is preferred among the two, not how much one outcome is preferred over another. By using different utility functions, we can capture the notion of how much one outcome is preferred over another. It turns out that this notion is important when considering decision problems where randomness plays a role, as we discuss next.

When the decision problem involves randomness, the actions can lead to distributions over outcomes. In such a scenario, in order for a decision maker to choose among the actions, not only should she be able to compare outcomes, she should be also able to compare two distributions over outcomes. However, note that currently the decision maker's preference (represented either by a relation $\preceq$ or by an utility function $u : \mathcal{O} \rightarrow \mathbb{R}$) can only compare outcomes. To compare distributions over outcomes, we need a preference over the set $\Delta(\mathcal{O})$.

A naive way to do this would be to ask for the decision maker's preferences over the entire $\Delta(\mathcal{O})$. However, even for a game with finitely many outcomes, the set $\Delta(\mathcal{O})$ is infinite, and it is an unreasonable and arduous task to specify this preference.

An easier approach would be to take a decision maker's preference over $\mathcal{O}$ (represented by a utility function) and use it to infer her preferences over $\Delta(\mathcal{O})$. The *expected utility framework* does precisely this, as we describe below.

Suppose a decision maker's preference over $\mathcal{O}$ is given by a utility function $u : \mathcal{O} \rightarrow \mathbb{R}$. For any distribution $\bar{p} = (p_1, p_2, \dots, p_m) \in \Delta(\mathcal{O})$, we denote the *expected utility* with respect to $\bar{p}$ to as follows:

$$\mathbf{E}[u(X) | X \sim \bar{p}] = \sum_{i=1}^m p_i u(o_i).$$

Note that to compute the expected utility with respect to a distribution $\bar{p}$, it is sufficient to know the decision maker's utility function over $\mathcal{O}$ (a much smaller, often a finite, set).

The *expected utility framework* then suggests that to compare two distributions $\bar{p}$ and $\bar{q}$, we just compare the expected utility with respect to the two distributions. Formally, have the following definition:

Definition 2 (Expected utility representation). *We say a utility function $u : \mathcal{O} \rightarrow \mathbb{R}$ represents a decision maker's preference over $\Delta(\mathcal{O})$, if for any two distributions $\bar{p}, \bar{q} \in \Delta(\mathcal{O})$, we have*

$$\bar{p} \preceq \bar{q}, \quad \text{if and only if,} \quad \mathbf{E}[u(X) | X \sim \bar{p}] \leq \mathbf{E}[u(X) | X \sim \bar{q}],$$

When the decision maker's preference $\preceq$ is represented by an utility function, one can say that the decision maker prefers a distribution $\bar{q}$ over distribution $\bar{p}$ if and only if the expected utility with respect to $\bar{q}$ is more than the expected utility with respect to $\bar{p}$. In this scenario, if one allows the decision maker to choose her most preferred distribution, she will choose the distribution that maximizes her expected utility. Hence, in this case, the decision maker is said to be an "expected utility maximizer".

It turns out that there are very good theoretical reasons for using the expected utility framework to represent a "rational" decision maker's preferences over distributions. This follows from the celebrated work of John von Neumann and Oskar Morgenstern in 1947. (We will not go into details here.) Due to this, the utility function that represents a decision maker's preference is sometimes called the decision maker's von Neumann-Morgenstern utility function.

The (expected) utility framework is very useful in two ways:

1. We no longer have to work with outcomes, distribution over outcomes, etc. Instead, one can work with expected utilities, which are just real numbers.
2. Finding the *best* outcome (or the best distribution over outcomes) for a decision maker, i.e., an outcome (or a distribution) she prefers at least as much as any other outcome (distribution, resp.), reduces to the problem of finding an outcome that has the highest (expected) utility. This latter problem can be solved using well studied techniques from optimization.

Due to these advantages, we will henceforth make the assumption that a decision maker's preference over $\Delta(\mathcal{O})$ can be represented by a utility function. We will call such preferences "rational".

Example 1. Consider a decision problem faced by a student the night before the final exam. The student has to decide whether to stay awake all night and study for the exam, or to get a good night's sleep before the exam. Denote the action to stay awake and study by a , and the action to sleep by s . Then the set of actions is given by $\mathcal{S} = \{a, s\}$. The outcome relevant to the student is the grade she obtains in the exam, which can be one of $\mathcal{O} = \{A, B, C\}$. The student's preference over the outcomes is given by $A \succ B \succ C$, implying she strictly prefers A over B over C .

If she chooses action s , the outcome is realized according to a distribution $\bar{p}(s) = (0.6A, 0.3B, 0.1C)$. If she chooses action a , the outcomes is realized according to a distribution $\bar{p}(a) = (0.7A, 0.1B, 0.2C)$. Which action should the student choose?

To answer this question, we will use the expected utility framework. For this, we first need a utility representation of the student's preferences. One possible utility function that is consistent with the student's preference over $\mathcal{O}$, is given by $u(A) = 100, u(B) = 50, u(C) = 20$. Then, the expected utility obtained by the student on choosing action s is given by

$$\begin{aligned} \mathbb{E}[u(X) | X \sim \bar{p}(s)] &= 0.6u(A) + 0.3u(B) + 0.1u(C) \\ &= 0.6(100) + 0.3(50) + 0.1(20) \\ &= 77. \end{aligned}$$

Similarly, the expected utility obtained by the student on choosing action a is given by

$$\begin{aligned} \mathbb{E}[u(X) | X \sim \bar{p}(a)] &= 0.7u(A) + 0.1u(B) + 0.2u(C) \\ &= 0.7(100) + 0.1(50) + 0.2(20) \\ &= 79. \end{aligned}$$

Thus, since choosing action a gets the student higher expected utility, the student should choose to stay awake and study for the final.

However, note that here we could have chosen a different utility function consistent with the student's preference over $\mathcal{O}$. For example, the utility function $v(A) = 100, v(B) = 90, v(C) = 60$ is also consistent with the student's preference over $\mathcal{O}$. With this utility function, we see the expected utility for the two actions as

$$\begin{aligned} \mathbb{E}[v(X) | X \sim \bar{p}(s)] &= 0.6v(A) + 0.3v(B) + 0.1v(C) \\ &= 0.6(100) + 0.3(90) + 0.1(60) \\ &= 93 \\ \mathbb{E}[v(X) | X \sim \bar{p}(a)] &= 0.7v(A) + 0.1v(B) + 0.2v(C) \\ &= 0.7(100) + 0.1(90) + 0.2(60) \\ &= 91. \end{aligned}$$

Thus, with the utility function v , we see that the student prefers to get a good night's sleep before the exam. (It can be checked that for the utility function w where $w(A) = 100, w(B) = 90, w(C) = 80$, the student obtains the same expected utility for the two actions, and hence is indifferent between them.)

In this example, the two actions differ in their "risk" profile. The action a has a higher chance of getting A , but also a higher chance of getting C in comparison to action s . To know how a student prefers these two actions, one must know the student's attitude toward risk – is she a risk taker or is she conservative towards risk? The specific choice of the utility function encodes this risk attitude of the student.

The above example highlights that although many different utility function can represent a decision maker's preference over $\mathcal{O}$, they lead to different conclusions regarding her preferences over $\Delta(\mathcal{O})$. This is because when comparing outcomes, it is enough to know which outcome is preferred among any two, whereas while comparing distributions over outcomes, one must exactly know *how much one outcome is preferred over another*. Different utility functions encode this information differently, and hence lead to different conclusions when comparing distributions.

2 Decision theory framework

Summarizing our discussion until now, a decision problem (with or without uncertainty) can precisely specified by stating

1. A set of actions $\mathcal{S} = \{s_1, s_2, \dots, s_n\}$
2. A set of outcomes $\mathcal{O} = \{o_1, o_2, \dots, o_m\}$
3. The rules $\bar{p} : \mathcal{S} \rightarrow \Delta(\mathcal{O})$ (or $x : \mathcal{S} \rightarrow \mathcal{O}$ in the case without uncertainty) specifying how actions lead to outcomes, and
4. The utility function $u : \mathcal{O} \rightarrow \mathbb{R}$ representing the decision maker's preferences over $\Delta(\mathcal{O})$.

In this representation, we have made two assumptions (first one explicit and the second one implicit), that

1. the decision maker is “rational” — this allowed us to represent her preferences over $\Delta(\mathcal{O})$ by a utility function $u : \mathcal{O} \rightarrow \mathbb{R}$ using the expected utility framework; and
2. the decision maker *knows* the above structure of the decision problem.

The second assumption is referred to as the *complete information* setting. We will later expand our framework to settings where the decision maker may not completely know the structure, i.e., settings with “incomplete” information.

With these two assumptions, it follows that the decision maker decides by choosing the action $s \in \mathcal{S}$ that leads to the highest expected utility. Such an action is referred to as an optimal action for the decision problem under consideration. (There might be multiple actions that are optimal for a decision problem, all achieving the same expected utility.)

2.1 Payoff function

Having developed this framework, we can slightly simplify it, especially during the analysis of the decision problem to identify an optimal action.

Suppose a decision maker's preferences over the set of outcomes $\mathcal{O}$ is given by a utility function $u : \mathcal{O} \rightarrow \mathbb{R}$. Then, when the decision has no randomness, given an action $s \in \mathcal{S}$, the outcome is $x(s)$, and the utility obtained by the decision maker is equal to $u(x(s))$. Thus, we can define a function $\pi : \mathcal{S} \rightarrow \mathbb{R}$ by $\pi(s) = u(x(s))$, that takes as input an action s , and gives as output the utility, $u(x(s))$, obtained on choosing that action s . We refer to this function π as the decision maker's payoff function.

Analogously, when the decision problem involves randomness, an action $s \in \mathcal{S}$ leads to a distribution over outcome $\bar{p}(s)$, with expected utility equal to $\mathbb{E}[u(X)|X \sim \bar{p}(s)]$. Again, we can define a payoff function $\pi : \mathcal{S} \rightarrow \mathbb{R}$ by $\pi(s) = \mathbb{E}[u(X)|X \sim \bar{p}(s)]$, that takes as input an action $s \in \mathcal{S}$ and gives as output the (expected) utility the decision maker obtains on choosing that action.

The payoff function is derived from the decision maker's preferences (represented by the utility function), and the rules of the game (represented by $x(\cdot)$ or $\bar{p}(\cdot)$). Note that the payoff function is defined over the set of actions $\mathcal{S}$, whereas the utility function is defined over the set of outcomes $\mathcal{O}$.

Given our discussion until now, we can succinctly represent a decision problem (with or without uncertainty) as follows:

1. A set of actions $\mathcal{S} = \{s_1, s_2, \dots, s_n\}$
2. A payoff function $\pi : \mathcal{S} \rightarrow \mathbb{R}$, that represents the payoff a decision maker obtains for choosing different actions.