

Enrico Mattia Salonia

Game Theory - Spring 2026

The Ultimatum Game

Ann has 10 euros. She can choose whether to give all of them to Bob, split them between the two, or give nothing to Bob. Bob can later accept Ann's offer or reject it. If Bob accepts, Ann's offer is implemented. If he rejects, they both get nothing. The interaction is represented in Figure 1.

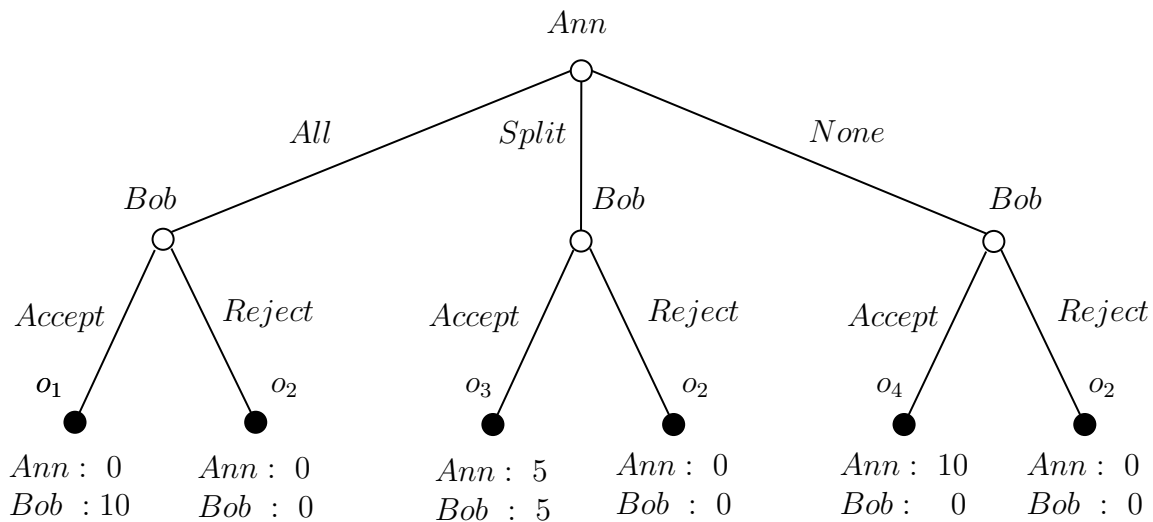


Figure 1: An extensive form for Ann and Bob.

The picture as an extensive form game with perfect information

This interaction is called the **Ultimatum Game**. Ann makes an offer to Bob, who can either accept or reject it; if he accepts, the offer is implemented; if he rejects, both players receive nothing. Let us study how elements of the definition of extensive form with perfect information map to this example.

The picture is a **rooted directed tree**.

- The *root* is the initial node, where *Ann* can choose what to do.
- Each arrow is a *directed edge* and points from a node to one of its *successors*.
- All the white dots are *decision nodes*. The black dots at the bottom are the *terminal nodes*.

As a **finite extensive form**, the tree comes with the following objects.

- The rooted directed tree defined above.
- The set of *players* is $I = \{Ann, Bob\}$.
- The *player function* assigns to each nonterminal node the player who moves there: *Ann* moves at the root, and *Bob* moves at the three nodes on the second level.

- For each nonterminal node, the set of *actions* available is given by the labels on the outgoing edges from the white dot:
 - at the root, where *Ann* moves: the available actions are the three offers *All*, *Split*, and *None*;
 - at each of *Bob*'s nodes: the available actions are *Accept* and *Reject*.
- The set of outcomes is $O = \{o_1, o_2, o_3, o_4\}$, where

$$o_1 = (0, 10), \quad o_2 = (0, 0), \quad o_3 = (5, 5), \quad o_4 = (10, 0),$$

with the first component giving *Ann*'s monetary outcome and the second component giving *Bob*'s monetary outcome.

- Each terminal node is assigned an *outcome*, shown next to the corresponding black dot.

This is a finite extensive form *with perfect information* because at each decision node the moving player knows exactly where in the tree they are. To complete the description of the game, we should add preferences. The following table records examples of preference relations.

Player	Preference ordering
Ann (selfish)	$o_4 \succ o_3 \succ o_1 \sim o_2$
Ann (fair, benevolent)	$o_3 \succ o_4 \succ o_1 \succ o_2$
Bob (selfish)	$o_1 \succ o_3 \sim o_4 \sim o_2$
Bob (fair, spiteful)	$o_3 \succ o_1 \succ o_2 \succ o_4$

We now have a **finite extensive game with perfect information**.

- The finite extensive form with perfect information we have above.
- A preference $\succsim_i$ over the set of outcomes O , one for each player $i \in I$.

Backward Induction

Say *Ann* is selfish and *Bob* is fair and spiteful, we can represent their preferences by the following utility functions.

Utility function	o_1	o_2	o_3	o_4
U_{Ann} (selfish)	2	2	3	4
U_{Bob} (fair, spiteful)	3	2	4	1

With these preferences, the game looks as in Figure 2.

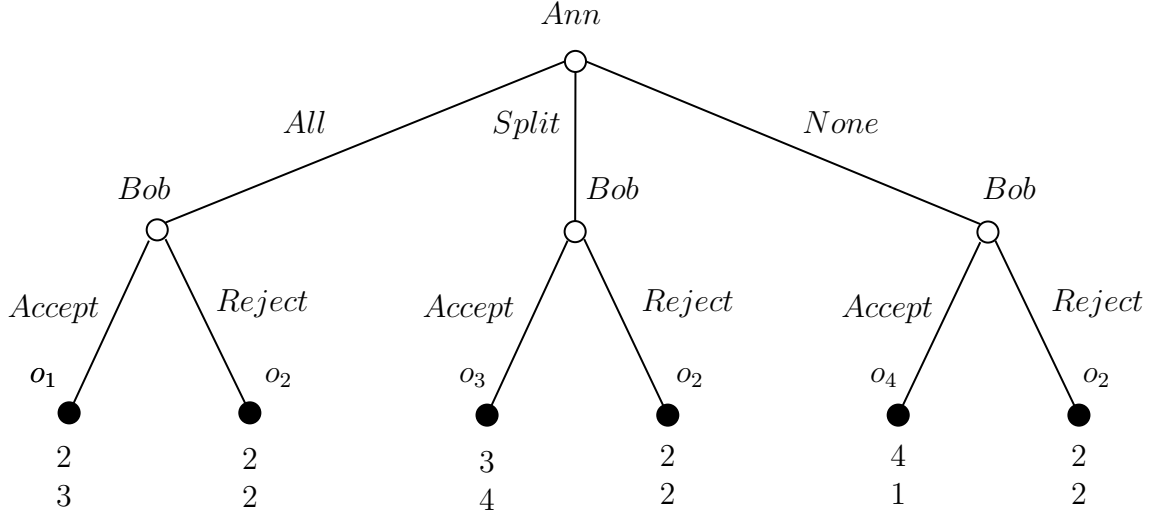


Figure 2: The ultimatum game if Ann is selfish and Bob is fair and spiteful.

Players could reason as follows. Start from each decision node immediately before a terminal node, and consider what the player assigned to that decision node would do. Then move to the decision nodes preceding the last decision node, and consider what the players assigned to those nodes would do, assuming they can predict what subsequent players will do. Such reasoning is called **Backward Induction**. Consider Bob's choice after Ann chooses to give *All* to him. He prefers o_1 to o_2 , and therefore chooses to *Accept*. After Ann chooses to *Split*, he chooses to *Accept*, because he prefers o_3 to o_2 . After Ann chooses *None*, Bob chooses to *Reject*, since he prefers o_2 to o_4 . The colored edges in Figure 3 represent Bob's choices.

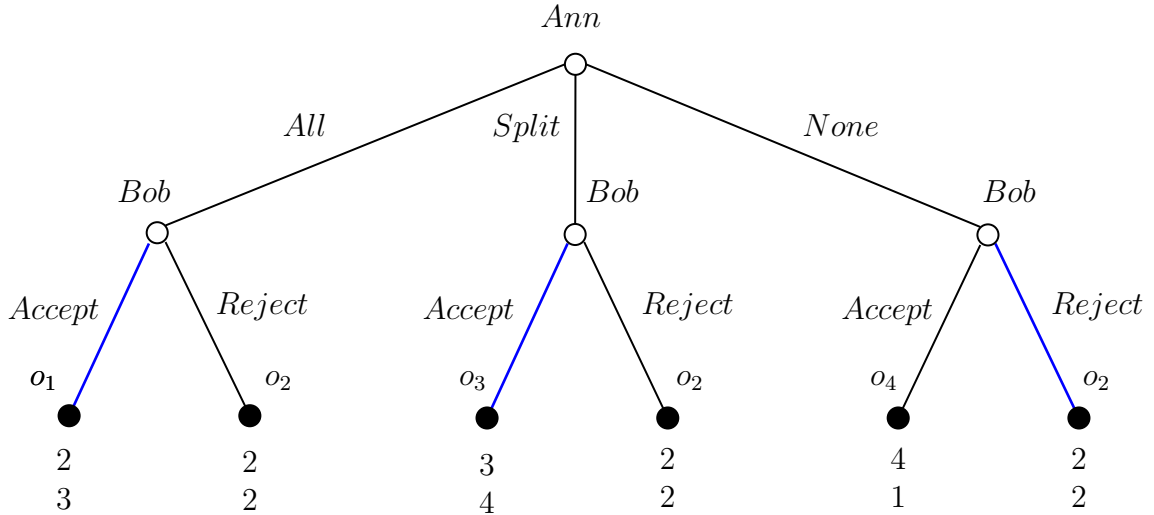


Figure 3: Bob's choice at each of his decision nodes in blue.

Ann can then anticipate Bob's choices. If she chooses *All*, Bob will induce outcome o_1 ; if she chooses *Split*, Bob will induce outcome o_3 ; and if she chooses *None*, Bob will induce outcome o_2 . Since she prefers o_3 to o_1 and o_2 , she chooses *Split*. The entire path we just described is represented in Figure 4.

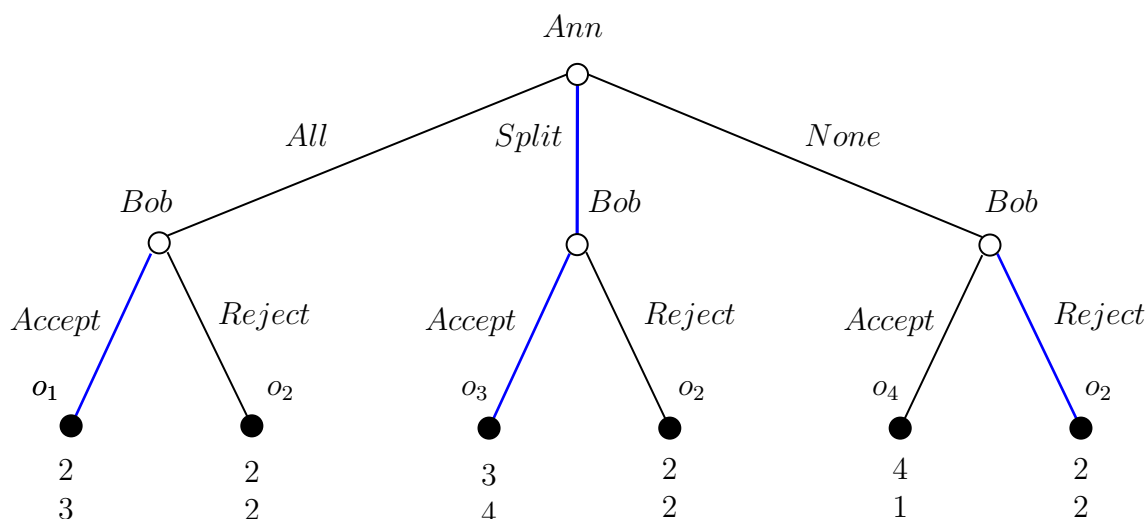


Figure 4: Ann's and Bob's choices at each decision node in blue.

What is the moral of the story? Bob's spitefulness induces Ann to split the money, even if she is selfish.

Exercises

"Human nature is too weak to acquire an art of which it has no experience."
(Plato, 1881, p. 111)

Exercise 1. Repeat the Backward Induction analysis under the assumption that both Ann and Bob are selfish. Do you notice any issue? Write the game in strategic form and find the Nash Equilibria.

Exercise 2. In this exercise, you study the **Dictator Game**. It is a variant of the ultimatum game in which the second player (Bob) cannot act; thus, the first player's (Ann) choice determines the outcome.

1. Describe this game as a finite extensive form and draw a picture of it.
2. Suppose you are a researcher. You do not know Ann's preferences, but you observe her choices in this game. What, if anything, can you infer about her preferences?
3. Elaborate on how, in general, observing players' choices in a game frame or extensive form allows researchers to infer their preferences.

References

Plato. (1881). *The theaetetus of plato* (B. H. Kennedy, Ed. & Trans.) [With translation and notes]. Cambridge University Press.

Enrico Mattia Salonia

Game Theory - Spring 2026

The Trust Game

Ann has 10 euros and can give any part of it to Bob. Bob receives three times what Ann gives him. He can then decide how much of what he has to give back to Ann.

Denote by a the amount Ann gives to Bob and by b the amount Bob gives back to Ann. Their monetary rewards at the end of the game are as follows:

$$m_A(a, b) = 10 - a + b \quad m_B(a, b) = 3a - b.$$

Exercise 1. Explain why this interaction constitutes a finite extensive-form game with perfect information.

From now on, assume Ann is selfish and greedy: she only wants to maximize her monetary reward. Bob, instead, might have one of two different preference types. He could be a *Beckerian altruist* (Becker, 1974), i.e., he cares about Ann's monetary reward; in that case, his preferences can be represented by the following utility function:

$$U_B(a, b) = m_B(a, b) + \alpha m_A(a, b).$$

Here, α is Bob's *altruism* parameter, which captures how much he cares about Ann's monetary reward. Alternatively, Bob might be *inequity averse* à la Fehr and Schmidt (1999), i.e., he prefers monetary rewards not to be too unequal. In that case, his utility function is:

$$U_B(a, b) = m_B(a, b) - \alpha \max(m_A(a, b) - m_B(a, b), 0) - \beta \max(m_B(a, b) - m_A(a, b), 0).$$

The parameter α represents Bob's *envy*, how much he dislikes Ann having more, and the parameter β represents Bob's *guilt*, how much he dislikes having more than Ann.

Exercise 2. Find the backward-induction solution under the assumption that Bob is a *Beckerian altruist*. How does the outcome depend on the altruism parameter α ?

Solution to Exercise 2. We solve the game using backward induction, starting with Bob's decision. Bob chooses the amount b to maximize his utility function:

$$U_B(a, b) = m_B(a, b) + \alpha m_A(a, b).$$

Substituting the monetary payoff functions, with Bob receiving three times what Ann gives, we get:

$$U_B(a, b) = (3a - b) + \alpha(10 - a + b) = 10\alpha + a(3 - \alpha) + b(\alpha - 1).$$

Since Bob chooses b , his optimal action depends entirely on the coefficient of b , which is $(\alpha - 1)$:

- If $\alpha < 1$, the coefficient is negative. Bob maximizes his utility by returning the minimum possible amount: $b^* = 0$.

- If $\alpha > 1$, the coefficient is positive. Bob maximizes his utility by returning the maximum possible amount: $b^* = 3a$.

Moving backward to Ann's decision, she anticipates Bob's response to maximize her own selfish monetary payoff, $m_A(a, b) = 10 - a + b$:

- **Case 1** ($\alpha < 1$): Ann knows $b^* = 0$. Her payoff is $10 - a$. To maximize this, she chooses $a^* = 0$.
- **Case 2** ($\alpha > 1$): Ann knows $b^* = 3a$. Her payoff is $10 - a + 3a = 10 + 2a$. To maximize this, she chooses the maximum possible amount: $a^* = 10$.

The outcome depends on whether α is greater or less than 1. If $\alpha > 1$, Bob is sufficiently altruistic, and full trust is achieved ($a^* = 10$, $b^* = 30$).

Exercise 3. Find the backward-induction solution under the assumption that Bob is *inequity averse* à la Fehr and Schmidt (1999). Proceed in steps:

1. First, assume β is equal to 0.
2. Second, assume α is equal to 0. In which of the two cases does Ann give something to Bob?

Solution to Exercise 3. Let us first determine the difference in their monetary payoffs, which is central to Bob's utility:

$$m_A(a, b) - m_B(a, b) = (3a - b) - (10 - a + b) = 10 - 4a + 2b.$$

$$m_B(a, b) - m_A(a, b) = (3a - b) - (10 - a + b) = 4a - 2b - 10.$$

1. **Assume** $\beta = 0$, Bob only feels envy, $\alpha > 0$:

$$U_B(a, b) = m_B(a, b) - \alpha \max(m_A(a, b) - m_B(a, b), 0).$$

$$U_B(a, b) = 3a - b - \alpha \max(10 - 4a + 2b).$$

Any increase in b directly reduces Bob's monetary payoff and increases Ann's. If Ann has more money than Bob, giving her money increases Bob's envy. If Bob has more money, giving her money reduces his wealth. Therefore, increasing b strictly decreases Bob's utility. Bob's best response is $b^* = 0$. Anticipating this, Ann maximizes $10 - a$ by choosing $a^* = 0$.

2. **Assume** $\alpha = 0$, Bob only feels guilt, $\beta > 0$:

$$U_B(a, b) = m_B(a, b) - \beta \max(m_B(a, b) - m_A(a, b), 0).$$

If Bob is ahead, i.e. $m_B > m_A$, his utility is:

$$U_B(a, b) = (3a - b) - \beta(4a - 2b - 10).$$

Taking the derivative of U_B with respect to b , we get $-1 + 2\beta$. This yields two sub-cases:

- If $\beta < 0.5$, the derivative is negative. Bob's greed outweighs his guilt, so he chooses $b^* = 0$. Anticipating this, Ann chooses $a^* = 0$.
- If $\beta > 0.5$, the derivative is positive. Bob's guilt outweighs his greed. He will increase b until the payoffs are equal, $m_A = m_B$, which occurs when $4a - 2b - 10 = 0$, or $b^* = 2a - 5$.

If $\beta > 0.5$, Ann knows Bob will equalize the payoffs. Equal payoffs mean they evenly split the total wealth generated in the game, which is $(10 - a) + 3a = 10 + 2a$. Ann's guaranteed payoff is half of this:

$$m_A = \frac{10 + 2a}{2} = 5 + a.$$

To maximize $5 + a$, Ann will give the maximum amount: $a^* = 10$.

Ann only gives something to Bob in the second case, when $\alpha = 0$ and $\beta > 0.5$. Because the multiplier is 3, Ann's final payoff when giving $a = 10$ is 15 euros, giving her a strict incentive to trust Bob when he is sufficiently inequity-averse.

This analysis illustrates how trust can emerge when individuals care about fairness or others' welfare, helping explain why trust and reciprocity are frequently observed in experiments.

References

- Becker, G. S. (1974). A theory of social interactions. *Journal of Political Economy*, 82(6), 1063–1093.
- Fehr, E., & Schmidt, K. M. (1999). A theory of fairness, competition, and cooperation. *The Quarterly Journal of Economics*, 114(3), 817–868.

Enrico Mattia Salonia

Game Theory - Spring 2026

Public Good Game

A group of $n \geq 2$ students is preparing a shared project for a game theory class they have together. Students are indexed by $i \in N = \{1, 2, \dots, n\}$. Each student has an effort endowment of $e > 0$ (think of hours of effort or energies that could be dedicated to other tasks) and simultaneously chooses a contribution $c_i \in [0, e]$ to the common project.

Total group effort is

$$C = \sum_{i=1}^n c_i.$$

Student i 's payoff is

$$\pi_i(c_i, c_{-i}) = e - c_i + rC, \quad \frac{1}{n} < r < 1.$$

Here, r is the marginal per-capita return: one extra unit contributed to the class project generates a benefit of r for each student.

1. Fix contributions of all students other than i , denoted by c_{-i} . Derive student i 's best response.
2. Find the set of Nash equilibria of the game.
3. Find the contribution profile that maximizes total class welfare $\sum_{i=1}^n \pi_i$.
4. Compare the Nash-equilibrium outcome and the socially optimal outcome, and briefly interpret the gap in this classroom context.
5. Suppose student i has *preferences for universalisation*: rather than taking others' contributions as fixed, she asks "what if everyone does the same as I do?" Formally, student i chooses c_i to maximise

$$\tilde{\pi}_i(c_i) = e - c_i + r \cdot nc_i = e + (rn - 1)c_i,$$

imagining that every other student $j \neq i$ also contributes $c_j = c_i$. What does student i choose? How does this compare to the Nash equilibrium and the social optimum?

Solution For student i , write

$$\pi_i = e - c_i + r \left(c_i + \sum_{j \neq i} c_j \right) = e + (r - 1)c_i + r \sum_{j \neq i} c_j.$$

Since $r < 1$, we have $r - 1 < 0$, so π_i is strictly decreasing in c_i . Hence the unique best response is

$$c_i^*(c_{-i}) = 0.$$

Therefore, since $c_i^* = 0$ is a *dominant strategy* for each student i —it is optimal regardless of what others contribute—the unique Nash equilibrium is

$$c_1^* = \dots = c_n^* = 0.$$

Total welfare is

$$W(c) = \sum_{i=1}^n \pi_i = ne - \sum_{i=1}^n c_i + nr \sum_{i=1}^n c_i = ne + (nr - 1) \sum_{i=1}^n c_i.$$

Because $r > 1/n$, we have $nr - 1 > 0$, so welfare is increasing in total contributions. Thus the social optimum is

$$c_i^{\text{SO}} = e \quad \text{for all } i.$$

Interpretation: this is a free-riding problem. Individually, each student prefers to keep their own endowment (private marginal return is $r < 1$); collectively, full contribution is best for the class (social marginal return is $nr > 1$).

If student i universalises her choice, she perceives her payoff as

$$\tilde{\pi}_i(c_i) = e - c_i + r \cdot nc_i = e + (rn - 1)c_i.$$

Since $r > 1/n$, we have $rn - 1 > 0$, so $\tilde{\pi}_i$ is strictly increasing in c_i . The universalising student therefore chooses

$$c_i^{\text{U}} = e.$$

This coincides with the socially optimal contribution. Intuitively, the universalisation preference internalises the positive externality of contributing: by imagining that her own choice is adopted by the whole group, student i effectively perceives the *social* marginal return $nr > 1$ rather than the private one $r < 1$. The free-rider problem dissolves not because incentives change, but because the student's reasoning criterion does.

Preferences for universalisation have been introduced by Laffont (1975) and further studied by Alger and Weibull (2013) and Roemer (2019).

References

- Alger, I., & Weibull, J. W. (2013). Homo moralis—preference evolution under incomplete information and assortative matching. *Econometrica*, 81(6), 2269–2302.
- Laffont, J.-J. (1975). Macroeconomic constraints, economic efficiency and ethics: An introduction to kantian economics. *Economica*, 42(168), 430–437.
- Roemer, J. E. (2019). *How we cooperate*. Yale University Press.

Enrico Mattia Salonia

Game Theory - Spring 2026

Guilt Aversion

A traveler (the *tipper*, Player 2) takes a taxi from the airport to the city center. The driver (Player 1) provides the service first, and then the tipper decides how much to tip. Let the tip be $t \in [0, 1]$, measured as the percentage of the ride price. The tipper prefers to keep more money, so tipping is costly. However, the tipper is **guilt-averse**: she dislikes disappointing what she believes the driver expects to receive. Let $\tau \in [0, 1]$ be the tip that the tipper **believes** the driver expects. Her utility is

$$U_2(t \mid \tau) = (1 - t) - \beta [\tau - t]^+,$$

where $[x]^+ = \max\{x, 0\}$ and $\beta > 0$ measures guilt sensitivity.

1. For fixed belief τ , find the tipper's best response $t^*(\tau)$.
2. Show that if $\beta \leq 1$, the unique best response is $t^*(\tau) = 0$ for every τ .
3. Show that if $\beta > 1$, the best response is

$$t^*(\tau) = \tau.$$

4. Briefly interpret the result: when does guilt sustain positive tipping?
5. Now require *belief consistency*: in equilibrium, the driver's expectation must be correct, i.e. $\tau = t^*(\tau)$. Find all equilibria as a function of β .
6. Use your answer to part (e) to explain why tipping norms can differ sharply across countries (e.g. 20% in the US, 0% in Japan) even if individuals have identical preferences.

Solution. Parts (a)–(c). For $t \geq \tau$, we have $[\tau - t]^+ = 0$, so

$$U_2(t \mid \tau) = 1 - t,$$

which is decreasing in t . Hence in this region the best choice is $t = \tau$.

For $t < \tau$, we have $[\tau - t]^+ = \tau - t$, so

$$U_2(t \mid \tau) = 1 - t - \beta(\tau - t) = 1 - \beta\tau + (\beta - 1)t.$$

Therefore:

- If $\beta \leq 1$, this is weakly decreasing in t , so the best point in $[0, \tau)$ is $t = 0$. Comparing with $t = \tau$:

$$U_2(0 \mid \tau) = 1 - \beta\tau \geq 1 - \tau = U_2(\tau \mid \tau),$$

so $t^*(\tau) = 0$.

- If $\beta > 1$, this is increasing in t , so the best point in $[0, \tau)$ is the upper boundary, i.e. approaching τ . In the region $t \geq \tau$, utility decreases from $t = \tau$. Hence $t = \tau$ is optimal, so $t^*(\tau) = \tau$.

Part (d). When guilt is weak ($\beta \leq 1$), monetary incentives dominate and the tipper leaves no tip regardless of what she thinks the driver expects. When guilt is strong ($\beta > 1$), the tipper matches what she believes the driver expects, so beliefs can sustain positive tipping.

Part (e). An equilibrium requires belief consistency: the driver's expectation equals the tip he actually receives, i.e. $\tau = t^*(\tau)$.

Case $\beta \leq 1$. From part (b), $t^*(\tau) = 0$ for every τ . Consistency requires $\tau = 0$. This is the unique equilibrium: the driver expects nothing and receives nothing.

Case $\beta > 1$. From part (c), $t^*(\tau) = \tau$. The consistency condition $\tau = t^*(\tau) = \tau$ is satisfied for every $\tau \in [0, 1]$. Therefore there is a *continuum* of equilibria: for any $\tau^* \in [0, 1]$, the profile $(\tau, t) = (\tau^*, \tau^*)$ is an equilibrium.

In particular, both $t = 0$ and $t = 1$ (and everything in between) can be sustained in equilibrium. The equilibrium tip level is indeterminate from preferences alone: it depends entirely on which self-fulfilling belief is selected.

Part (f). When $\beta > 1$, the model generates a continuum of equilibria indexed by τ^* . Each equilibrium is sustained by self-fulfilling beliefs: the driver expects τ^* , and the guilt-averse tipper matches that expectation, confirming the driver's belief.

This means that two societies with identical preferences (same $\beta > 1$) can settle on very different tipping norms. A 20% norm in the US and a 0% norm in Japan are both equilibria of the same game. What differs is not guilt sensitivity but the *belief* that coordinates behavior. Tipping norms are social conventions: arbitrary in origin but self-reinforcing once established, because deviating triggers guilt. The exercise thus explains cross-cultural variation without requiring any difference in underlying preferences.

A game where players' beliefs affect their utility over outcomes, like the one in this example, is called a psychological game. Psychological games were introduced by Geanakoplos, Pearce, and Stacchetti (1989) and further studied by Battigalli and Dufwenberg (2009).

References

- Battigalli, P., & Dufwenberg, M. (2009). Dynamic psychological games. *Journal of Economic Theory*, 144(1), 1–35.
- Geanakoplos, J., Pearce, D., & Stacchetti, E. (1989). Psychological games and sequential rationality. *Games and economic Behavior*, 1(1), 60–79.