

Rauf Vuqar ALIYEV

ANALYSING THE NOTION OF FAIRNESS THROUGH THE SCOPE OF THE ULTIMATUM GAME WITH REAL WORLD APPLICATIONS

Abstract

Traditional economic theory posits humans as purely self-interested, while other perspectives argue humans have an intrinsic preference for fairness. This paper puts forward the idea that we use the shared notion of ‘fairness’ as a coordination tool to help us maximise our long-term payoffs. We argue individuals cultivate reputations of fairness not out of altruism, but as a heuristic to maximize long-run payoffs. We use the ultimatum game to construct a model that formalizes the perceived reputational cost of accepting an unfair offer as a function of the type of player, the offer itself, and stage of the game. This reputational cost function decays exponentially as offers near an equitable point, explaining empirical findings. We then show how our model can be used to explain behaviours exhibited in real-life situations such as: salary negotiations, tipping, and in general situations where we exhibit inequity aversion.

Keywords: Fairness, Ultimatum Game, Reputational Cost, Coordination Tool, Long-term Pay-offs, Equitable Outcomes, Bounded Rationality, Behavioral Economics

UOT: 339

JEL: F1

DOI: <https://doi.org/10.54414/SDGO8130>

Introduction

‘Fairness’ is a concept that has been the subject of multidisciplinary study garnering attention from philosophers, ethicists, economists, psychologists and even mathematicians. In the realm of economics, fairness is often understood in the framework of a distribution of resources between parties (Rawls, 1971). Traditional economics is largely based around the concept of the Homo Economicus, a being of pure rationality and self-interest who will reject any action that is not utility maximizing. This concept is the foundation of the most important of classical and neoclassical works such as the Rational Choice Theory (Von Neuman & Morgenstern, 1944) and many of the ideas presented in *The Wealth of Nations* (Smith, 1776) from which comes the famous quote:

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

Critics of the traditional view argue that it oversimplifies human motivations and that there are intrinsic values, such as fairness and reciprocity, that guide human behaviour independently of self-interest (Bowles & Gintis, 2012). Fehr, E., & Schmidt, K. M. (1999) argue that a fraction of human’s value ‘fairness’ while another fraction act in the way of the Homo Economicus (as traditional economic theory posits). However, a key point to note is that they posit that instead of caring about objective ‘fairness’ per se, what this fraction really exhibit is self-centered inequity aversion, meaning that when they are faced with inequitable outcomes towards them, they will even sacrifice in order to veer towards a fairer to them outcome.

To illustrate this consider the following public good game from the Fehr, E., &

Schmidt, K. M. (1999) study: a group of players ($N \geq 2$) all get an endowment of Y , they then each have to choose an amount g ($0 \leq g \leq Y$) to donate to some public good, at the end each player gets a pay-off (X_i) of $Y - g_i + a(\sum g_i)$ where $1/n < a < 1$. This means that by contributing to the public good, the individual utility of each player is less than it would be if they set $g_i = 0$, as $a < 1$, however, the total utility in the system (or the sum utility for all the players) is maximized if $g_i = Y$, as $a > 1/n$. Now, consider an extension of this game, where after the initial allocation by the players, there comes a stage 2, where each player can choose to punish other players. Punishment comes in the form of a reduced pay-off for the punished player. Punishing a player comes at a cost of C per unit deducted from the punishing player. Then at this stage:

$$\text{Individual Payoff}(X_i) = Y_i - g_i + a(\sum g_i) - c(\sum P_{ij}) - (\sum P_{ji})$$

Where P_{ij} is the punishment conferred by the player on the other players, and P_{ji} is the punishment conferred by other players on the player. According to rational economic theory, and the model of the Homo Economicus, stage 2 should be irrelevant, as imposing punishment comes at a cost to the punishing player. Therefore, it follows that in the case of the Homo Economicus both P_{ij} and P_{ji} should be 0. However, empirical studies show that, cooperators will punish free-riders (for under contributing to the public good) even at the cost of their own pay-off, achieving a more equitable outcome (by reducing both their and other pay-offs to reach an equal level), however, in variations of the game where punishment is impossible (where only stage 1 exists without stage 2), almost all agents, will veer towards the free-riding model in-line with traditional economic theory. In real world examples, the Fehr, E., & Schmidt, K. M. (1999) study talks about the success of corporation, and equity-seeking in bilateral agreements as in bilateral agreements it is more likely that an unfairly acting agent can be punished, however, in the market setting the self-interest model success-

fully predicts behaviours for most agents, suggesting that in a situation where blame can be attributed and punishment can be carried out, agents will veer towards the 'fairness' model, but if actions are anonymous and blame cannot be easily attributed, agents will act in a purely selfish manner. They use this evidence to suggest that the actions of this selfish faction, in situations where punishment is impossible, corrupt the selfless into acting in purely selfish ways. In this paper, through the use of the repeated ultimatum game, we will argue that it is not out of the drive for equity, even if self-centred, that causes this behaviour, but instead that this behaviour is part of a larger super-game strategy, that we as given our bounded rationality, have adopted to achieve favourable outcomes as the consequences of our own behaviours.

The Ultimatum Game:

The ultimatum game is one of the most frequently used tools in academic literature analysing the concept of 'fairness'. Originally introduced by John Harsanyi (Harsanyi, 1961), the ultimatum game consists of 2 players, the proposer and the responder. The proposer, is given an endowment of Y , the proposer must then split the endowment between themselves and the responder, so that their respective pay-offs are $(Y-r, r)$, with r being the amount the proposer is offering the responder. In the next stage, after the offer has been made, the responder has a choice between accepting the offer or rejecting the offer, in the case that they accept the offer the proposer gets $Y-r$, and the responder gets r , as offered. In the case that the responder rejects however, both parties get 0.

In the classic version of the game, there is 1 round, where players are matched randomly each round against a different player. According to classical game-theory and the rational choice model the Nash equilibrium for this game is for the proposer to offer $(Y-r_{\min}, r_{\min})$ with $r_{\min}$ being the smallest denomination of the whole endowment that the proposer can offer, for example if the minimum denomination is 1 cent, then the Nash equilibrium should be $(Y-0.01\$, 0.01\$)$. The best possible option then for the responder would be to accept, as by rejecting they simply get a pay-off of 0, and $r_{\min} > 0$.

However, in empirical studies individuals are likely to reject any offer below 40% of the total endowment (Güth and Tietz, 1990; Bergstrom, 1996). In traditional behavioural economics this effect is explained by inequity aversion, as discussed in our earlier example and as outlined Fehr, E., & Schmidt, K. M. (1999), additionally past research has shown that when individuals work for the right to be a proposer (i.e. when the roles are not randomly assigned but there has to be some work done to be assigned the proposal role), they are more likely to offer unfair offers, and the responders are much more likely to accept those unfair offers (Hoffman et al., 1996), as both parties perceive that it is fair for the proposer to get more of the endowment as they had to bear some cost, not in pay-off but in utility, as the work that had to be done to 'earn' the role of the proposer caused some disutility. This is interesting because it shows that the concept of fairness isn't predicated on just simple equity and that utility regulation plays into our considerations for what is 'fair' and what is unfair. Furthermore studies have shown that cultural aspects play a great role in bargaining behavior as different cultures have different societal norms which in turn influence what each individual considers as 'fair' (Simon Gächter & Jonathan F. Schulz, 2016), which further reinforces our idea that 'fairness' is more of a dynamic concept than it is a static one, and it is predicated on more than equality, moreover agents adapt their decision making based on their expectations (Sanfey, 2009; Chang and Sanfey, 2011; Xiang et al., 2013), for example, if agents are told beforehand that they will be receiving low offers, then they are more likely to accept lower offers, this contradicts inequity aversion because being told you are going to get less shouldn't impact the disutility you experience if you truly are inequity averse, and as such shouldn't impact your behaviour.

The Model.

To demonstrate our model, we make use of a derivative of the standard ultimatum game, the repeated ultimatum game. In the repeated ultimatum game, instead of having 1 round against random opponents, the ultimatum game is played repeatedly against the same opponent for

N number of rounds. Theoretically, the sub-game Nash equilibria do not change in the repeated ultimatum game, however one key difference is, that in the repeated ultimatum game, there exists an opportunity for a super-game strategy (Slembeck, T., 1999). A super-game strategy can be defined as a strategy that is not confined to the current stage or round of a game but instead considers the entire sequence of play across multiple stages. It's a way to strategize in the context of repeated interactions with the same opponent, rather than treating each round of the game in isolation. As decisions made in individual sub-games can affect future sub-game behaviour and as such this must be considered when thinking about Nash equilibria. In our original Nash equilibria in the standard ultimatum game, we made the implicit assumption that accepting an offer does not incur a cost on the responder, and as such they should accept any non-null offer, now however, in a super game setting, if accepting low offers might signal to the proposer that you are willing to accept low offers and as such would introduce another a 'reputational' cost to accepting offers, as if you accept a low offer, it might anchor future offers from the proposer to be low. As such the new pay-offs in each

sub-game for the proposer, and the responder respectively becomes $(Y - r_{min} - \delta, r_{min} + \delta)$, where δ is the reputational effect of accepting an offer, if the reputation effect of accepting an offer is negative ($\delta < 0$) this means that, the proposer gains a reputational benefit from the responder accepting such offer and thus with a negative δ the proposer's pay-off increases, whereas the responder's pay-off decreases. To illustrate this, imagine the proposer gave the most unequitable offer possible the Nash equilibrium offer, this signals to the proposer that the responder is more likely to accept unequitable offers, and as such gives the proposer/responder an additional reputational benefit/cost in their total payoff, as they can then adjust their super game strategy to reflect this. This view is supported by the study done by (Slembeck, T., 1999) which found the presence of 'fair' and 'tough' players. Although in our model we are going to look at the responder

side of things, I want to note that although in this paper we are only going to look at the reputational effect from the side of the responder, i.e the reputational effect of *accepting* an offer, there is also a reputational effect of *proposing* an offer, that we are not going to analyse in this paper, to demonstrate this idea, imagine if a repeated ultimatum game, a respondent employs the strategy of accepting offers that amount to 80% of the total endowment, and rejecting any other offer, then let's say the proposer starts with an equitable 50/50 offer, if the proposer is then quick to be coerced into the demands of the responder, the very act of *offering* a 20/80 offer might put the proposer at a reputational disadvantage. Although the exact mechanisms of this effect are left to be studied in future studies.

In our paper we are going to try and formalize this δ variable and introduce a model which explains the role of fairness in the repeated ultimatum game, and where the fundamental model can be applied to other games as-well, we will show that the behaviour that our model predicts in-line with previous literature.

The model introduced in this paper, proposes that we as humans do not have a preference for equity, nor does it that some fraction of people have a preference for equity and others don't, instead we argue that in situations which presuppose competition, or in any situation where our and others behaviour can impact our pay-off, such as in market-based games or in bilateral bargaining situations, we use our shared notions of fairness as a coordination device to help us achieve cooperative behaviour, especially when non-cooperative behaviour might be costly.

We can formalize the payoffs for the respondent as such:

$$\text{Payoff}(r \mid \text{Accept}) = r + \delta$$

Whereas if they reject

$$\text{Payoff}(r \mid \text{Reject}) = 0.$$

With r being the amount offered to the responder, and δ being the reputational cost of accepting said offer, therefore the respondent should only accept the offer if $r > \delta$. Therefore to understand better whether a respondent will accept or reject the offer we need to delve

deeper into the determinants of δ , and what δ depends upon, we propose that δ depends largely on 3 factors, that being how far along in the game you are, how many more rounds there are in the future, who you are playing, the unique characteristics of who you are going against, and lastly the offer itself, not only the amount being offered, but also how that amount relates to both the full endowment and how that amount relates to real world sums, as past studies have found that when the stakes are increased rejection rates fall dramatically (Steffen Andersen et al. 2011).

Therefore, we can formalise our model as such:

$$\text{Payoff}(r) = r + \delta(\text{Player}, r/r + X, \text{round})$$

The Player

To further continue deconstructing the δ term we are going to focus on the first vector $\delta(\text{Player})$. Who we are playing against is an important factor in the reputational cost of accepting an offer as for example, different people have different perceptions of bargaining and different societal norms on what they consider fair. For the purposes of this model, we are going to assume that there exists two types of proposers, 'elastic' and 'inelastic' proposers. 'elastic' proposers are ones that are quick to react to signals (rejections) by responders and are likely to adjust their offers to achieve equitable outcomes, inelastic proposers on the other hand are 'tough' proposers who are likely to keep offering their initial offer to try and strong-arm the responder into accepting. In an ideal model we would be able to measure the exact elasticity of the proposer, however for the sake of simplicity in this model we are going to assume that the proposer is either elastic or inelastic, we are then going to assume that given an inequitable offer, accepting it is more reputationally costly when going against an elastic proposer than it is going against an inelastic proposer. This is because if you are to reject an inequitable offer from an elastic proposer, the chance that they then adjust future offers to be more

equitable is higher, therefore the reputational cost of accepting is higher, as such:

$$\text{Payoff}(r) = r + \Pr(\text{Elastic}) \cdot \delta\left(\frac{r}{r+X}, \text{round} \mid \text{Elastic}\right) + \Pr(\text{Inelastic}) \cdot \delta\left(\frac{r}{r+X}, \text{round} \mid \text{inelastic}\right)$$

$$\text{Payoff}(r) = r + \Pr(\text{Elastic}) \cdot \delta\left(\frac{r}{r+X}, \text{round} \mid \text{Elastic}\right) + \Pr(\text{Inelastic}) \cdot \delta\left(\frac{r}{r+X}, \text{round} \mid \text{inelastic}\right)$$

$$\text{Payoff}(r) = r + \Pr(\text{Elastic}) \left(\delta\left(\frac{r}{r+X}, \text{round} \mid \text{Elastic}\right) - \delta\left(\frac{r}{r+X}, \text{round} \mid \text{inelastic}\right) \right) + \delta\left(\frac{r}{r+X}, \text{round} \mid \text{inelastic}\right)$$

Therefore:

$$r > \Pr(\text{Elastic}) \left(\delta\left(\frac{r}{r+X}, \text{round} \mid \text{Elastic}\right) - \delta\left(\frac{r}{r+X}, \text{round} \mid \text{inelastic}\right) \right) + \delta\left(\frac{r}{r+X}, \text{round} \mid \text{inelastic}\right)$$

Now if we assume that the reputation cost of accepting an offer from an inelastic proposer is 0.

$$r > \Pr(\text{Elastic}) \cdot \delta\left(\frac{r}{r+X}, \text{round} \mid \text{Elastic}\right)$$

It must be acknowledged that this is a very stringent assumption, as in real life it might not be a case that the reputational cost of accepting offers for inelastic consumers is 0 but is simply less than the cost of accepting offers from elastic and future research should look at relaxing this assumption to get more nuanced results. However, to be able to derive theoretical implications about behaviour from this model without over-complicating the model.

The Offer.

The offer value itself plays a large role in the reputational cost attached to accepting, and it is perhaps the most salient factor. It is important to note that unless

$$r \geq Y - r_{\min}$$

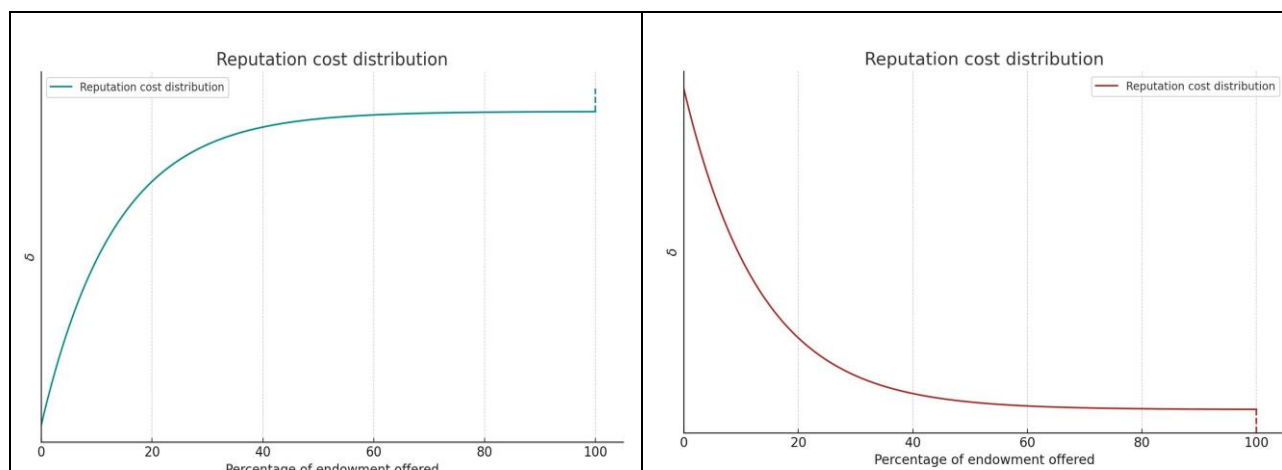
$\Delta+$ will always be equal to or less than 0, this is because accepting any offer of R_i signals to the proposer a definitive boundary. Specifically, it communicates that, all other things being equal, there is no need for the proposer to exceed an offer higher than R_i , as the recipient's willingness to accept R_i is evidenced by rational preferences. i.e if the responder has already accepted an offer of R_i offering any other amount is illogical as they have signalled that they are willing to accept R_i .

Therefore the reputational cost of accepting any offer that is not either the full endowment, or $Y - r_{\min}$ is below 0. The cost drops to 0 at $Y - r_{\min}$ as this is the greatest amount any rational proposer would propose as when proposing anything over this amount the proposers pay-off is 0 thus they have no incentive to offer the full endowment. This signalling effect thus defines a strategic constraint within the model, impacting the proposer's behaviour and expectations.

The crux of this model lies in the following analysis, the distribution of $\delta()$ is highly dependent on the offer itself, and more importantly it is determined by $r/r+X$. We posit that the absolute value of the reputational cost follows a pattern of exponential decay.

Reputational cost distribution:

Absolute reputational cost distribution (for illustration purposes): As we can see in the graphs, we posit that the reputational effect decreases (increases technically) exponentially, plateauing at around $x=50$ (this is an assumption that we have made for the sake of the illustration). The idea is that the closer any offer comes to being recognized as 'fair' the less of a reputational cost accepting that offer incurs, and therefore the less of a reputational gain can be made from rejecting said offer. The idea is, if you reject offers that are unfair, the proposer is more likely to adjust their offer to a fairer one, but if you are to reject an offer that is already recognized as 'equitable' the proposer is unlikely to shift their stance and offer you a more favourable offer.



Furthermore, the higher an offer is from what is considered 'equitable' the bigger the reputational cost of accepting said offer, for example, if I am playing the ultimatum game against a proposer and I am offered a split of (99,1) if I reject this offer I am more likely, in the following rounds to see the proposer offer a split of (99-Adjustment, 1+Adjustment), and the closer this offer is to being 'fair', the lower the probability that an adjustment will happen as well as the smaller the adjustment will be. So if instead let's say I reject an offer that is (60,40) I am less likely to see an adjustment, and even if I do it won't be to the degree that I would see if I rejected a (99,1) offer. Furthermore, any marginal reputational gain to be had after the 'fair' point is negligible, therefore rejection rates will fall dramatically after what is considered an 'equitable' point as there is no **more** reputational benefit to be had rejecting offers after that amount. Furthermore, this analysis then suggests a 'fairness' isn't an endogenous moral standard that all, or 'some' humans have, but is instead a game-theoretic-esque coordination tool that may differ in definition over different human societies, that we use to signal our preferences and achieve equitable outcomes. This is supported by the fact that in situations where coordination is impossible, such as in market-based situations, fairness considerations are dismissed, supported by the study by Fehr, E., & Schmidt, K. M. (1999), coming back to our earlier conclusion where we said that 'suggesting that in a situation where blame can be attributed and punishment can be carried out,

agents will veer towards the 'fairness' model, but if actions are anonymous and blame cannot be easily attributed, agents will act in a purely selfish manner.' Another way to look at this is that instead of it being the fact that blame cannot be attributed, it's the fact that situations where blame cannot be attributed coincide with situations where coordination is impossible, and if coordination is impossible, subsequently coordination devices, such as fairness, are thrown out the window and not considered in our behaviour.

Finally looking at the final variable, rounds, in our model works in a rather straightforward way, the more rounds left to play, the incentive to reject offers must be higher, coming down in a linear fashion. To illustrate, if we are playing 100 rounds, and you offer an unequitable offer (99,1) in round 1, I have a greater incentive to reject this offer as compared to if you offer an inequitable offer in round 99, and I have no incentive to reject the offer on round 100. This is because if I can achieve a shift in the proposer's behaviour in round 1, I get to enjoy this reputational benefit for the next 99 rounds, whereas if I gain a reputational benefit in round 99, I only get to consume for 1 round, therefore the reputational benefit I can gain from rejecting an offer in round 1, should be greater than the reputational benefit I can get from rejecting an offer in round 98, and this effect should decrease linearly over the number of rounds played, δ should theoretically reach 0 in the final round of play. However, this often-times in empirical studies is not the case.

To summarize, our model predicts that agents exhibit uncooperative behaviour in a bit to coerce others in game-situations into behaving in a more favourable, and cooperative manner. Although this interpretation is elegant and fits nicely with previous literature such as Fehr, E., & Schmidt, K. M. (1999) findings, we want to address a seeming short-coming of the paper, and that is that if rejection behaviour is a solely selfish and utility maximizing behaviour exhibited by agents to maximise super-game pay-offs, why is it that studies have found that even in one-shot interactions, where there is no potential for any future games, do users *still* exhibit such behaviour. To answer this question, we are going to rely on the interpretation given by (Camerer & Thaler (1995); Gale Et Al. (1995); Slembeck, T. (1999)) that agents carry over their “repeated-game impulses” (Hoffman et al. 1994) that the learnt through-out their life, into an experimental setting, that is to say, we as humans aren’t exactly calculating our reputational benefits and costs at every-turn, but instead due to our bounded rationality live by principles we have learnt throughout our lives to achieve favorable outcomes.

Another way to look at it is this, imagine we as humans look at the duration of our lives as one large ‘super-game’, and in every interaction we follow the model we have proposed to ensure equitable outcomes *on average*, then it follows that even in one-shot games, we are still going to act according to our larger principles. The largest piece of evidence for this interpretation is that past literature has consistently shown that as people get older their approaches to bargaining become more relaxed.

Research by Lim and Yu (2015) used experimental methods to investigate bargaining behaviour across different age groups. They found that older adults were more likely to accept unfair offers in a bargaining game, demonstrating a less stringent and more accommodating attitude towards bargaining. Furthermore, Yeung and Fung (2007) in their research on lifespan changes in social decision-making suggest that age-related changes in social motivation may

lead to an emphasis on maintaining positive social interactions, potentially overriding the strict competitive behaviour often seen in bargaining scenarios (Yeung, D. Y., & Fung, H. H., 2007). The idea that as people get older, they put a higher emphasis on maintaining positive social interactions and employ a more relaxed attitude towards bargaining is endogenous with the findings of our model. To illustrate think about it this way, think about a person’s life as a large set of ‘games’ or negotiations, with each negotiation being a sub-game, with this interpretation, it follows that as we have less sub-games to go, we are less worried about incurring any reputational disadvantage, of-course you can make the argument that many of the sub-games within our life aren’t connected to each other and therefore curating a reputation consistently throughout all of them is not logical, this is a valid criticism of the model, however, we can make the argument that actually a lot of sub-games in life are interconnected, for example, decisions you make at work when negotiating with let’s say consultants, might affect the way that people perceive you (your reputation) which might in-turn affect your future job-prospects, in a way akin to the butterfly effect, therefore with our bounded rationality it is hard or near impossible to predict which ‘sub-game’ might have an effect on our future livelihood and pay-offs and which won’t as such people live their lives according to heuristics and principles.

Conclusion

In conclusion, this paper has presented a novel theoretical framework for understanding our seemingly inherent desire to fairness and equitable outcomes. We argue that instead of being motivated by an altruistic need for fairness, we instead use our shared concept of fairness as a coordination device to achieve equitable and favourable outcomes. We also work to curate a reputation of ‘fairness’ in order to ensure that ensure possible cooperators that we are unwilling to accept unequitable propositions and offers. In a way a-kin to the strategy of commitment in classical game-theory, by making a commitment that you are unwilling to

accept inequitable offers, using the shared notion of fairness, we maximize our payoff in the long term.

REFERENCES

1. Bergstrom, J.C. (1996). Kagel, John H., and Alvin E. Roth, eds. *The Handbook of Experimental Economics*. Princeton NJ: Princeton University Press, 1995, vi + 328 pp,.
2. \$@@-@@55.00. *American Journal of Agricultural Economics*, 78(3), pp.834–834. doi:<https://doi.org/10.2307/1243316>.
3. Chang, L.J. and Sanfey, A.G. (2011). Great expectations: neural computations underlying the use of social norms in decision-making. *Social Cognitive and Affective Neuroscience*, 8(3), pp.277–284. doi:<https://doi.org/10.1093/scan/nsr094>.
4. Dubreuil, B. (2012). *A Cooperative Species: Human Reciprocity and its Evolution*, S. Bowles and H. Gintis. Princeton University Press, 2011, xii + 262 pages. *Economics and Philosophy*, 28(3), pp.423–428. doi:<https://doi.org/10.1017/s0266267112000302>.
5. Fehr, E. and Schmidt, K.M. (1999). A Theory of Fairness, Competition, and Cooperation. *The Quarterly Journal of Economics*, 114(3), pp.817–868. doi:<https://doi.org/10.1162/003355399556151>.
6. Gächter, S. and Schulz, J.F. (2016). Intrinsic honesty and the prevalence of rule violations across societies. *Nature*, [online] 531(7595), pp.496–499. doi:<https://doi.org/10.1038/nature17160>.
7. Güth, W. and Tietz, R. (1990). Ultimatum bargaining behavior. *Journal of Economic Psychology*, 11(3), pp.417–449. doi:[https://doi.org/10.1016/0167-4870\(90\)90021-z](https://doi.org/10.1016/0167-4870(90)90021-z).
8. Harsanyi, J.C. (1961). On the rationality postulates underlying the theory of cooperative games. *Journal of Conflict Resolution*, 5(2), pp.179–196. doi:<https://doi.org/10.1177/002200276100500205>.
9. Hoffman, E., McCabe, K.A. and Smith, V.L. (1996). On expectations and the monetary stakes in ultimatum games. *International Journal of Game Theory*, 25(3), pp.289–301. doi:<https://doi.org/10.1007/bf02425259>.
10. K., M.G., von Neumann, J. and Morgenstern, O. (1944). *Theory of Games and Economic Behaviour*. *Journal of the Royal Statistical Society*, 107(3/4), p.293. doi:<https://doi.org/10.2307/2981222>.
11. Rawls, J. (1971). *A Theory of Justice*. Harvard University Press.
12. Sanfey, A.G. (2009). Expectations and social decision-making: biasing effects of prior knowledge on Ultimatum responses. *Mind & Society*, 8(1), pp.93–107. doi:<https://doi.org/10.1007/s11299-009-0053-6>.
13. Slembeck, T. (1999). Reputations and Fairness in Bargaining - Experimental Evidence from a Repeated Ultimatum Game With Fixed Opponents. [Review of Reputations and Fairness in Bargaining - Experimental Evidence from a Repeated Ultimatum Game With Fixed Opponents.].
14. Smith, A. (1776). *The Wealth of Nations*. London: W. Strahan and T. Cadell.
15. Xiang, T., Lohrenz, T. and Montague, P.R. (2013). Computational Substrates of Norms and Their Violations during Social Exchange. *Journal of Neuroscience*, 33(3), pp.1099–1108. doi:<https://doi.org/10.1523/jneurosci.1642-12.2013>.

Рауф Вугар АЛИЕВ

АНАЛИЗ ПОНЯТИЯ СПРАВЕДЛИВОСТИ В РАМКАХ СОВЕРШЕННОЙ ИГРЫ С ПРИМЕНЕНИЕМ В РЕАЛЬНОЙ МИРЕ

Резюме

Традиционная экономическая теория утверждает, что люди руководствуются исключительно личными интересами, в то время как другие точки зрения утверждают, что людям присуще предпочтение справедливости. В этой статье выдвигается идея о том, что мы используем общее понятие «справедливости» в качестве инструмента координации, который поможет нам максимизировать наши долгосрочные выгоды. Мы утверждаем, что люди культивируют репутацию честных людей не из альтруизма, а как эвристический подход, направленный на максимизацию долгосрочных выгод. Мы используем игру «Ультиматум» для построения модели, которая формализует предполагаемые репутационные издержки принятия несправедливого предложения как функцию типа игрока, самого предложения и стадии игры. Эта функция репутационных издержек убывает экспоненциально по мере того, как предложения приближаются к равноправной точке, что объясняет эмпирические результаты. Затем мы покажем, как нашу модель можно использовать для объяснения поведения, проявляющегося в реальных ситуациях, таких как переговоры о зарплате, чаевые и в общих ситуациях, когда мы проявляем неприятие неравенства.

Ключевые слова: *справедливость, игра «Ультиматум», репутационные издержки, инструмент координации, долгосрочные выплаты, справедливые результаты, ограниченная рациональность, поведенческая экономика.*

Rauf Vüqar ƏLİYEV

ƏDALƏT DÜŞÜNCƏSİNİN HƏQİQİ DÜNYA TƏTBİQLƏRİ İLƏ ULTIMATUM OYUNUN ƏTRAFINDA TƏHLİL EDİLMƏSİ

Xülasə

Ənənəvi iqtisadi nəzəriyyə insanları sırf şəxsi mənafeyini güdür, digər perspektivlər isə insanların ədalətliyyə daxili üstünlük verdiyini iddia edir. Bu sənəd bizim uzunmüddətli gəlirlərimizi maksimum dərəcədə artırmağa kömək etmək üçün koordinasiya vasitəsi kimi ortaqlar "ədalətlik" anlayışından istifadə etdiyimiz fikri irəli sürür. Biz iddia edirik ki, fərdlər ədalətlik reputasiyasını altruizmdən deyil, uzunmüddətli gəlirləri artırmaq üçün bir evristik olaraq inkişaf etdirirlər. Biz ultimatum oyunundan, oyunçunun növündən, təklifin özündən və oyunun mərhələsindən asılı olaraq ədalətsiz təklifi qəbul etmək üçün qəbul edilən reputasiya dəyərini rəsmiləşdirən bir model qurmaq üçün istifadə edirik. Bu reputasiya dəyəri funksiyası tapıntıları izah edən empirik ədalətli nöqtəyə yaxın təkliflər kimi eksponent olaraq azalır. Daha sonra biz modelimizin real həyat danışıqlarında nümayiş olunan davranışları izah etmək üçün necə istifadə oluna biləcəyini göstəririk: maaş danışıqları, bahalaşma və ədalətsizlikdən çəkiniyimiz ümumi vəziyyətlərdə.

Açar sözlər: *Ədalət, Ultimatum Oyunu, Reputasiya Xərcləri, Koordinasiya Aləti, Uzunmüddətli Ödəmələr, Ədalətli Nəticələr, Məhdud Rəsonallıq, Davranış İqtisadiyyatı*