

Results of the First Annual Human-Agent League of the Automated Negotiating Agents Competition

Mell, Johnathan; Gratch, Jonathan ; Baarslag, Tim; Aydogan, Reyhan; Jonker, Catholijn M.

DOI

[10.1145/3267851.3267907](https://doi.org/10.1145/3267851.3267907)

Publication date

2018

Document Version

Final published version

Published in

IVA '18 Proceedings of the 18th International Conference on Intelligent Virtual Agents

Citation (APA)

Mell, J., Gratch, J., Baarslag, T., Aydogan, R., & Jonker, C. M. (2018). Results of the First Annual Human-Agent League of the Automated Negotiating Agents Competition. In *IVA '18 Proceedings of the 18th International Conference on Intelligent Virtual Agents* (pp. 23-28). ACM.
<https://doi.org/10.1145/3267851.3267907>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' – Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.

Results of the First Annual Human-Agent League of the Automated Negotiating Agents Competition

Johnathan Mell
University of Southern California
Los Angeles, USA
mell@ict.usc.edu

Jonathan Gratch
USC Institute for Creative Technologies
Los Angeles, USA
gratch@ict.usc.edu

Tim Baarslag
Centrum Wiskunde & Informatica
Amsterdam, Netherlands
T.Baarslag@cw.nl

Reyhan Aydoğan
Özyeğin University
Istanbul, Turkey
reyhan.aydogan@ozyegin.edu.tr

Catholijn M. Jonker
Delft University of Technology
Delft, Netherlands
C.M.jonker@tudelft.nl

ABSTRACT

We present the results of the first annual Human-Agent League of ANAC. By introducing a new human-agent negotiating platform to the research community at large, we facilitated new advancements in human-aware agents. This has succeeded in pushing the envelope in agent design, and creating a corpus of useful human-agent interaction data. Our results indicate a variety of agents were submitted, and that their varying strategies had distinct outcomes on many measures of the negotiation. These agents approach the problems endemic to human negotiation, including user modeling, bidding strategy, rapport techniques, and strategic bargaining. Some agents employed advanced tactics in information gathering or emotional displays and gained more points than their opponents, while others were considered more “likeable” by their partners.

CCS CONCEPTS

• Human-centered computing~Empirical studies in HCI

KEYWORDS

Human-Agent Negotiation; IAGO Negotiation Platform;

ACM Reference format:

Johnathan Mell, Jonathan Gratch, Tim Baarslag, Reyhan Aydoğan, and Catholijn M. Jonker. 2018. Results of the First Annual Human-Agent League of the Automated Negotiating Agents Competition. In *International Conference on Intelligent Virtual Agents (IVA '18)*, November 5–8, 2018, Sydney, NSW, Australia. ACM, New York, NY, USA, 8 pages. <https://doi.org/10.1145/3267851.3267907>

1 INTRODUCTION & BACKGROUND

Automated negotiation has been presented in previous works as a key challenge problem for the advancement of virtual hu-

mans—agents that possess human-like characteristics. Negotiation has been described as an “indispensable skill for any social creature”, and the fruits of research into automated negotiating agents have myriad benefits [5]. Specifically, automated negotiating agents are capable of training humans to be better negotiators, and insights gleaned from the design of such agents can lead to the development of more general teaching agents [3,8]. Additionally, automated agents provide benefits to empirical studies, as they can provide a consistent confederate, unchanging and tireless over the course of a human-subjects study [5].

Beyond training humans to be better negotiators, negotiating agents can also serve as virtual assistants for a variety of applications. Google’s recent Duplex demo demonstrates this astutely, as an agent acts on behalf of a human to create an appointment, negotiating with a human partner to decide on an agreeable time [7]. Legal scholarship has long examined the ethics of acting on behalf of others [12,15], and psychology and computer science have both explored the mechanisms by which humans instruct their representatives, be they human or virtual [4,11].

Of course, such negotiating agents require many components to render them fully capable of acting as negotiating partners. The design of negotiating agents presents problems in strategy, opponent modelling, preference elicitation, rapport-building, natural language generation/understanding, non-verbal behavior generation, use of emotional affect, to name just a few. When creating agents, designers must address, at the bare minimum, how an agent will model its opponent (through their utterances and/or offer patterns), how the agent will make offers, and how it will respond to offers.

Designing automated negotiating agents has been an ongoing area of research. One such way in which this research has been driven is by researcher collaboration during the Automated Negotiating Agents Competition (ANAC) [6]. While ANAC has been a recurrent, successful competition for 8 years (2018 has marked the 9th annual ANAC), it has been focused primarily on agent-agent negotiation. Human-agent negotiation is fundamentally different than agent-agent negotiation, and the Human-Agent Track of ANAC was added to promote research into this promising area.

Agents that are developed with humans in mind need not emulate their behaviors (although that may be a goal in certain

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. For all other uses, contact the Owner/Author.

IVA '18, November 5–8, 2018, Sydney, NSW, Australia

© 2018 Copyright is held by the owner/author(s). Publication rights licensed to ACM.

ACM 978-1-4503-6013-5/18/11...\$15.00

<https://doi.org/10.1145/3267851.3267907>

circumstances), but do need to be able to respond human behaviors intelligently. First and foremost, this presents a problem for agents to accurately model their opponents. Since negotiation is not a fully-visible scenario, agents must make educated decisions about user preferences, user alternatives to agreement (referred to as “Best Alternative to Negotiated Agreement, or “BATNA”), and individual user personality types/strategies. Agents may use a variety of information sources to come to their conclusions, including the user history of offers, the natural language utterances they send, and even the emotions they express.

Beyond opponent modeling, agents must decide their own behaviors in a negotiation. Although starting with a tough offer and conceding slowly is a known negotiating tactic in both human and agent-agent negotiations, it does have its risks. Certain negotiating partners are more disagreeable than others, and may fight back, lowering the joint value possible if an agent is too aggressive early on. Agents may also decide how optimistic or pessimistic to be in the face of uncertainty. And they must decide how to utilize their BATNA strategically, and if they want to follow certain strategies that may be considered unethical (such as lying).

Finally, agents must also be able to adapt to individual differences in negotiation. While a single strategy may be helpful in many cases, the best agents (and negotiators of any provenance) must be able to read the situation and adjust accordingly. These problems may be addressed with a variety of solutions, many of which are on display here, in the results from ANAC 2017.

2 COMPETITION DESIGN

2.1 IAGO Negotiation Platform

The IAGO Negotiation platform was proposed and designed by Mell et al. and was selected to be used for the Human-Agent League of ANAC [10]. IAGO provides a front-facing GUI for the human-participants (see Figure 1). This GUI is web-based, and does not require plugins beyond JavaScript support. As such, it is supported on multiple browsers and systems, and can be utilized over the web without lengthy installation or training. In particular, this feature allows subjects to be recruited using online platforms, such as Amazon’s Mechanical Turk (MTurk).

Additionally, IAGO provides the features necessary for simulating the characteristics of human negotiation. These include an expanded set of channels for communication between both sides of negotiation, such as by sending text, expressing preferences, and transmitting emotions. This is in addition to the traditional methods supported by agent-agent negotiation platforms, such as exchanging offers. IAGO also allows offers to be sent that do not involve all the items in the negotiation (“partial offers”).

These features of IAGO mean that it provides a platform to address the basic features that intelligent negotiating agents require. It provides information that allows for robust user modeling, and allows multiple channels for communicating in different ways. IAGO provides information that agents require to reason about their own preferences, and allows them to pursue a number of more complex strategies that require specific features (such as partial offers).

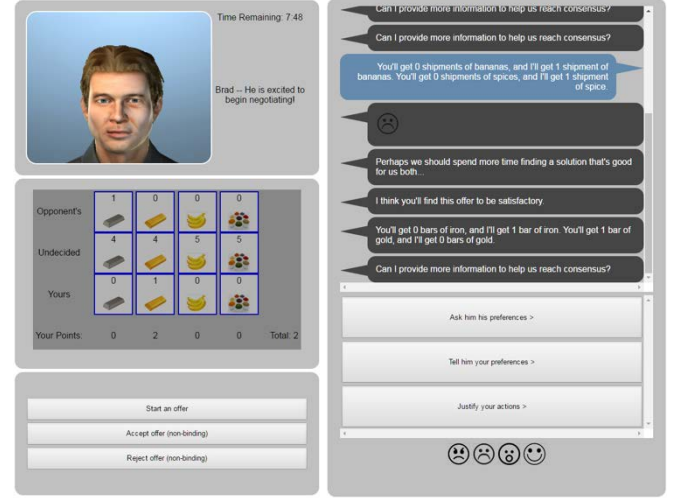


Figure 1: IAGO Research Platform (Client View)

2.2 Human-Agent Competition Design

2.2.1 General Information

The competition featured an array of participant-submitted agents competing against humans in a single, 10-minute, multi-issue negotiation. Participants were also asked a series of questions, ranging from demographic information to reviews of the agent behavior. The submitted agents were judged along two categories, and prizes were awarded to the best two agents in each category. The first category, agent points, was purely performance-based—agents that received more points in the negotiation were scored higher. The second category, agent likeability, was determined by user-submitted responses to Likert-scale questions after the end of the negotiation. See Section 4.1 for more details on this measure.

Agent designers were provided with a set of guidelines that restricted the domain of the negotiation within moderate bounds, but they were not given the details of the task itself, which was determined secretly prior to agent submission. The guidelines were summarized per the following equation:

$$\sum_{i=1}^k \text{Agent utility}(i) * (\text{num_levels}(i) - 1) = \sum_{i=1}^k \text{Human utility}(i) * (\text{num_levels}(i) - 1)$$

where k is the total number of issues. Succinctly, this means that the total for each side would be the same if that side got every item. In this way, the agent designers were given some idea of the scope of the negotiation.

Designers were also provided with a limited set of natural language utterances that the humans could use in the negotiation. Human players could also send messages that contained information about their preferences, in addition to using the emotional and offer channels. Agents were unrestricted in the types of messages they could send back to players. Finally, agent developers were provided with the source code for a baseline agent (“Pinocchio”) which was provided with the IAGO platform.

The competition’s negotiation had 4 issues, with a varying amount of levels to each. Respectively, each issue had 3, 2, 6, and 3 items. The task was partially integrative, with both sides gaining the most points by receiving the 6-item issue, but differing on their preferences for the two 3-item issues. Both sides also included a BATNA, which gave both players a minimum number of points should they fail to reach agreement.

2.2.2 Participant Information

Competition subject participants were selected from the MTurk subject pool. Subjects were adults, and asserted that they were permanent residents of the US (verified with IP address). Restriction to the US was chosen in order to reduce cross-cultural variance. Each submitted competition agent was tested against 25 participants. Participants were not re-used or matched against more than one agent. Subjects were asked a set of verification questions/attention checks to ensure they comprehended and were engaged in the negotiation.

All participants were presented with a tutorial of the system before use. Participants were paid regardless of their success in the negotiation. However, they were also awarded “lottery tickets” based on their performance. These lottery tickets then entered them into a prize drawing for one of several \$10 MTurk credits, incentivizing good performance during the negotiation. Subjects were removed from the pool if the agent against which they were matched encountered an unrecoverable crash. In this case, the subjects were not rerun. 10 total subjects were removed in this fashion. This design allowed the competition to follow best practices for subject recruitment and handling, in line with other research [1].

3 AGENT DESIGN

Table 1: Agent Provenance

Agent Name	Institution	Primary Author
Pinocchio	University of Southern California, USA	Johnathan Mell
Elphaba	College of Mgmt Academic Studies, Israel	Dor Nisim
Cena	Southwest University, China	Siqi Chen
Wotan	Southwest University, China	Lichun Yuan
Murphy	Özyeğin University, Turkey	Celal O.B. Yavuz
LyingAgent	University of Southern California, USA	Zahra Nazari

3.1.1 Pinocchio (Baseline)

The baseline agent was provided as source code to all participants of the competition. Pinocchio followed a straightforward strategy—it attempted to gain information about the human’s preferences, then made fair partial offers, giving away one item at a time until all items were assigned. These offers were made in response to user offers, and Pinocchio did not generate offers on its own outside of these counter-offers. It also used only positive emotions. Pinocchio has been discussed in detail in previous work, as a part of the IAGO toolkit [9].

3.1.2 Elphaba

Elphaba used a mixture of positive and negative emotions (anger, happiness, and sadness), depending on a user “reliability” score. Reliability increased when user preferences were detected, and decreased when contradictions in user statements were exposed. From this, Elphaba grew value for both human and agent.

3.1.3 Agent Cena

Agent Cena was an aggressive agent, utilizing tough statements and negative emotions to attempt to browbeat opponents into giving up value. Cena started with unreasonably high offers (90% of the total value), and slowly conceded, accelerating in the last 3 minutes of the negotiation. Cena did attempt to give the human player their most valuable item first.

3.1.4 Agent Wotan

Wotan was similar to Cena in that it started with a high offer and conceded, although it would concede further and faster than Cena. Furthermore, Wotan only made full offers, in an attempt to save time. Wotan also conceded more if it thought the total value would increase. Wotan used neutral text and emotions.

3.1.5 Agent Murphy

Agent Murphy utilized several innovations, including a preference graph which was updated after user statements. It also took a pessimistic view of the human’s preferences, tending to assume the worst possible outcome. Murphy also attempted to lighten the mood with occasional jokes.

3.1.6 LyingAgent

LyingAgent focused on perpetrating a type of lie often referred to as the “fixed-pie lie”. By misleading the human player into believing they were seeking the same items, it was able to concede items that seemed valuable, but were actually nearly worthless to it. This, plus its generally friendly demeanor, allowed it to appear fair while claiming more value than its opponent.

3.1.7 Overview

As with human-human negotiation, there are many effective tactics that can lead to success in agent negotiation. Among the agents that were submitted to this competition, for example, there are several that use emotion in an attempt to influence their opponent. This strategy (particularly the use of negative emotions to gain concession) has been well-documented both in human-human and human-agent contexts [16]. Agent Cena, the runner-up agent in terms of points scored, applied this tactic. Conversely, some agents in the competition attempted to build rapport with the human-participant through positive emotion, with the hope that it would lead to greater value. The top-rated agent for likeability, Agent Wotan, used this strategy.

Another major tactic used in the competition was deception. The aptly-named LyingAgent was able to employ this strategy to “grow the pie” of the negotiation and claim more than its fair share of the result. Lying in negotiation is a well-established tactic, and indeed the LyingAgent has been subsequently described in detail in publications both before and after the competition [2,13,14]. The ethics of these tactics have also been explored [15].

All the agents also model their opponent to an extent, although Agent Murphy and Elphaba take the greatest steps in doing so, both incorporating a reliability metric for how certain they are of opponent preferences. Agent Murphy, additionally, takes a pessimistic view of the potential outcomes.

4 RESULTS & ANALYSES

4.1 Method

For the purposes of the competition, all agents' score and likeability rating (see below) averages were compared one-to-one. Dunnett's 2-sided test confirmed any significant differences for one-way contrasts against the baseline agent, Pinocchio. Significant differences between submitted agents were determined with post-hoc analysis, using Bonferroni correction. For winners, selection of first-place vs. runner up was broken in the direction of the trend for prize-awarding purposes, if no significant difference was observed. Likeability was determined by a series of self-reported 7-point Likert questions that were answered by the participants after negotiation:

- How satisfied were you with the final agreement?
- How much do you like your opponent?
- Would you negotiate with this opponent again?

Cronbach's α for these was 0.880, indicating high reliability.

4.2 Likeability

All agents were less likeable than the baseline agent, although the difference between Agent Wotan and the baseline was only marginally significant ($p = .068$). Specifically, the baseline had a mean likeability rating that was 1.18 points higher than Wotan. Agent Wotan and Elphaba were the two most likeable, after the baseline (Figure 2). Wotan and Elphaba were also not significantly different from each other, with a mean likeability difference of only .287 ($p = 1.00$). However, Elphaba was significantly worse than the baseline (differing by 1.47, $p = .010$).

4.3 Agent Score

Agent score took into account the total agent points earned only. Here, the top performing agent was the LyingAgent. Agent Cena was the runner-up in terms of agent point score, and both it and the LyingAgent scored significantly more than the baseline per ANOVAs with post-hoc Bonferroni corrections. Agent Cena featured a mean difference of 6.86 points over baseline ($p < .01$). Similarly, the LyingAgent beat the baseline with a mean difference of 9.64 points ($p < .001$). The difference between Agent Cena and the LyingAgent was not significant (Figure 3).

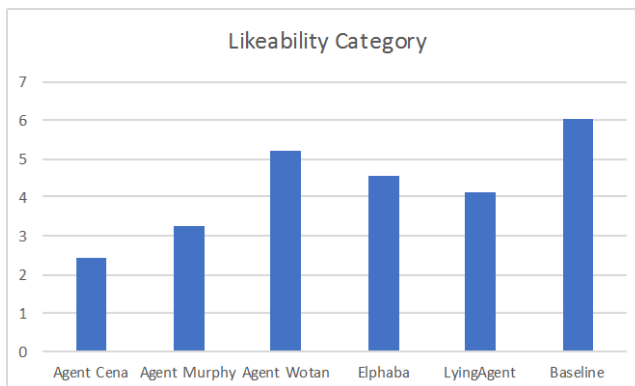


Figure 2: Agent Likeability (7-Point Likert)

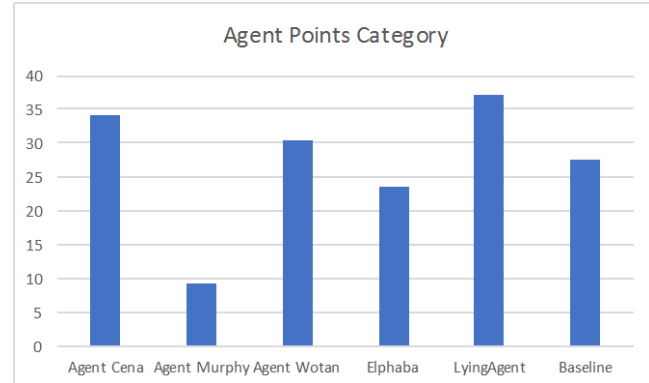


Figure 3: Agent Points (Averages)

4.4 Agent Behaviors' Effects on Score

One goal of competitions such as ANAC is to try to determine the specific antecedents to good outcomes when possible—can we determine what sort of behaviors from the agents do lead to increased points, for example? Post-hoc regression analysis of some of the more salient factors leads to takeaways for agent design.

The number of offers that the agents make varies largely between agents. And while more offers mean more information exchanged, such a torrent of offers may reveal more about how disagreeable the human partner was, especially if the agents often counter-offer. And indeed, this is the case, although the effects differed even when within the same category of agents (i.e., the "Likeability winners" vs. the "Point winners"). Specifically, the LyingAgent tended to receive more points when it made fewer offers ($t = -3.378$, $N = 21$, $p = .003$). Elphaba also received more points when it made fewer offers ($t = -2.143$, $N = 25$, $p = .043$). Conversely, Agent Cena received more points as it made more offers ($t = 2.354$, $N = 25$, $p = .027$). The other agents did not present significant relationships.

These results can be further examined by combining them with the amount of information the agent had available to it. All the agents attempted to model user preferences—indeed Agent Murphy and Elphaba had explicit measures of reliability. Human participants, however, tend to make few offers themselves, so agents often must rely on explicit statements of preferences by the human users themselves. We can find a simulacrum of user model accuracy from the data by examining the number of (truthful) statements that users made about their own preferences. Sadly, the majority of the submitted agents did not adequately utilize this information. No significant correlations were found between the number of user preference statements made and the agent's point outcome, except for Elphaba's case, where the result was distinctly *negative* ($t = -2.222$, $N = 25$, $p = .036$).

Nevertheless, information is only as good as it is used. If an agent has a good user model (due to having preference information), then it might only be useful if the agent actually uses this information to make strategically informed decisions. Specifically, the LyingAgent required a very accurate user model in able to successfully perpetrate its lie and claim value. When examining the relationship between the number of offers made by

the agent and the amount of user information is gathered, the LyingAgent shows a significant gain. Explicitly, there exists a two-way interaction between the number of offers made by the agent and the number of preference statements received from the user for the LyingAgent.¹ Specifically, when the LyingAgent makes many offers, it performs better if it *also* has access to the user's preferences, while it performs much worse if it does not. This interaction is detailed in Figure 4. There was a main, negative effect of the number of offers made by the agent, controlling for the number of preferences expressed by the human ($t = -3.792$, $N = 21$, $p = .001$). There was a main, positive effect of the number of preferences, controlling for the number of offers ($t = 2.424$, $N = 21$, $p = .027$). Finally, there was an interaction between offers and preferences ($t = 2.172$, $N = 21$, $p = .044$).

4.5 Agent Behaviors' Effects on Likeability

While the point value of a single negotiation represents a good metric for an agent's success, over time, this may be eclipsed by how much the agent is liked by its opponent. Given the choice, people do not often return to negotiations with people (or agents) that they dislike. Therefore, likeability provides another valid metric for measuring success. For the top-scoring agent in likeability (Agent Wotan), the number of messages sent by the agent was correlated to its likeability ($t = 2.564$, $N=23$, $p = .018$). However, for one of the top-scoring agents in points (Agent Cena), the correlation was significant, and negative ($t = -2.125$, $N=25$, $p = .045$). In short, when Agent Cena sent messages, it reduced the likeability rating, while when Agent Wotan sent messages, the likeability rating increased (Figure 5). These results are perhaps not surprising, since in negotiation, the content of the message is important. Agent Cena, whose dialogue consisted of aggressive remarks, had a markedly different effect on players than did its more conciliatory cousin, Wotan.

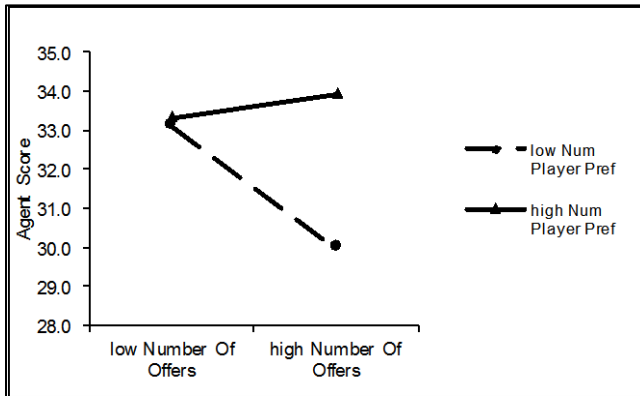


Figure 4: Winning Agent (LyingAgent) – Score by User Preference Statements and Offers Made by Agent

¹ Note that while the users CAN lie about their preferences, this behavior is tracked in IAGO's data, and is uncommon.

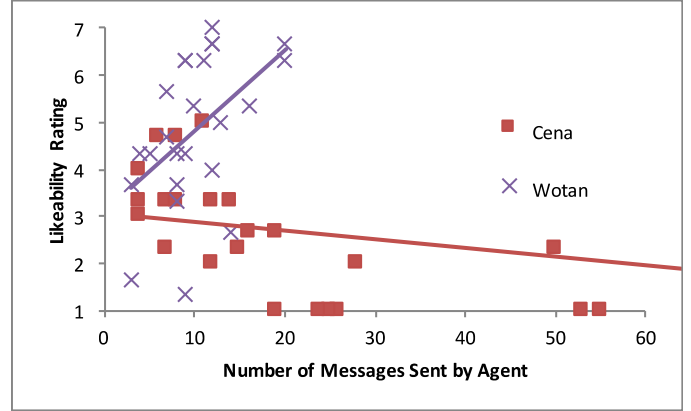


Figure 5: Significant Likeability & Agent Messages Interaction

Because human negotiators are often concerned with ideas of fairness, agents that can successfully seem “fair” most likely have a good user model (and good strategic use of it). The final distribution of the points in the negotiation can be examined to determine the actual fairness of a deal (although this does NOT correlate exactly to the perceived fairness of a deal, a fact exploited by the LyingAgent). By examining the percentage of the total points that went to the agent, the relationship between this outcome and likeability can be analyzed. For two of the top-scoring agents (Cena and Wotan), this relationship was significant: ($t = -3.37$, $N = 25$, $p = .003$) for Cena and ($t = -4.34$, $N = 24$, $p < .001$) for Wotan. When either agent took more than its fair share of the points in a negotiation, its likeability suffered. Interestingly, this negative trend existed for all agents, except the LyingAgent, where the trend was slightly positive (but not significant). See Figure 6.

5 DISCUSSION & CONCLUSIONS

The agents submitted to this competition attempted to solve the problems required of intelligent negotiators: they modeled their opponents, determined their own strategies, and adapted to human behavior. In doing this, they followed a number of varied strategies which utilized all the channels supported within IAGO: messaging, offer exchange, and emoting. The two top-performing agents relied on using advanced tactics such as expression of anger (Agent Cena) or lying about the agent's own preferences in order to claim value (LyingAgent). Furthermore, the two most likeable agents (beyond the baseline), also utilized emotion and dialogue to create rapport with the human participant (Agent Wotan & Elphaba).

What is perhaps most striking about these results is how varied the effective tactics seem to be. Agents that send a great deal of messages to the player can lead to increased likeability ratings (in the case of Agent Wotan), or reduced ratings (Cena) but this is not universal. And while the relationship between scoring well and liking one's opponent is certainly predictable, these results illustrate a very important point about how humans perceive outcomes. The LyingAgent worked by convincing humans

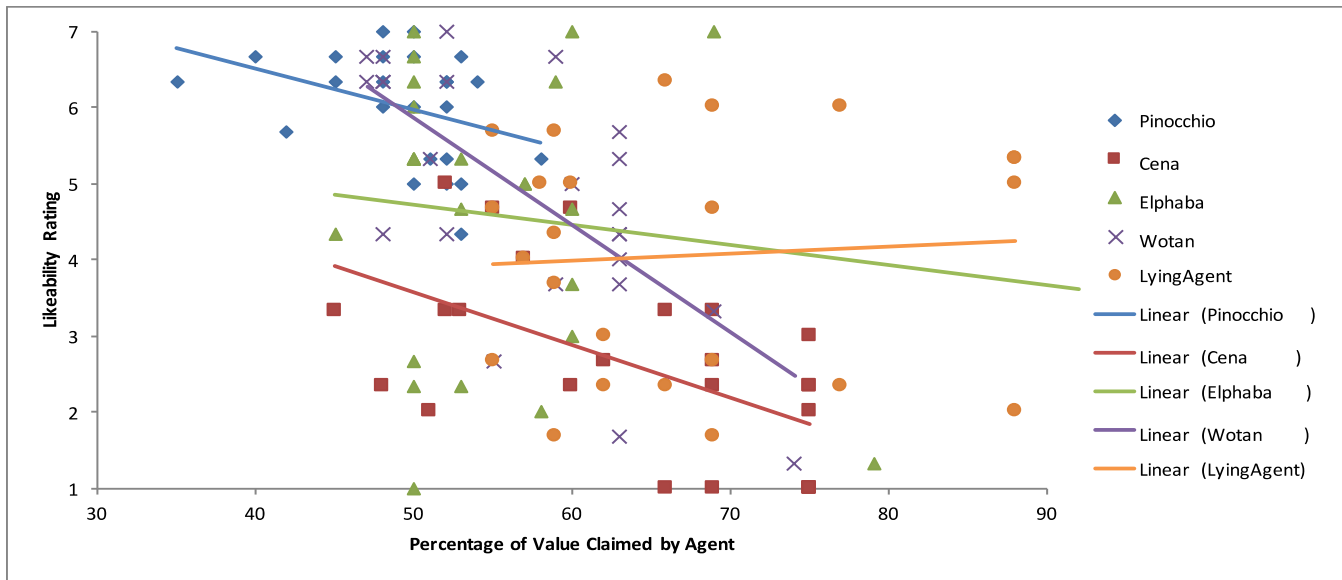


Figure 6: Scatterplot of Likeability Rating vs. Percentage of Value Claimed by Agent (Trend-lines included)

that they were getting a fair deal (when they were actually much closer to a 70/30 split). The LyingAgent had a great deal of variance in how it was perceived (see Figure 6), but when it was able to successfully sell its lie, it appeared to do well. In the future, we hope to reduce the divide between the categories of agents into “likeability” agents and “point-scoring” ones. Future competitions within ANAC in the Human-Agent League will involve repeated negotiations, in which reputational effects may cause the long-term benefits of likeable strategies to shine. Furthermore, short-sighted strategies such as those employed by Agent Cena and LyingAgent may backfire, with scores plummeting in later rounds. For these reasons and more, we look forward to the continued evolution of negotiating agents, and their performance in future competitions.

ACKNOWLEDGMENTS

This paper is sponsored by U.S. Air Force OSR FA9550-14-1-0364 as well as by NSF 1419621. Statements and opinions expressed do not necessarily reflect the position of the US Government. This work is part of the Veni research program with project number 639.021.751, which is financed by the Netherlands Organization for Scientific Research (NWO). Finally, we wish to thank the participants of the competition and the ANAC Board.

REFERENCES

- [1] Kathleen M. Alto, Keiko M. McCullough, and Ronald F. Levant, 2018. Who is on Craigslist? A novel approach to participant recruitment for masculinities scholarship. *Psychology of Men & Masculinity*, 19(2), p.319.
- [2] Karl Aquino and Thomas E. Becker, 2005. Lying in negotiations: How individual and situational factors influence the use of neutralization strategies. *Journal of Organizational Behavior*, 26(6), pp.661-679.
- [3] Joost Broekens, Maaïke Harbers, Willem-Paul Brinkman, Catholijn M. Jonker, Karel Van den Bosch, and John-Jules Meyer, 2012, September. Virtual reality negotiation training increases negotiation knowledge and skill. In *International Conference on Intelligent Virtual Agents* (pp. 218-230). Springer, Berlin, Heidelberg.
- [4] Celso M. de Melo, Stacy Marsella, and Jonathan Gratch, 2016, May. Do as I say, not as I do: Challenges in delegating decisions to automated agents. In *Proceedings of the 2016 International Conference on Autonomous Agents & Multiagent Systems* (pp. 949-956). International Foundation for Autonomous Agents and Multiagent Systems.
- [5] Jonathan Gratch, David DeVault, Gale M. Lucas, and Stacy Marsella, 2015, August. Negotiation as a challenge problem for virtual humans. In *International Conference on Intelligent Virtual Agents* (pp. 201-215). Springer, Cham.
- [6] Catholijn M. Jonker, Reyhan Aydogan, Tim Baarslag, Katsuhide Fujita, Takayuki Ito, and Koen V. Hindriks, 2017. Automated Negotiating Agents Competition (ANAC). In *AAAI* (pp. 5070-5072).
- [7] Yanviv Leviathan, and Yossi Matias, 2018, May. Google Duplex: An AI System for Accomplishing Real-World Tasks over the Phone. Blog post, downloaded from <https://ai.googleblog.com/2018/05/duplex-ai-system-for-natural-conversation.html>.
- [8] Raz Lin, Yinon Oshrat, and Sarit Kraus, 2009, May. Investigating the benefits of automated negotiations in enhancing people's negotiation skills. In *Proceedings of The 8th International Conference on Autonomous Agents and Multiagent Systems-Volume 1* (pp. 345-352). International Foundation for Autonomous Agents and Multiagent Systems.
- [9] Johnathan Mell, and Jonathan Gratch, 2017, May. Grumpy & Pinocchio: Answering Human-Agent Negotiation Questions through Realistic Agent Design. In *Proceedings of the 16th Conference on Autonomous Agents and Multiagent Systems* (pp. 401-409). International Foundation for Autonomous Agents and Multiagent Systems.
- [10] Johnathan Mell, and Jonathan Gratch, 2016, May. IAGO: interactive arbitration guide online. In *Proceedings of the 2016 International Conference on Autonomous Agents & Multiagent Systems* (pp. 1510-1512). International Foundation for Autonomous Agents and Multiagent Systems. IAGO
- [11] Johnathan Mell, Gale Lucas, and Jonathan Gratch, 2018. Welcome to the Real World: How Agent Strategy Increases Human Willingness to Deceive. In *Proceedings of the 2018 International Conference on Autonomous Agents and Multiagent Systems*. International Foundation for Autonomous Agents and Multiagent Systems.
- [12] Robert H. Mnookin, and Lawrence E. Susskind, eds., 1999. Negotiating on behalf of others: advice to lawyers, business executives, sports agents, diplomats, politicians, and everybody else (Vol. 1). Sage Publications.
- [13] Zahra Nazari, Gale M. Lucas, and Jonathan Gratch, 2017, August. Fixed-pie Lie in Action. In *International Conference on Intelligent Virtual Agents* (pp. 287-300). Springer, Cham.
- [14] Mara Olekalns, and Philip L. Smith, 2009. Mutually dependent: Power, trust, affect and the use of deception in negotiation. *Journal of Business Ethics*, 85(3), pp.347-365.
- [15] James J. White, 1980. Machiavelli and the bar: Ethical limitations on lying in negotiation. *Law & Social Inquiry*, 5(4), pp.926-938.
- [16] Gerben A. Van Kleef, Carsten KW De Dreu, and Antony SR Manstead, 2004. The interpersonal effects of anger and happiness in negotiations. *Journal of personality and social psychology*, 86(1), p.57.