

Pre-proceedings of the Twelfth International Workshop on

MULTI-AGENT BASED SIMULATION

MABS 2011

Held as a satellite event of the

Tenth International Conference on Autonomous Agents
and Multiagent Systems (AAMAS 2011)

Taipei, Taiwan, 2 May 2011

Workshop Organizers

Jordi Sabater-Mir
Jaime Sichman
Daniel Villatoro

Programme Committee

Diana Francisca Adamatti
Shah Jamal Alam
Fred Amblard
Joao Balsa
Riccardo Boero
Tibor Bosse
Fr. Bousquet
Jean-Pierre Briot
Cristiano Castelfranchi
Helder Coelho
Nuno David
Paul Davidsson
Gennaro Di Tosto
Bruce Edmonds
Armando Geller
Joseph Giampapa
Nigel Gilbert
Francisco Grimaldo
Laszlo Gulyas
Samer Hassan
Rainer Hegselmann
Cesareo Hernandez
Mark Hoogendoorn
Marco Janssen
Catholijn Jonker
Maciej Latek
Michael Lees
Jean-Pierre Muller
Pablo Noriega
Emma Norling
Paulo Novais
Mario Paolucci
Isaac Pinyol
Juliette Rouchier
Keith Sawyer
Elizabeth Sklar
Keiki Takadama
Oswaldo Terán
Klaus G Troitzsch
Manuela Veloso

External Reviewers

Inacio Guerberooff
Eric Guerci
Narine Udumyan

Contents

Tongkui Yu and Shu-Heng Chen. <i>Agent-based Model of the Political Election Prediction Market</i>	1
Ondrej Vanek, Michal Jakob, Ondrej Hrstka and Michal Pechoucek. <i>Using Multi-Agent Simulation to Improve the Security of Maritime Transit</i>	12
Benoit Gaudou, Andreas Herzig, Emiliano Lorini and Christophe Sibertin-Blanc. <i>How to do social simulation in logic: modelling the segregation game in a dynamic logic of assignments</i>	24
Nardine Osman, Jordi Sabater Mir and Carles Sierra. <i>Simulating Research Behaviour</i>	36
Francisco Grimaldo, Mario Paolucci and Rosaria Conte. <i>Agent Simulation of Peer Review: The PR-1 Model</i>	48
Lorenza Manenti, Sara Manzoni, Giuseppe Vizzari, Kazumichi Ohtsuka and Kenichiro Shimura. <i>An Agent-Based Proxemic Model for Pedestrian and Group Dynamics: Motivations and First Experiments</i>	60
Xing Pengfei, Michael Lees, Hu Nan and Vaisagh Viswanthathn T.. <i>Validation of Agent-Based Simulation through Human Computation : An Example of Crowd Simulation</i>	72
Gildas Morvan, Alexandre Veremme and Daniel Dupont. <i>Observation of large-scale multi-agent based simulations</i>	84
Jason Tsai, Emma Bowring, Stacy Marsella and Milind Tambe. <i>Modeling Emotional Contagion</i>	94
H. Van Dyke Parunak. <i>Between Agents and Mean Fields</i>	106

Agent-based Model of the Political Election Prediction Market

Tongkui Yu^{1,2} and Shu-Heng Chen¹ *

1. AI-ECON Research Center, Department of Economics, National Chengchi University, Taipei 11605

2. Department of Systems Science, School of Management, Beijing Normal University, Beijing 100875

Abstract. We propose a simple agent-based model of the political election prediction market which reflects the intrinsic feature of the prediction market as an information aggregation mechanism. Each agent has a vote, and all agents' votes determine the election result. Some of the agents participate in the prediction market. Agents form their beliefs by observing their neighbors' voting disposition. In this model, the mean price can be a good forecast of the election result. We study the effect of the agents' neighbor scope and information distribution on the prediction accuracy. In addition, we also find the mechanism which can replicate the favorite-longshot bias, a stylized fact in the prediction market. This model can provide a framework for much further analysis on the prediction market such as complex agent behavioral patterns.

Keywords: Prediction market, Agent-based simulation, Information aggregation mechanism, Prediction accuracy, favorite-longshot bias

1 Introduction

Prediction Markets, sometimes referred to as “information markets,” “idea futures” or “event futures”, are markets where participants trade contracts whose payoffs are tied to a future event, thereby yielding prices that can be interpreted as market-aggregated forecasts [1]. To predict whether a particular event (say, some candidate winning the election) will happen, a common approach is to create a security that will pay out some predetermined amount (say, 1 dollar) if the event happens, and let agents trade this security until a stable price emerges; the price can then be interpreted as the consensus probability that the event will happen. There is mounting evidence that such markets can help to produce forecasts of event outcomes with a lower prediction error than conventional forecasting methods [2].

To explain the efficiency of the prediction market in relation to aggregate information, many researchers use the efficient market hypothesis which attributes the market efficiency to a pool of knowledgeable traders who are capable of setting prices and acting without bias [3]. While Manski proposed a model [4] to

* Corresponding author: chen.shuheng@gmail.com

find that price is a particular quantile of the distribution of traders' beliefs, price does not reveal the mean belief that traders hold, but does yield a bound on the mean belief. It can explain both the market efficiency and another stylized fact in prediction markets - favorite-longshot bias - which means that likely events (favorite) are underpriced and unlikely events (longshot) are overpriced [5]. Wolfers and Zitzewitz provided sufficient conditions under which prediction market prices coincide with average beliefs among traders in a model with log-utility agents [6]. Snowberg and Wolfers found evidence that misperceptions of probability drive the favorite-long shot bias, as suggested by prospect theory [7]. Instead of complex assumptions and technique modeling, we provide an intuitive model of a political election prediction market which reflects the intrinsic feature of prediction market as an information aggregation mechanism. This model can replicate the prediction accuracy and favorite-longshot bias of the prediction market. Based on this model, we analyze the effect of agent information scope, i. e., the distribution of information among agents on market accuracy. We also diagnose the origin of the favorite-longshot bias. This model can also provide a framework for much further analysis on the prediction market such as the complex agent behavioral patterns.

2 Basic Model

In a society represented by a two-dimensional torus grid, each element is a person, and each person has his own disposition on the vote for some candidate, blue for supporting (1) and green for not supporting (0) as in Figure 1. This is just a illustration, we use the grid size 200×200 in the simulation. Furthermore, in the simulations, the agents' voting disposition is randomly initialized with a given overall support ratio *SupportRatio*.

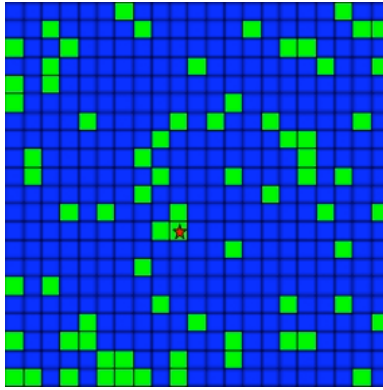


Fig. 1. Illustration of a voting society in a grid.

Some agents randomly sampled from the entire population will participate in a prediction market which provides a winner-take-all contract that pays a dollar if the candidate wins, and pays nothing if the candidate loses. Each agent has a belief $b_i \in [0, 1]$ that a candidate will win. An agent forms his belief by observing the voting disposition of all agents in his neighborhood. The neighborhood consists of the center node (the agent himself) and the nearest Moore neighbors. The neighborhood has 9 agents for neighborhood radius $r = 1$, and 25 agents for neighborhood radius $r = 2$, as in Figure 2. For the agent marked with the red star, his own vote disposition is to support the candidate, and there are two other agents in his neighborhood with radius $r = 1$ who also support the candidate (Figure 2, left panel), so his belief is $b_i = 3/9 = 1/3$. If the radius $r = 2$ (Figure 2, right panel), his belief is $b_i = 5/25 = 1/5$.



Fig. 2. Moore neighborhood and belief formation

An agent decides to either place a buy (bid) or sell (ask) order with equal probability. Then the agent with belief b_i places a bid order for one share of the event at a price *uniformly* on $[0, b_i]$, or an ask order for one share at a price *uniformly* on $[b_i, 1]$, as in figure 3. Agents do not observe current market prices, and do not react to any result of their actions; they keep no record of previous orders or if those orders resulted in trades. This is a device of the *zero-intelligence agent* initiated by Gode and Sunder [8], which is now widely used in agent-based models. The zero-intelligence agent is a randomly-behaving agent or, more precisely speaking, an entropy-maximizing agent. Since normal traders would not propose or accept a deal which would obviously lead to economic loss or not lead to welfare improvement, under no further information on what else they will do, the design of the zero-intelligence agent is *minimally prejudiced* in the sense of entropy maximization. The *uniform distribution* is employed here to realize maximum entropy. Othman also carried out this design in his pioneering study on the agent-based prediction market [9].

Following [9], the transactions are closed by the mechanism of continuous double exchange. Agents place buy or sell orders continuously. Once the highest-priced bid exceeds the lowest-priced ask, a trade occurs at the price of the order which was placed first. The paper, however, differs from [9] by explicitly embedding the agent-based prediction market within the network structure, a checkerboard as demonstrated in Figure 1. We consider this as a first attempt to

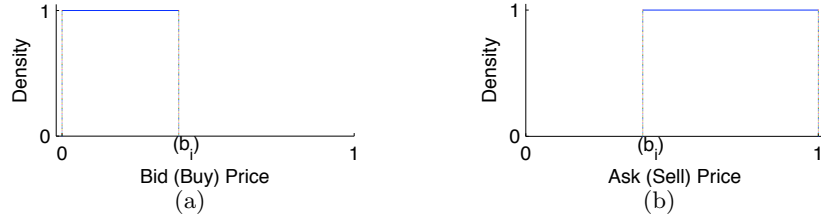


Fig. 3. Order price formation

hybridize the spatial agent-based political election models and the agent-based prediction markets.

The procedure of the model is described by the pseudocode in Algorithm 1.

3 Simulation and Analysis

We perform a great deal of simulations and try to find the regularity in the simulated data.

3.1 Prediction power and favorite-longshot bias

In the agent-based model of the political election prediction market, the mean of the transaction price is a good predictor of the real support ratio in the overall population. Figure 4 provides the relationship between the real support ratio (vertical axis) and the mean of the transaction price (horizontal axis) in typical simulations with the neighborhood radius $r = 1$. This simulation is conducted in a 200 by 200 checkerboard with 40,000 agents ($N = 40,000$), one in each cell (a checkerboard with full size). For each given support ratio, the continuous double auction is run once for 40,000 rounds. We then try support ratios from 0.1 to 0.9 with an increment of 0.01; in other words, a total of 91 support ratios are tried. We then plot the mean of the transaction price over the 40,000 rounds for each support ratio in Figure 4. As we can see from that figure, the mean transaction price can trace the true support ratio to some degree.

At the same time, the simulation replicates the favorite-longshot bias, a stylized fact in the prediction market. We can find that likely events (bottom-left in Figure 4) are underpriced, and unlikely events (top-right in Figure 4) are overpriced.

3.2 Neighborhood scope and prediction accuracy

We investigate the effect of the neighborhood scope on the prediction accuracy. Figure 5 provides the relationship between the real support ratio (vertical axis) and the mean of the transaction price (horizontal axis) in typical simulations with different neighborhood radii $r = 1, 2$ and 3. We can find that the prediction

```

Input: Total number of agents  $N$ ; Rounds of the market  $M$ ; The overall
        support ratio  $SupportRatio$ ;
Output: Transaction price history  $TransactionPrice$ 
// INITIALIZATION
Generate a random number  $rndNum$  uniformly from 0 to 1;
foreach agent in the population do
    if  $rndNum < SupportRatio$  then
        | Set his voting disposition as supporting (1);
    else
        | Set his voting disposition as not supporting (0);
    end
end
// RUNNING THE MARKET
foreach  $round \in [1, M]$  do
    Choose an agent randomly from the whole population ;
    Get the voting dispositions of his neighbors ;
    Set the agent's belief  $b$  as the number of supporters over neighborhood size ;
    Generate a random number  $rndNum$  uniformly from 0 to 1 ;
    if  $rndNum < 0.5$  then
        | Set the order side to 1 (buy) ;
        | Set the  $OrderPrice$  as a random number uniformly drawn from 0 to  $b$  ;
        | Get the  $MinSellPrice$  in the  $SellOrderList$  ;
        if  $OrderPrice > MinSellPrice$  then
            | Insert a transaction in the  $TransactionList$  with  $MinSellPrice$  ;
        else
            | Insert a buy order with  $OrderPrice$  in the  $BuyOrderList$  ;
        end
    else
        | Set the order side to 0 (sell) ;
        | Set the  $OrderPrice$  as a random number uniformly drawn from  $b$  to 1 ;
        | Get the  $MaxBuyPrice$  in the  $BuyOrderList$  ;
        if  $OrderPrice < MaxBuyPrice$  then
            | Insert a transaction in the  $TransactionList$  with  $MaxBuyPrice$  ;
        else
            | Insert a sell order with  $OrderPrice$  in the  $SellOrderList$  ;
        end
    end
end

```

Algorithm 1: Simulation Procedure

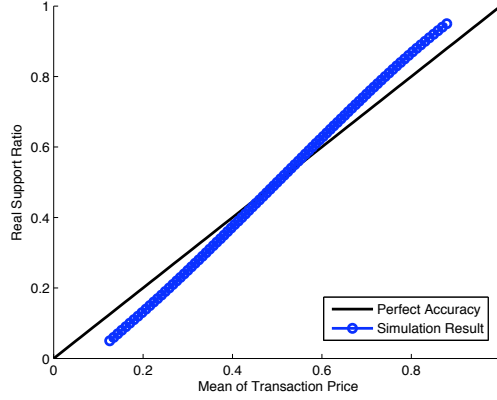


Fig. 4. Typical simulation result and favorite-longshot bias.

accuracy of the prediction market increases as the neighborhood scope increases. It is easy to understand that the more information that each participant has, the more accurate the prediction that the market can provide. This result is similar to the work of Othman [9]. However, we obtain the result with a totally different basic assumption. We use the most intuitive assumption that the agent forms his belief by observing his neighbors' voting deposition, while Othman's work just specifies the belief distribution arbitrarily. We attribute the extent of the bias to the amount of individual information, while Othman's work attributes it to the arbitrarily specified belief distribution.

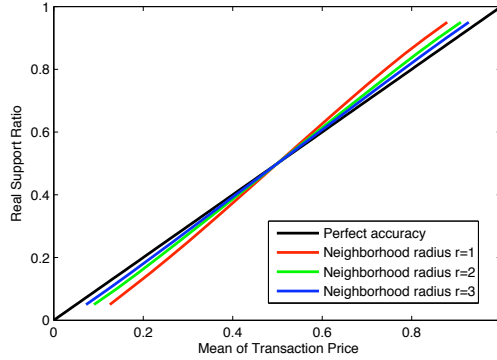


Fig. 5. Neighborhood radius and prediction accuracy.

To measure the prediction accuracy of a prediction market, we define a variable referred to as the mean squared error $\rho = \frac{\sum_1^N (\text{MeanPrice}_i - \text{SupportRatio}_i)^2}{N}$, where $i = 1 : N$ is the number of simulations. A larger mean squared error implies less accurate prediction, while a smaller one implies more accurate prediction, and $\rho = 0$ implies perfect accuracy. With the same neighborhood radius, we simulate $N(N = 200)$ times with *SupportRatio* linearly spaced between 0 and 1 and calculate the prediction accuracy. Figure 6 presents the effect of the neighborhood radius r (horizontal) on the prediction accuracy (vertical axis). The larger the neighborhood radius r , the more information agents have, and the more accurate prediction the market can provide.

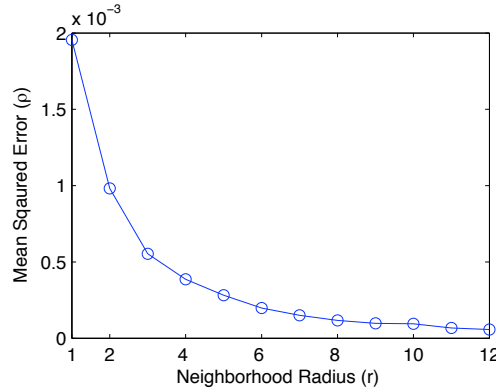


Fig. 6. Effect of neighborhood radius on the prediction accuracy.

3.3 Information distribution and prediction accuracy

In a real political election, the voting disposition may be clustered. Some places may be dominated by the supporter of one candidate, and most of the population may support the other candidate elsewhere. So the information is not well scattered. We wish to study the effect of information distribution on the prediction accuracy. To this end, we use the voting disposition block size to represent different information distributions.

When initializing the voting disposition of the agents, we use different granularities or block size s . If the block size $s = 1$, we initialize the agents' voting disposition one by one; if the block size $s=2$, we initialize the agents' voting disposition four by four, i.e., all four agents in the 2×2 sub-grid have the same voting disposition randomly generated according to the specified support ratio. Figure 7 illustrates a comparison of information distributions with block sizes $s = 1$ and 2 under the same support ratio 0.3.

Figure 8 depicts the effect of voting disposition block size on prediction accuracy with neighborhood radii $r = 2, 3, 5$ and 7. We can find the nonlinear

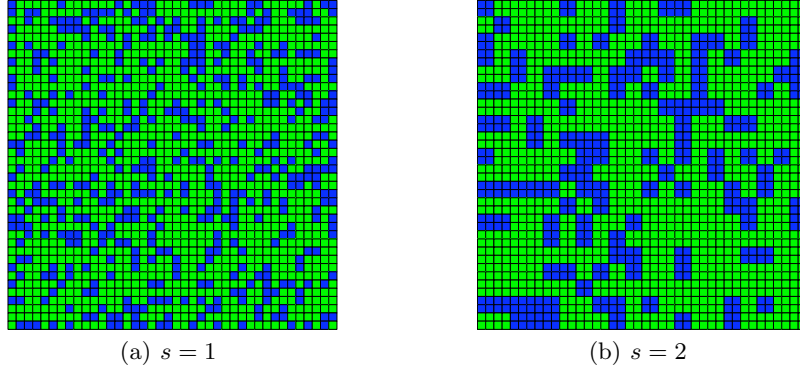


Fig. 7. Comparison of information distributions with different voting disposition block sizes

relationship between the voting disposition block size and mean squared error. The mean squared error is the largest when the voting disposition block size is close to the neighborhood radius, and the market prediction power is the least accurate. The farther away the voting disposition block size is from the neighborhood radius, the smaller is the mean squared error, the more accurate is the market prediction. Further research is needed to understand this phenomenon.

3.4 The origin of the favorite-longshot bias

This model can only replicate the favorite-longshot bias, but cannot answer the question as to what is the origin of the favorite-longshot bias. To find the mechanism that produces the favorite-longshot bias, we perform some experiments. One possible origin of the favorite-longshot bias is the order price formation mechanism, which closes transactions at the price of the order first placed. We thus try another order-matching mechanism that uses the mid-price of the buy and sell order, but we find that it makes no difference. The other possible origin of the favorite-longshot bias is the order-formation mechanism, where an agent places a buy order at a price between zero and his belief b_i , and places a sell order at a price between b_i and 1 (as in Figure 3). When the agents' belief is not equal to 0.5, there is an inherent asymmetry in the order prices. This may lead to the favorite-longshot bias. We test the hypothesis by trying to incorporate a symmetric order price formation mechanism where an agent places a buy order at a price drawn randomly from $[b_i - \delta, b_i]$, and places a sell order at a price drawn randomly from $[b_i, b_i + \delta]$ as in Figure 9(a). We find that the favorite-longshot bias disappears in this specification as in 9(b). Moreover, we perform a simulation using the designed asymmetric order price formation mechanism, where the prices of buy orders are uniformly on $[b_i - \delta, b_i]$ and the prices of sell orders are uniformly on $[b_i, b_i + 2\delta]$ as in Figure 9(c), and we find that the prices are all overvalued as in figure 9(d). Furthermore, if the prices of buy orders are

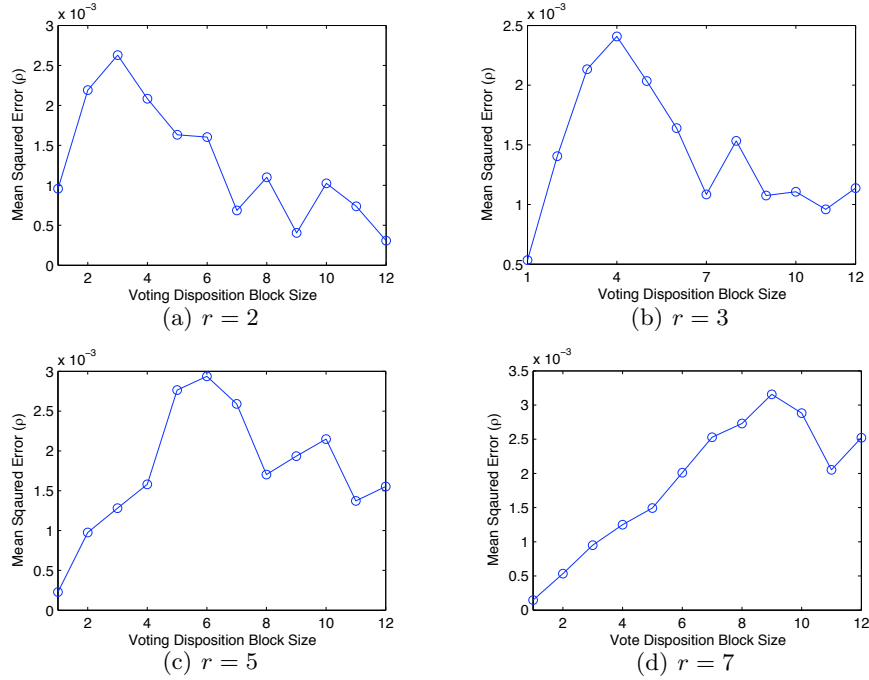


Fig. 8. The effect of voting disposition block size on prediction accuracy

uniformly on $[b_i - 2\delta, b_i]$ and the prices of sell orders are uniformly on $[b_i, b_i + \delta]$ as in Figure 9(e), the prices are all undervalued as in Figure 9(f). So we can come to the conclusion that the asymmetry of the order-price mechanism is the origin of the favorite-longshot bias.

4 Conclusion

In this paper, we build a simple agent-based model of the political election prediction market which reflects the intrinsic feature of the prediction market as an information aggregation mechanism. Each agent has a vote, and all agents' votes determine the election result. Some randomly chosen agents participate in the prediction market. They form their beliefs by observing neighbors' voting disposition. By a great amount of simulation, we find that the mean price can be a good forecast of the election result. We also study the effect of the agents' neighborhood scope and distribution of information on the prediction accuracy. The larger the neighborhood radius, the more information agents have, and the more accurate prediction the market can provide. The farther away the voting disposition block size is from the neighborhood radius, the smaller is the mean squared error, and the more accurate is the market prediction. The model can replicate the favorite-longshot bias in the prediction market. Moreover, we find

that the asymmetry of the order-price mechanism is the origin of the favorite-longshot bias based on the simulation of this model. Furthermore, this model provides a framework for many other studies on the prediction market, such as more complex agent behavioral patterns, the evolution of participation, etc.

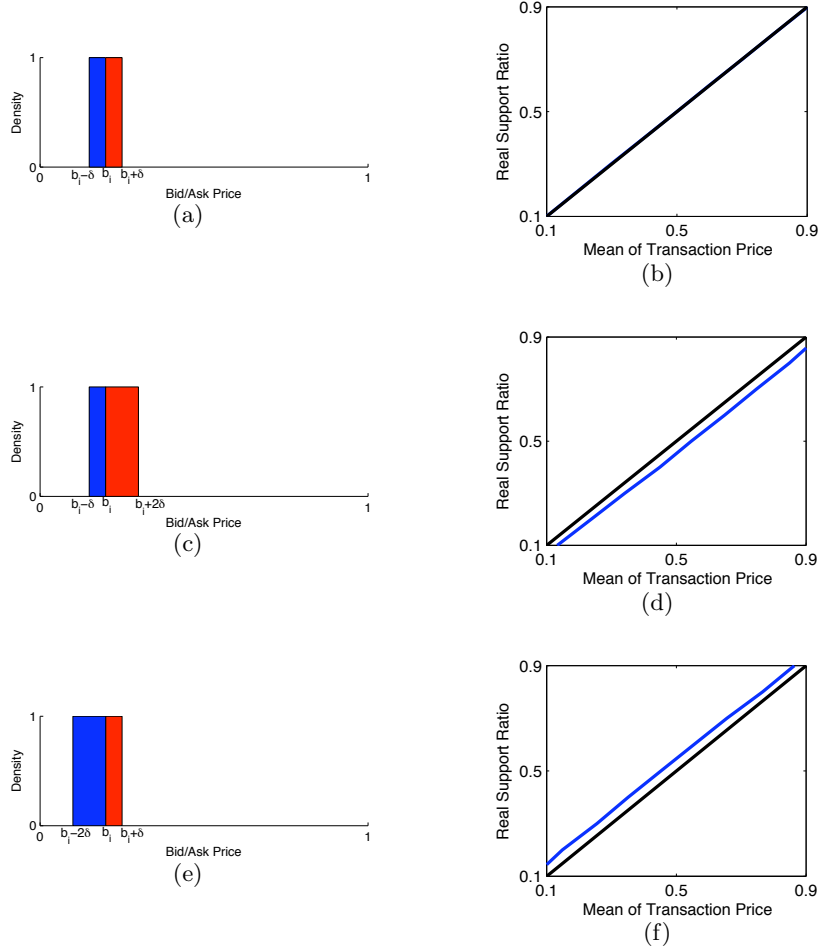


Fig. 9. Order price formation mechanism and favorite-longshot bias

References

1. Wolfers J., Zitzewitz E.: Prediction Markets. J. Econ. Perspect. 18(2), 107-126 (2004)

2. Arrow K. J., Forsythe R., Gorham M., et al.: The Promise of Prediction Markets. *Science*. 320, 877-878 (2008)
3. Forsythe R., Nelson F., Neumann G., Wright J.: Anatomy of an Experimental Political Stock Market. *Am. Econ. Rev.* 82(5), 1142-1161 (1992)
4. Manski C. F.: Interpreting the Predictions of Prediction Markets. *Econ. Lett.* 91(3), 425-429 (2006)
5. Wolfers J., Zitzewitz E.: Five Open Questions About Prediction Markets. In: R. W. Hahn (eds.) *Information Markets: A New Way of Making Decisions*. AEI Press (2008)
6. Wolfers J., Zitzewitz E.: Interpreting Prediction Market Prices as Probabilities. NBER Working Paper No. 10359 (2005)
7. Snowberg E., Wolfers J.: Explaining the Favorite-Longshot Bias: Is it Risk-Love or Misperceptions? *J. Polit. Econ.* 118(4), 723-746 (2010)
8. Gode D., Sunder S.: Allocative Efficiency of Markets with Zero-Intelligence Traders: Market as a Partial Substitute for Individual Rationality. *J. of Polit. Econ.* 101(1), 119-137 (1993)
9. Othman A.: Zero-intelligence Agents in Prediction Markets. In: *Proceedings of AAMAS 2008*, pp. 879-886 (2008)

Using Multi-Agent Simulation to Improve the Security of Maritime Transit

Ondřej Vaněk, Michal Jakob, Ondřej Hrstka, and Michal Pěchouček

Agent Technology Center, Dept. of Cybernetics, FEE, Czech Technical University,
Technická 2, Praha 6, Czech Republic

{vanek, jakob, hrstka, pechoucek}@agents.felk.cvut.cz

<http://agents.felk.cvut.cz>

Abstract. Despite their use for modeling traffic in ports and regional waters, agent-based simulations have not yet been applied to model maritime traffic on a global scale. We therefore propose a fully agent-based, data-driven model of global maritime traffic, focusing primarily on modeling transit through piracy-affected areas. The model uses finite state machines to represent the behavior of several classes of vessels and can accurately replicate global shipping patterns and approximate real-world distribution of pirate attacks. The application of the model to the problem of optimizing the Gulf of Aden group transit demonstrates the usefulness of agent-based modeling in evaluating and improving counter-piracy measures.

1 Introduction

Due to its inherent dynamism, distribution and complexity of dependencies, traffic and transportation is a domain particularly suitable for the application of multi-agent techniques. This has been reflected in the number of multi-agent simulations developed in the field of air traffic (e.g. [15]) and ground transportation [1, 11]. The key motivation behind these models is to better understand the traffic dynamics and to evaluate novel mechanisms for improving its properties.

In the maritime domain, existing models focus on traffic in ports and regional, nearshore waters [2, 12, 8]. High-level equation-based models are typically used, which have difficulties capturing vessel interactions and more complex dynamics of maritime traffic. In contrast, our model focuses on modeling *global* maritime traffic and employs the fully agent-based, microsimulation approach. Such a model is pivotal in the development of measures for countering the complex problem of maritime piracy, which presents a growing threat to global shipping industry and consequently international trade. In 2010 alone, 53 cargo vessels were hijacked and 1181 crew members held hostage [9] and the numbers continue to rise.

The proposed simulation model is the first agent-based model focusing on maritime traffic and piracy. It models the operation of all key actors in piracy scenarios, i.e., the long-range cargo vessels, pirate vessels and navy patrols. Although the simulation is geared towards maritime piracy, it is rather universal

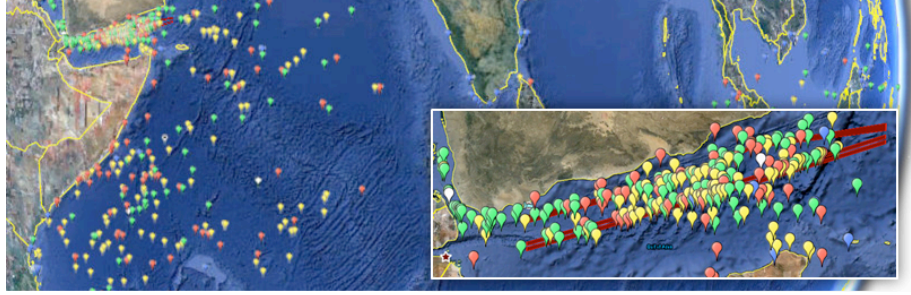


Fig. 1: Locations of pirate incidents over the last 5 years.

and could be used for other applications too. As an example of potential applications, we show how it can be used to optimize the group transit scheme established to improve the security of transit through the notorious pirate waters of the Gulf of Aden.

2 Domain Background

The agent-based model developed reflects the current state of affairs regarding maritime piracy.

2.1 Piracy around the Horn of Africa

Over the last years, waters around the Horn of Africa have experienced a steep rise in piracy. For approximately 20 thousand vessels¹ that annually transit the area, insurance rates have increased more than tenfold and the costs of piracy were estimated at up to US\$16 billion in 2008.

Even though attacks and hijackings used to be concentrated into the Gulf of Aden, in the last two years the pirates have been expanding further from the coast and attacking vessels on main shipping lanes in the Indian Ocean, more than 1500 nm from the Somali coast (see Figure 1).

2.2 Existing Antipiracy Measures

In order to counter the rising piracy threat, a number of measures have been put into effect in the Horn of Africa Area, which encompasses the Gulf of Aden and the West Indian Ocean. Since 2008, transit through the Gulf of Aden itself has been organized – transiting vessels are grouped according to their speed and directed through a narrow transit corridor which is patrolled by international naval forces.

International Recommended Transit Corridor Initially introduced in 2008, the *International Recommended Transit Corridor (IRTC)* was amended in 2009 to reflect the revised analysis of piracy in the Gulf of Aden and to incorporate

¹ about 40% of the world fleet

Speed	Entry point A – time	Entry point B – time
10 kts	04:00 GMT+3	18:00 GMT+3
12 kts	08:30 GMT+3	00:01 GMT+3
14 kts	11:30 GMT+3	04:00 GMT+3
16 kts	14:00 GMT+3	08:30 GMT+3
18 kts	16:00 GMT+3	10:00 GMT+3

Table 1: IRTC group transit schedule – entry times for vessels travelling at different speeds.

shipping industry feedback. The new corridor has been positioned further from established fishing areas, resulting in a reduction of false piracy alerts².

Navy Patrols Several naval task forces from various countries and allied forces³ operate in the Gulf of Aden (see [10] for details) to protect the transit corridor and prevent attacks on transiting vessels. Detailed information about the strategy of the naval vessels is classified. On a high-level, their coordination is based on a *4W Grid* on which areas of responsibility are assigned [4].

Group Transit Scheme In August 2010, the *Group Transit Scheme* was introduced to further reduce the risk of pirate attacks to vessels transiting the Gulf of Aden [14]. *Group transits* are designed to group ships into several speed groups in order to leverage additional protection and assurance of traveling in a group. Each transit follows a recommended route through the IRTC at a published speed and schedule, designed to avoid the highest-risk areas and time periods. There is one transit per day for each speed group (see the Table 1).

3 Model Description

We employ the agent-based modeling approach. Each vessel is implemented as an autonomous agent with its distinct behavior, capable of interacting with the simulated maritime environment and other vessels. Three categories of vessels form the core of the model: (1) long-haul cargo vessels, (2) pirate vessels and (3) navy vessels. Each behavior model is based on real-world data. Below, we describe models for each vessel category.

3.1 Cargo Vessel Model

Cargo vessels (also called merchant vessels), traveling repeatedly between the world’s large ports, form the bulk of international maritime traffic. Our agent-based model aims to achieve the same spatio-temporal distribution as the real traffic. We do not simulate the physical dynamics of vessels and we do not take into account external geographical conditions such as weather or sea currents.

² MSCHOA, The Maritime Security Centre, Horn of Africa, strongly recommends transiting vessels to follow the corridor to ensure protection from naval forces.

³ NATO, Combined Maritime Forces including Japan, China, India, Korea and others.

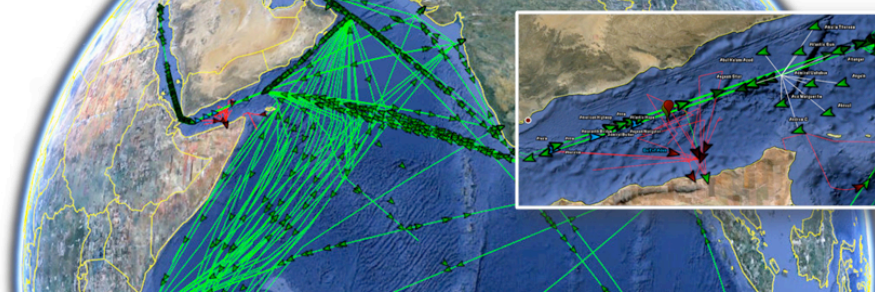


Fig. 2: Visualization of the simulation output, depicting the cargo shipping traffic (green) and pirate vessels (red).

Data Used Data from the Automatic Identification System (AIS) – the most widely used vessel position tracking system – are publicly available⁴ and contain time-stamped records of vessel GPS positions along their routes. We use AIS data to estimate the main shipping lanes and traffic density. Furthermore, we use data about geography, including locations as ports (and their capacities), straits, channels or capes through which vessels have to pass while avoiding land, islands and shallow waters.

Cargo Vessel Model The above data are combined with route planning algorithms in order to generate realistic cargo vessel traffic. As long as it is safe, vessels follow shortest paths between desired locations, avoiding obstacles. Transit through high-risk areas is handled separately to reflect vessel’s tendency to avoid such areas or to use special tactics to make such transit more secure

To approximate global properties of the real-world cargo shipping traffic, cargo-vessel agents are generated randomly in ports and are assigned a destination port to reach (taking into account port capacities). Vessel agents plan their routes between ports using the combination of the following planners:

Shortest Path Planning The shortest route planner is based on the A^* algorithm in a spherical environment with polygonal obstacles. The algorithm searches for the shortest path on a *visibility graph*, which is proven to correspond to the shortest path in the free space [13].

Risk-aware Route Planning The risk-aware planner is invoked for planning routes through high-risk areas. The planner searches for a route which minimizes a weighted sum of route length and the overall piracy risk along the route. To quantify risk along a route, we use a *risk map* representing the number of piracy incidents in a given area over a given period of time. By integrating incident density along a vessel route, we obtain the estimate of the number of incidents N that can be expected when following the route. Using the Poisson probability

⁴ Large databases are available on <http://www.vesseltracker.com> or <http://www.aislive.com>.

distribution, we can then estimate the probability of at least one attack as $1 - e^{-N}$ (for details, see [10]).

Strategic Route Planning In contrast to the previous two route planners, which reflect practices currently used in the field, the strategic route planning provides an experimental technique for further reducing piracy risk. This game-theoretic route planner explicitly accounts for pirate adaptivity and ability to reason about cargo vessel routes and produces routes in an optimally randomized manner which minimizes the chance of successful pirate attack. See [16] for details.

3.2 Pirate Vessel Model

Due to the lack of data on real-world pirate vessel trajectories, an indirect approach is used in modeling pirate behavior.

Data Used Primary source of information for pirate vessel model are anecdotal descriptions of typical pirate strategies (found e.g. in [6]) which are translated into executable behavior models (see below). These are then used in conjunction with additional data about real-world pirate activity. Specifically, we use a public dataset [5] with the information about places on the Somali coast which serve as pirate hubs or ransom anchorages. We also use piracy incident records published since 2005 by the IMB Piracy Reporting Centre⁵. Each record contains the incident location in GPS format, the type of attacked vessel, incident date and the type of the attack (depicted on Figure 1). These data are used to validate the pirate behavior model (see Section 5.2).

Pirate Behavior Model In order to capture the diversity of real-world pirate strategies (see e.g. [6]) we have implemented four different pirate vessel models: (1) *Simple pirate* with a small boat without any means of vessel detection except direct line of sight (5 nm), (2) *Radar pirate* with a radar extending the vessel detection range to 20-50 nm, (3) *AIS pirate* with an AIS interception device monitoring AIS broadcasts and (4) *Mothership pirate* with a medium-size vessel and several boats able to attack vessels up 1500 nm from the Somali coast. The last type can be combined with the radar or AIS interception device to achieve more complex behavior. Moreover, the pirate models can be extended with learning ability to model adaptation to changes in transit routing and patrolling.

Each pirate agent is initialized in its home port and based on the position of the port, it is assigned a piracy zone (Gulf of Aden, Northern or Southern part of Indian Ocean) in which it then executes its strategy.

3.3 Naval Vessel Model

The lack of data and general complexity of patrolling strategies – which can vary from active search for pirates to protecting transiting groups of cargo vessels

⁵ <http://www.icc-ccs.org/home/piracy-reporting-centre>

– makes proper modeling of navy vessel agents very difficult. We have thus proposed a minimal model based on the information available. As a potential improvement to the current practices, we have also proposed a game-theoretic policy for patrolling (see [3] for details).

Data Used The *4W Grid* [4], dividing the Gulf of Aden into square sectors, is used for anti-piracy operation coordination. The highest risk sectors are assigned to anti-piracy operation forces as their areas of responsibility.

Patrolling Behavior We have implemented a set of naval vessel agents, ranging from static reactive patrols to deliberative agents evaluating the situation in the assigned area and deciding where to patrol and which vessels to protect. Naval vessel agents are controlled by a hierarchical structure of authority agents, reflecting the real-world chain-of-command. A central *Navy Authority* agent controls a set of *Task forces* (i.e. group of agents), each disposing of a number of agents controlling warships with helicopters, able to patrol the area and respond to reported attacks. More details can be found in [10].

4 Model Implementation

We have implemented the proposed multi-agent model in Java, employing selected components of the A-lite⁶ simulation platform and using Google Earth for geo-referenced visualization. Scripts written in Groovy are used for scenario description. Below we describe the implementation of agent behavior – see [10] for a detailed description of the (rest of) implementation.

Agent Behavior Implementation Behavior model implementation has to be efficient enough to allow simulation of hundreds or even thousands of vessel agents and expressive enough to model complex behavior and interaction between vessels. The agents should be able to execute multi-step plans while handling possible interruptions. The agents share many behavior similarities (such as trajectory planning or pirate reasoning) so we want reusable models. Moreover, we want to implement the ability to learn and adapt directly in the behavior model.

Finite state machines (FSM) fit our needs well. Individual states correspond to the principal mental states of the vessel agent (such as move, attack, hijacked, patrol etc.). Transitions between the states are defined by unconditional state transitions (a pair $\{s_{from}, s_{to}\}$, e.g. $\{wait, move\}$) or by conditional transitions triggered by external events (a tuple $\{s_{from}, event, s_{to}\}$, e.g. $\{move, shipSpotted, attack\}$). Each state stores its context when deactivated so that when reactivated, the context can be restored to continue the previously interrupted plan. Transitions between states can be internally or externally triggered, thus allowing interruption of plans or actions. FSMs are easily extensible

⁶ <http://agents.felk.cvut.cz/projects/#a-lite>

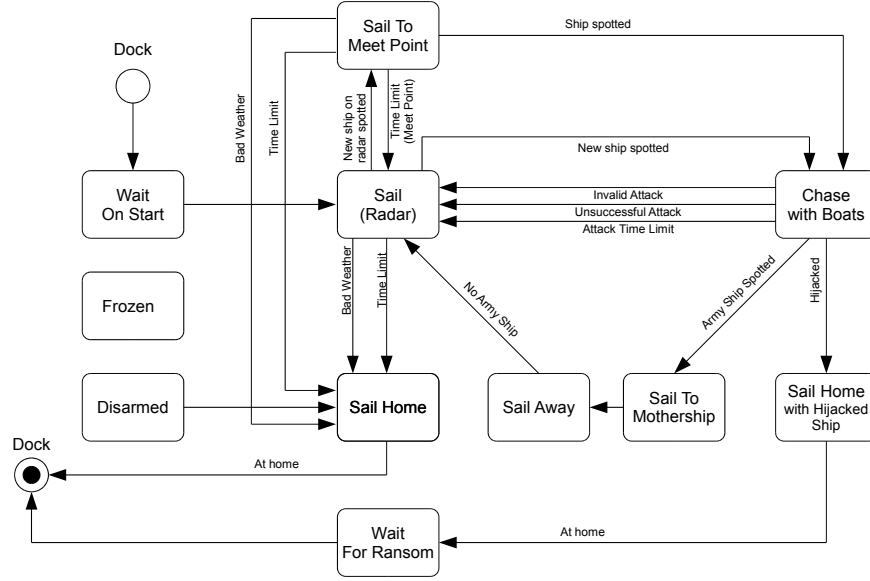


Fig. 3: FSM implementing the behavior of a pirate with a mothership and a radar.

and state implementations can be easily reused. It is possible to create abstract skeletons of various FSM and then enrich them with a specific behavior. The ability to learn can be implemented using internal state variables.

The problem of incorporating time into FSMs is solved by granting the agent each turn a time quantum which is used for performing the activity associated with each state. Nevertheless, one notable disadvantages of the FSM-based approach is the inability to perform multiple concurrent actions.

Pirate FSM Example Figure 3 depicts a FSM of a pirate equipped with a radar and a mothership. The pirate waits an arbitrary amount of time in its home port and then sails onto the open sea, scanning with the radar for a potential target. If it detects a cargo vessel, it approaches it and launches an attack. If the attack is successful and it hijacks the vessel, it gains full control over the vessel and forces it to sail to its home port and wait for the payment of ransom. This cycle can be disrupted by several external events, such as the proximity of a patrolling vessel or depletion of the pirate’s resources. The pirate can be also detained and disarmed by a naval patrol.

5 Model Validation

To validate the propose model, we compare simulation output with real-world data. We first validate the model of long-range shipping (represented by cargo vessels) alone and then look at the full model. Because the real-world data are

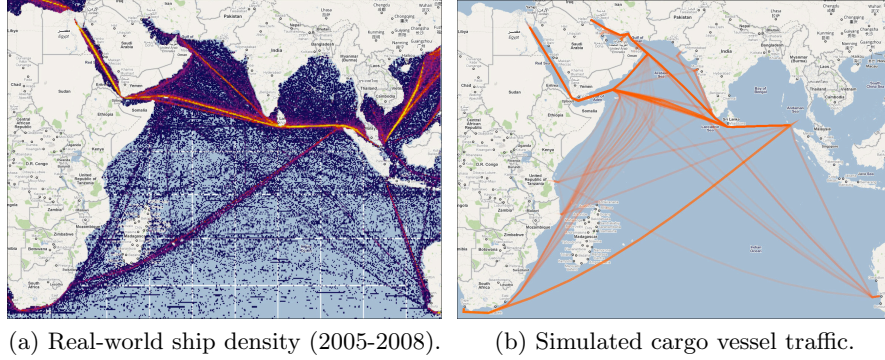


Fig. 4: Comparison of the real-world shipping density (left) and simulated cargo shipping traffic (right).

not sampled from any distribution and are result of the interaction of thousands of independent actors, the most straightforward way of comparison is visual inspection.

5.1 Long-range Shipping Traffic Model Validation

To verify the accuracy of long-range shipping model, we compare the traces of simulated cargo vessels with the aggregated AIS-based maritime shipping density map [7] (gathered from 2005 till the beginning of 2008, see Figure 4a). A few differences are visible: (1) Main corridors are narrower in the simulation. This difference could be removed by adding perturbations to simulated vessel routes. (2) The simulated traces do not exactly correspond to the corridors near the Socotra island. This difference is caused by the introduction of the IRTC corridor which is not yet reflected in the density map (our traces are more up-to-date). (3) The routes along the Somali coast extend farther from the coast in the simulated data. Again, the routing in these waters underwent major changes in the last years and the density map does not yet fully reflect the tendency of vessels to stay farther from the dangerous Somali coast. Our model uses the risk-based planner which takes this factor into consideration.

The reference data [7] also contain samples of trajectories of vessels not considered in our simulation, such as fishing vessels, maritime-research vessels or private vessels. These samples account for the irregular traffic outside the main corridors.

Overall the agreement between the real and simulated traffic is very good. In the future, we aim to introduce and evaluate a formal measure of model accuracy.

5.2 Piracy Model Validation

Due to the lack of real-world pirate movement data, we are not able to directly validate pirate behavior models. Instead we compare the pirate attack density, as

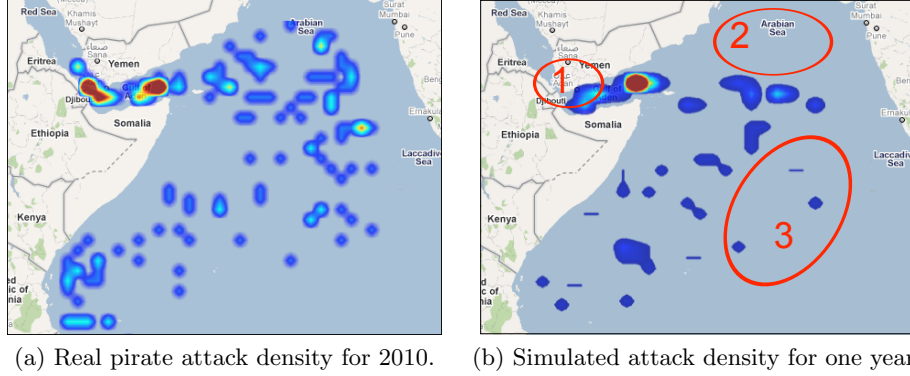


Fig. 5: Comparison of real and simulated attack densities for a 1-year period.

produced by the pirate models in conjunction with models of cargo and navy vessels, and compare this density with the data provided by IMB Piracy Reporting Centre (see Section 3.2).

For comparison, we have simulated one year of cargo vessel traffic through the region with realistic traffic density (approximately 60 vessels a day in IRTC). We do not have any estimates of real-world density of pirates in the area. We have therefore set their number so that the overall number of incidents corresponds to the number observed in the real-world. We have simulated 15 pirates – 5 simple, 5 radar equipped pirates and 5 pirates with a mothership and a radar. Finally, we have placed 3 naval task forces, each equipped with 2 warships and a helicopter, into the IRTC according to the 4W grid.

Figure 5 shows the comparison of the real and simulated density of pirate attacks. The red circles denote main differences in incident distribution: (1) Simulated pirates do not attack vessels in the Red Sea; this is because the current configuration focuses on the IRTC corridor. (2) There are no attacks in the Arabian Sea in the simulation. This is because the simulated pirates which sail far from the coast are equipped with a radar, and are thus almost always able to detect and attack vessels sailing in the shipping lane from the Gulf of Aden to Malaysia. (3) The density of attacks in the central Indian Ocean is lower; this is because of the lower density of the simulated cargo vessel traffic itself.

Overall, the agreement between the real-world and simulated pirate attack density is satisfactory, especially considering the random fluctuations in the real-world attack density and the fact that the simulated attack density is the result of the interaction of three types of vessels: cargo vessels, pirates and patrols. Each deviation in the behavior model is amplified through these interactions and can have a disproportional effect on the resulting incident density. After fixing the specific issues mentioned, the agreement between the model and reality should be further improved.

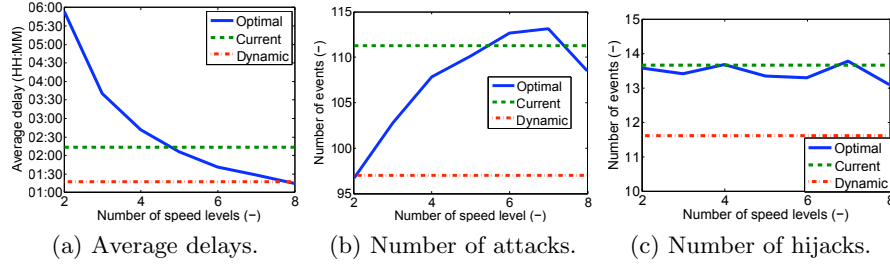


Fig. 6: Evaluation of different group transit schemes.

6 Optimization of Group Transit Schedules

As an example application of the developed model, we show how it can be used to optimize the Gulf of Aden group transit scheme (described in Section 2.2). The current transit scheme is based on fixed transit schedules designed to split transiting vessels into five groups travelling at different speeds, uniformly spaced between 10 and 18 knots. Our motivation is to explore whether the number of attacks and the transit delay caused by participating in group transit can be reduced by (1) proposing a different set of speed levels for the fixed-schedule scheme and (2) employing a dynamic grouping scheme, which takes into account not only the speeds of transiting vessels but their time of arrival too.

Optimal Fixed Schedule Group Transit The real-world distribution of cargo vessel speeds is approximated by a histogram (see Figure 7) with bins of a fixed width (can be set e.g. to 0.1, 0.5 or 1 knot). The optimization problem can be formulated as partitioning the histogram into N groups, corresponding to group transit speeds, which minimize the average transit delay (i.e. the delay caused by traveling at a group speed which might be lower than the vessel's normal cruising speed).

Dynamic Group Transit Instead of assigning vessels to groups according to predefined schedules and speed levels, dynamic grouping forms groups on the fly. This allows to form groups that better reflect actual arrival times and speed distribution of incoming vessels, at the expense of a more complex coordination scheme required. Our current implementation uses a greedy technique which assigns incoming vessels with similar speeds to the same group (see [10] for details).

Evaluation We have used the simulation to evaluate the average transit delay (Figure 6a), the average number of transit groups and the number of total and successful pirate attacks for different grouping schemes. Figure 7 shows the optimal speed levels for the fixed-schedule group transit with 5 speed levels ($\{10, 12.2, 14, 15.4, 17\}$ knots). Note that in contrast to the current scheme, new

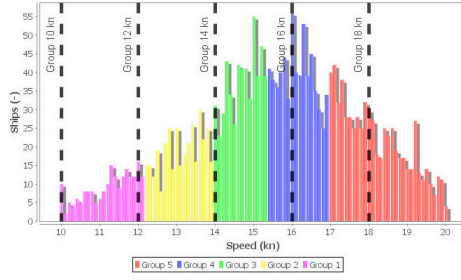


Fig. 7: Vessel speed histogram and its optimum partitioning into 5 speed groups.

scheme	ΔT /ship	ΔT /year	#groups
current	0 min	0 days	10.73
opt. 5	7 min	97 days	9.16
opt. 6	31 min	430 days	9.88
dynamic	56 min	778 days	12.14

Fig. 8: Time savings and average number of transit groups for selected grouping schemes. ΔT is the average transit time reduction compared to the current scheme (per one transit and aggregated for all transits in one year).

speed levels are not distributed uniformly and are concentrated around the mean vessel speed. As shown in Table on Figure 8, new speed levels result in shorter transit times and lower average number of transit groups.

The grouping for 6 speed levels ($\{10, 11.7, 13.3, 14.6, 16, 17.6\}$ knots) further reduces group transit delay and saves over a year of total vessel travel time. The number of pirate incidents does not significantly vary with the increasing number of speed levels (see Figures 6b, 6c), because even though there are potentially more groups to attack, the speed of the groups is higher on average. The dynamic scheme is by far the most efficient in reducing transit delay although at the expense of higher number of groups (i.e. smaller average group size).

7 Conclusion

Multi-agent simulations have a great potential for designing and evaluating solutions to a range of issues in international shipping, including the threat of contemporary maritime piracy. To this end, we have developed a first fully agent-based model of global traffic that accounts for the effects of maritime piracy on global shipping. The model is based on a range of real-world datasets and represents the operation of three types of vessels and their interactions. Despite the lack of hard data on some phenomena related to maritime piracy, the implemented model shows good correspondence in areas where validation data are available. As an example of potential applications of the developed model, we have showed how it can be used to optimize the current Gulf of Aden group transit scheme.

Overall, the implemented model demonstrates the viability of agent-based modeling in the maritime domain and opens a number of promising research directions. In the future, we would like to use the model to evaluate novel game-theoretic transit-patrol coordination and route planning techniques required to counter the recent expansion of piracy to the vast open areas of the Indian Ocean, where measures employed in the Gulf of Aden are no longer applicable.

Acknowledgments

Funded by the Office of Naval Research (grant no. N000140 910537) and by the Czech Ministry of Education, Youth and Sports under grant "Decision Making and Control for Manufacturing III" (grant no. MSM 6840770038).

References

1. K. Bernhardt. Agent-based modeling in transportation. *Artificial Intelligence in Transportation*, E-C113:72–80, 2007.
2. S. Bourdon, Y. Gauthier, and J. Greiss. MATRICS: A Maritime Traffic Simulation. Technical report, Defence R&D Canada, 2007.
3. B. Bošanský, V. Lisý, M. Jakob, and M. Pěchouček. Computing time-dependent policies for patrolling games with mobile targets. In *AAMAS*, 2011 (to appear).
4. R. CAPT George McGuire. Combined maritime forces (cmf) - who we are and wider military counter piracy update. online, 12 2009. http://www.cusnc.navy.mil/marlo/Events/DEC-MARLO-DubaiConference_files-2009/CAPT%20McGuire.ppt.
5. Expedition. Somalia pirate attacks. online, 4 2008. http://bbs.keyhole.com/ubb/ubbthreads.php?ubb=showflat&Number=1242871&site_id=1.
6. R. Gilpin. Counting the Costs of Somali Piracy. Technical report, US Institute of Peace, 2009.
7. B. Halpern, S. Walbridge, K. Selkoe, C. Kappel, F. Micheli, C. D'Agrosa, J. Bruno, K. Casey, C. Ebert, and H. Fox. A global map of human impact on marine ecosystems. *Science*, 319(5865):948, 2008.
8. K. Hasegawa, K. Hata, M. Shioji, K. Niwa, S. Mori, and H. Fukuda. Maritime Traffic Simulation in Congested Waterways and Its Applications. In *The 4th Conference for New Ship and Marine Technology (New S-Tech 2004)*, Chine, pages 195–199, 2004.
9. ICC Commercial Crime Services. Hostage-taking at sea rises to record levels, says IMB. online, 2010. <http://www.icc-ccs.org/news/429-hostage-taking-at-sea-rises-to-record-levels-says-imb>.
10. M. Jakob, O. Vaněk, B. Bošanský, O. Hrstka, and M. Pěchouček. Adversarial modeling and reasoning in the maritime domain year 2 report. Technical report, ATG, CTU, Prague, 2010.
11. D. Krajzewicz, G. Hertkorn, C. Ressel, and P. Wagner. Sumo (simulation of urban mobility) – an open-source traffic simulation. In *MESM2002*, 2002.
12. E. Kse, E. Basar, E. Demirci, A. Gneroglu, and S. Erkebay. Simulation of marine traffic in istanbul strait. *Simulation Modelling Practice and Theory*, 11(7-8):597–608, 2003.
13. T. Lozano-Perez and M. A. Wesley. An algorithm for planning collision-free paths among polyhedral obstacles. In *Communication of the ACM*, 1979.
14. MSCHOA. Gulf of aden internationally recommended transit corridor and group transit explanation. online, August 2010. http://www.shipping.nato.int/GroupTrans1/file/_WFS/GroupTransit12aug2010.pdf.
15. M. Pěchouček and D. Šišlák. Agent-based approach to free-flight planning, control, and simulation. *IEEE Intelligent Systems*, 24(1):14–17, Jan./Feb. 2009.
16. O. Vaněk, M. Jakob, V. Lisý, B. Bošanský, and M. Pěchouček. Iterative game-theoretic route selection for hostile area transit and patrolling. In *AAMAS*, 2011.

How to do social simulation in logic: modelling the segregation game in a dynamic logic of assignments

Benoit Gaudou, Andreas Herzig, Emiliano Lorini, and Christophe Sibertin-Blanc

Institut de Recherche en Informatique de Toulouse (IRIT), CNRS, France

Abstract. The aim of this paper is to show how to do social simulation in logic. In order to meet this objective we present a dynamic logic with assignments, tests, sequential and nondeterministic composition, and bounded and non-bounded iteration. We show that our logic allows to represent and reason about a paradigmatic example of social simulation: Schelling’s segregation game. We also build a bridge between social simulation and planning. In particular, we show that the problem of checking whether a given property P (such as segregation) will emerge after n simulation moves is nothing but the planning problem with horizon n , which is widely studied in AI: the problem of verifying whether there exists a plan of length at most n ensuring that a given goal will be achieved.

1 Introduction

In a recent debate Edmonds [7] attacked what he saw as “empty formal logic papers without any results” that are proposed in the field of multi-agent systems (MASs). He opposed them to papers describing social simulations: according to Edmonds the latter present many experimental results which are useful to better understand social phenomena, while the former kind of papers only aim at studying some relevant concepts and their mathematical properties (axiomatization, decidability, etc.), while not adding anything new to our understanding of social phenomena. In response to Edmonds’s attack, some researchers defended the use of logic in MAS in general, and in particular in *agent-based social simulation* (ABSS) [8, 4, 5]. For example, in [4] it is argued that logic is relevant for MAS because it can be used to construct a much needed formal social theory. In [8] it is argued that logic can be useful in ABSS because a logical model based on (a) a philosophical or sociological theory, (b) observations and data about a particular social phenomenon, and (c) intuitions —or a blend of them— can be considered to provide the requirements and the specification for an ABSS system and more generally a MAS. Moreover, a logical system might help to check the validity of the ABSS model and to adjust it by way of having a clear understanding of the formal model underpinning it. All these researchers consider logic and ABSS not only as compatible, but also as complementary methodologies.

The idea we defend in this paper is much more radical than those of the above advocates of logic-based approaches. Our aim is to show that ABSS can be directly done in logic and that a logical specification of a given social phenomenon can be conceived as an ABSS model of this phenomenon. We believe that the use of adequate

theorem provers will allow to obtain results that are beyond the possibilities of existing simulators. As a first step towards our aim we present in this paper a simple logic called $DL^\leftarrow$ (Dynamic Logic of Assignments). $DL^\leftarrow$ is an extension of propositional logic with dynamic operators. These operators allow to reason about *assignments* $p \leftarrow \top$ and $p \leftarrow \perp$ changing the truth value of a propositional variable p to ‘true’ or ‘false’ and about *tests* $\Phi?$ of the truth of a Boolean formula Φ . More generally, $DL^\leftarrow$ allows to reason about those facts that will be true after complex events ϵ that are built from assignments and tests by means of the operators of sequential composition ($\epsilon_1; \epsilon_2$), nondeterministic composition ($\epsilon_1 \cup \epsilon_2$), bounded iteration ($\epsilon^{\leq n}$), and unbounded iteration ($\epsilon^{<\infty}$).

In order to illustrate the power of our logic we show that a paradigmatic ABSS model can be represented in our logic: Schelling’s segregation game [15]. The problem of checking whether (under some initial conditions) a given property P such as segregation will *possibly emerge* after n simulation moves is reduced to the problem of checking in our logic whether the initial conditions imply that formula φ encoding property P will be true at the end of at least one sequence of events ϵ of length at most n . Similarly, the problem of checking whether P will *necessarily emerge* after n simulation moves is reduced to the problem of checking whether the initial conditions imply that φ will be true at the end of every sequence of events ϵ of length n . Actually the latter is nothing but the planning problem with horizon n , which is widely studied in AI and is for example at the base of the state of the art planner **SatPlan** [13]: the problem of verifying whether there exists a plan of length at most n ensuring that a given goal φ will be achieved. In the general case this problem is known to be in PSPACE, i.e. decidable in polynomial space [3]. We show that our logic fits these boundaries. In the past such PSPACE hard decision problems were considered to be out of reach of automated theorem provers. However, in the last 20 years huge progress was made on that kind of problems: state-of-the-art theorem provers for PSPACE complete problems were shown to be of practical interest in particular in semantic web applications even for realistic problem instances with thousands of clauses [12].

One might wish to go beyond the simple existential and universal quantifications that we mentioned above. This can be achieved by means of modal operators with counting (stemming from graded modal logics [9, 19] and description logics [1]). We are going to briefly discuss this, and show that complexity stays in PSPACE for such extensions.

The rest of the paper is organized as follows. In Section 2 we describe the segregation game. In Section 3 we define our basic logic, and in Section 4 we show how it allows to reason about the segregation game. Finally we sketch some extensions (Section 5) and conclude (Section 6).

2 The segregation game

In this section we give an informal description of the segregation game. A formal description is in Section 4.

The original model. Thomas C. Schelling in [15] studied the phenomenon of segregation and in particular the conditions of its occurrence due to “discriminatory individual choices” in groups with recognizable distinctions such as sex, age, colour, etc. The

best-known example is the formation of color-dependent residential areas, under the influence of the individual preference of being surrounded by at most a threshold number of neighbours with different colour: above the threshold inhabitants are unhappy and will move to another location.

One of the main results of Schelling's work is to show that the segregation phenomenon emerges even with a quite high tolerance threshold. For example, even if each inhabitant accepts that the majority of the neighbours surrounding him has a colour different from his, there will nevertheless be a tendency to form groups of inhabitants with the same colour.

The implemented model. The segregation model has been implemented in many languages and formalisms, in particular in almost all agent-based simulation platforms. Good examples are NetLogo [22, 21] and GAMA [17]. Every inhabitant (or family) is represented by an agent from a finite set of agents $\mathbb{A} = \{1, \dots, |\mathbb{A}|\}$. Typical elements of $\mathbb{A}$ are noted i, j , etc. Each of these agents can move on a chessboard-like grid with $N \times N$ locations, for some integer N such that $|\mathbb{A}| < N^2$. Each agent is characterized by his colour (e.g. red or blue) and his location on the grid. The latter is described by a couple of integers $(k, l) \in [1..N] \times [1..N]$, for $1 \leq k, l \leq N$. The two main parameters of the simulation are:

- the number of inhabitants $|\mathbb{A}|$ and
- the tolerance threshold: the number of different agents from which on an agent is unhappy, which is supposed to be the same for every agent.

(Alternatively the parameters may be the density of inhabitants and the percentage of different agents. We also note that most simulation models rather use the inverse of the tolerance threshold, called the similarity threshold.)

There is a scheduler which at each simulation step generates a random ordering of the set of agents and then activates the agents according to that ordering during the step. Upon activation an agent checks his happiness: an agent is happy iff the percentage of different neighbour agents (having a different colour) is below his tolerance threshold. If the agent is unhappy then he moves to another free cell on the grid.

The simulation stops when a stable state is reached, i.e. when every agent is happy.

A simulation typically has three different outcomes: (1) the simulation loops because the system does not reach a stable state where every agent is happy (typically when the density of agents on the grid is high and the tolerance threshold is low, which means agents are very intolerant); (2) the simulation stops but we cannot observe any kind of segregation (this is typically the case when the similarity threshold is very low, i.e. tolerance is high and/or density is very low); or (3) clusters of inhabitants with the same colour emerge.

3 Dynamic logic with assignments

This section introduces the syntax and the semantics of the logic $DL^{\leftarrow}$.

3.1 Language

We suppose given a countable set of propositional variables $\mathbb{P}$ with typical elements $p, q, \dots$. Remember that the set of Boolean formulas is built from $\mathbb{P}$ by means of the Boolean operators of negation and disjunction. (The other connectives are defined by means of abbreviations.) In our running example Boolean formulas allow to express things such as $\bigwedge_{(k,l) \neq (k',l')} \neg(\text{At}(i, k, l) \wedge \text{At}(i, k', l'))$, representing the fact that an agent i cannot be both at location (k, l) and at (k', l') .

The set of *events* $\mathcal{E}$ is defined by the following BNF:

$$\epsilon ::= p \leftarrow \top \mid p \leftarrow \perp \mid \Phi? \mid \epsilon; \epsilon \mid \epsilon \cup \epsilon \mid \epsilon^{\leq n} \mid \epsilon^{< \infty}$$

where p ranges over $\mathbb{P}$, Φ ranges over the set of Boolean formulas, and n ranges over the set of natural numbers $\mathbb{N}$.

The events $p \leftarrow \top$ and $p \leftarrow \perp$ are assignments modifying the truth value of the propositional variable p : the event $p \leftarrow \top$ sets p to true, and the event $p \leftarrow \perp$ sets p to false. $\Phi?$ is the test of Φ , which fails if Φ is false. $\epsilon_1; \epsilon_2$ denotes a sequence of events. $\epsilon^{\leq n}$ denotes iteration of ϵ exactly n times, $\epsilon^{< \infty}$ denotes iteration of ϵ up to n times, and $\epsilon^{< \infty}$ denotes arbitrary iteration of ϵ . Assignments and tests are atomic events, while the other events are called complex. An example of a complex event is agent i 's move from location (k, l) to location (k', l') , written $\text{At}(i, k, l) \leftarrow \perp; \text{At}(i, k', l') \leftarrow \top$.

The set of *formulas* $\mathcal{F}$ is defined by the following BNF:

$$\varphi ::= q \mid \top \mid \perp \mid \neg\varphi \mid \varphi \vee \varphi \mid \exists\epsilon.\varphi$$

where q ranges over $\mathbb{P}$ and ϵ ranges over the set of events $\mathcal{E}$. (Observe that we use Φ for Boolean formulas and φ for formulas of $\text{DL}^{\leftarrow}$.) The formula $\exists\epsilon.\varphi$ reads “there is an execution of the event ϵ after which φ is true”. Hence $\exists\epsilon.\top$ has to be read “ ϵ may occur”.

The operators ‘?’ , ‘;’ , ‘ $\cup$ ’ and ‘ $< \infty$ ’ are familiar from propositional dynamic logic PDL. (We could as well use the Kleene star and write ϵ^* instead of $\epsilon^{< \infty}$, as customary in PDL.) In uninterpreted PDL these operators combine abstract atomic programs, while in interpreted PDL they combine assignments of object variables to values from some domain. In contrast, atomic programs are here assignments of truth values to propositional variables, as previously studied in the dynamic epistemic logic literature [20, 6, 2].

The formula $\exists(q \leftarrow \top \cup q \leftarrow \perp).\varphi$ has the same interpretation as the quantified Boolean formula (QBF) $\exists q.\varphi$ [10]. The language of $\text{DL}^{\leftarrow}$ can therefore be viewed as a generalisation of quantification over Boolean variables to complex ‘quantification programs’.

The *length* of a formula φ , noted $|\varphi|$, is the number of symbols used to write down φ —without ‘ $\langle$ ’, ‘ $\rangle$ ’, dots, and parentheses—, where integers are supposed to have length 1.¹ For example $|\exists(q \leftarrow \top)^{\leq 3}.(q \vee r)| = 1 + |q \leftarrow \top| + 1 + 1 + |q \vee r| = 1 + 3 + 2 + 3 = 9$.

The logical operators $\wedge$ and $\rightarrow$ are defined as abbreviations; for example $\varphi \rightarrow \psi$ abbreviates $\neg\varphi \vee \psi$. Moreover, $\forall\epsilon.\varphi$ abbreviates $\neg\exists\epsilon.\neg\varphi$. Hence $\forall\epsilon.\perp$ has to be read “ ϵ cannot occur”.

¹ Precisely, the length of an integer n should be $\log n$. Our hypothesis is however without harm.

Remark 1. Note that could define programs ϵ^k as abbreviations of the sequential composition of ϵ , k times,² and then define $\epsilon^{\leq n}$ as abbreviations of $\bigcup_{0 \leq k \leq n} \epsilon^k$. However, the expansion of these abbreviations would exponentially increase formula length. For the same reason we avoided to introduce ' $\leftrightarrow$ ' as an abbreviation.

We use τ as a placeholder for either $\top$ or $\perp$, and write $q \leftarrow \tau$ in order to talk about $q \leftarrow \top$ and $q \leftarrow \perp$ in an economic way.

3.2 Semantics

A *valuation* is nothing but a model of classical propositional logic, viz. a subset of the set of propositional variables $\mathbb{P}$. The truth conditions are the usual ones for $\top$, $\perp$, negation and disjunction, plus:

$$\begin{aligned} V \models p & \quad \text{iff } p \in V \\ V \models \exists \epsilon. \varphi & \quad \text{iff there is } V' \text{ such that } VR_\epsilon V' \text{ and } V' \models \varphi \end{aligned}$$

where R_ϵ is a binary relation on valuations that is defined by:

$$\begin{aligned} R_{p \leftarrow \top} &= \{(V, V') : V' = V \cup \{p\}\} \\ R_{p \leftarrow \perp} &= \{(V, V') : V' = V \setminus \{p\}\} \\ R_{\epsilon_1; \epsilon_2} &= R_{\epsilon_1} \circ R_{\epsilon_2} \\ R_{\epsilon_1 \cup \epsilon_2} &= R_{\epsilon_1} \cup R_{\epsilon_2} \\ R_{\varphi?} &= \{(V, V) : V \models \varphi\} \\ R_{\epsilon^n} &= (R_\epsilon)^n \\ R_{\epsilon^{\leq n}} &= \bigcup_{0 \leq m \leq n} (R_\epsilon)^m \\ R_{\epsilon^{< \infty}} &= \bigcup_{0 \leq m} (R_\epsilon)^m \end{aligned}$$

The valuations V' such that $(V, V') \in R_\epsilon$ are called the *possible updates of V by ϵ* .

Validity and satisfiability are defined as usual.

3.3 Reduction axioms for the star-free fragment

The fragment of DL^+ without arbitrary iterations (called the 'star-free fragment' in PDL) can be axiomatised by means of reduction axioms. These axioms allow to eliminate all the dynamic operators from formulas. However, that elimination might result in an exponential blowup due to the operator of nondeterministic composition $\cup$. We will therefore later on characterize the complexity of validity checking by other means.

Proposition 1. *The following equivalences are DL^+ valid.*

² The precise definition is inductive: $\epsilon^0 = \text{skip}$, and $\epsilon^{k+1} = \epsilon^k; \epsilon$.

$$\begin{aligned}
\exists \epsilon^{=n}. \varphi &\leftrightarrow \begin{cases} \varphi & \text{if } n = 0 \\ \exists \epsilon^{=n-1}. \exists \epsilon. \varphi & \text{if } n > 0 \end{cases} \\
\exists \epsilon^{\leq n}. \varphi &\leftrightarrow \begin{cases} \varphi & \text{if } n = 0 \\ \exists \epsilon^{\leq n-1}. (\varphi \vee \exists \epsilon. \varphi) & \text{if } n > 0 \end{cases} \\
\exists \varphi?. \psi &\leftrightarrow \varphi \wedge \psi \\
\exists \epsilon_1 \cup \epsilon_2. \varphi &\leftrightarrow \exists \epsilon_1. \varphi \vee \exists \epsilon_2. \varphi \\
\exists \epsilon_1; \epsilon_2. \varphi &\leftrightarrow \exists \epsilon_1. \exists \epsilon_2. \varphi \\
\exists p \leftarrow \tau. \neg \varphi &\leftrightarrow \neg \exists p \leftarrow \tau. \varphi \\
\exists p \leftarrow \tau. (\varphi_1 \vee \varphi_2) &\leftrightarrow \exists p \leftarrow \tau. \varphi_1 \vee \exists p \leftarrow \tau. \varphi_2 \\
\exists p \leftarrow \tau. \top &\leftrightarrow \top \\
\exists p \leftarrow \tau. \perp &\leftrightarrow \perp \\
\exists p \leftarrow \tau. q &\leftrightarrow \begin{cases} \tau & \text{if } q = p \\ q & \text{if } q \neq p \end{cases}
\end{aligned}$$

These equivalences provide a complete set of reduction axioms for $^{\infty}$ -free dynamic operators $\exists \epsilon$. Call red the mapping on $\text{DL}^{\leftarrow}$ formulas which iteratively applies the above equivalences from the left to the right, starting from one of the innermost modal operators. It allows to first eliminate complex events, then push the dynamic operators inside the formula, and finally eliminate them when facing an atomic formula.

Proposition 2. *Let φ be a formula in the language of $\text{DL}^{\leftarrow}$ without the arbitrary iteration operator $\epsilon^{<\infty}$. Then*

1. $\text{red}(\varphi)$ has no modal operators;
2. $\text{red}(\varphi) \leftrightarrow \varphi$ is $\text{DL}^{\leftarrow}$ valid;
3. $\text{red}(\varphi)$ is $\text{DL}^{\leftarrow}$ valid iff $\text{red}(\varphi)$ is valid in classical propositional logic.

Proposition 2 indicates a close relationship between star-free $\text{DL}^{\leftarrow}$ and propositional logic. The merit of former over the latter is to provide a model theory and an intuitive and more succinct language. Both are valuable from a modelling perspective.

3.4 Complexity

Proposition 2 tells us that $\text{DL}^{\leftarrow}$ is not more expressive than classical propositional logic. However, reduction may exponentially increase the length of the formula because of the event operators $\cup$ and $\epsilon^{\leq n}$. But this is a suboptimal procedure, as we shall see now.

We first give the complexity result for model checking. The inputs of the model checking problem are a valuation V and a formula φ , and the problem is to decide whether $V \models \varphi$.

Theorem 1. *The problem of $\text{DL}^{\leftarrow}$ model checking is PSPACE-complete.*

PROOF. We first establish *hardness* by reducing the problem of validity of QBFs to $\text{DL}^{\leftarrow}$ model checking. Consider a fully quantified Boolean formula

$$\Phi = \exists q_1 \forall q_2 \exists q_3 \dots \exists q_{m-1} \forall q_m. \varphi$$

where $m \geq 0$ is even and where $\varphi(q_1, \dots, q_m)$ is a propositional formula containing no variables other than $q_1, \dots, q_m$. (The hypothesis that the number of quantifiers is even is without loss of generality: it suffices to add a dummy variable q_m not occurring in φ .) We define

$$\Phi^{\text{DL}^-} = \exists \epsilon_1. \forall \epsilon_2. \dots \exists \epsilon_{m-1}. \forall \epsilon_m. \varphi$$

where $\epsilon_k = q_k \leftarrow \top \cup q_k \leftarrow \perp$, for $1 \leq k \leq m$. Consider the valuation over the set $\mathbb{P} = \{q_1, \dots, q_m\}$ of propositional variables such that, say, $V(q_i) = \text{ff}$ for all q_i . It is readily checked that Φ is valid in Quantified Boolean Logic iff Φ^{DL^-} is true in V . Since both the size of Φ^{DL^-} and the size of the model are linear in the size of Φ , we conclude that DL^- model checking is PSPACE-hard.

Proof of *membership* is similar to that for DCL-PC [11]. It requires a recursive definition of the set of sequences of atomic events *admitted* by a complex event ϵ .

$$\begin{aligned} \text{adm}(p \leftarrow \tau) &= \{p \leftarrow \tau\} \\ \text{adm}(\Phi?) &= \{\Phi?\} \\ \text{adm}(\epsilon_1; \epsilon_2) &= \{\alpha_1; \alpha_2 : \alpha_1 \in \text{adm}(\epsilon_1) \text{ and } \alpha_2 \in \text{adm}(\epsilon_2)\} \\ \text{adm}(\epsilon_1 \cup \epsilon_2) &= \text{adm}(\epsilon_1) \cup \text{adm}(\epsilon_2) \\ \text{adm}(\epsilon^{=0}) &= \{\top?\} \\ \text{adm}(\epsilon^{=n+1}) &= \{\alpha_1; \alpha_2 : \alpha_1 \in \text{adm}(\epsilon^{=n}) \text{ and } \alpha_2 \in \text{adm}(\epsilon)\} \\ \text{adm}(\epsilon^{\leq n}) &= \bigcup_{m \leq n} \text{adm}(\epsilon^{=m}) \\ \text{adm}(\epsilon^{<\infty}) &= \bigcup_n \text{adm}(\epsilon^{\leq n}) \end{aligned}$$

(Clearly, the set $\text{adm}(\epsilon)$ is infinite iff ϵ contains $<\infty$.) The main point in the proof is that every possible update of a valuation V by a complex event ϵ can also be reached by a sequence of atomic events that is admitted by ϵ and that is of length exponential in the length of ϵ . This is the case despite the presence of the arbitrary iteration operator $<\infty$, the reason being that any $\epsilon^{<\infty}$ can only bring about a finite number of valuation updates. Based on that one can prove that membership of a couple of valuations (V, V') in R_ϵ can be decided in polynomial space. Finally one can prove that a formula φ can be evaluated in space polynomial in the size of φ ; in particular, when evaluating $\exists \epsilon^{<\infty}. \varphi$ one may decide in polynomial space whether one of the (exponentially many) valuations V' is accessible from V . ■

Theorem 2. *The problem of DL^- validity checking is PSPACE-complete.*

PROOF. Hardness can be proved by translating QBF formulas to DL^- in the same way as in Theorem 1.

Membership can be proved for the DL^- satisfiability checking problem as follows: given φ we guess a model V . (V can be supposed to be of polynomial size because we may restrict our attention to the propositional variables occurring in φ , and neglect those that don't.) Then we check whether $V \models \varphi$, which can be done in polynomial space by mirroring the truth conditions. This shows that DL^- satisfiability can be checked in NPSPACE. Now by Savitch's theorem $\text{NPSPACE} = \text{PSPACE}$, and therefore DL^- satisfiability can be checked in polynomial space. It follows that the complementary DL^- validity problem can be checked in polynomial space, too. ■

The above result shows that $DL^\leftarrow$ is more succinct than propositional logic: there are $DL^\leftarrow$ formulas (and even star-free $DL^\leftarrow$ formulas) such that every equivalent propositional formula is exponential longer.

4 The segregation game in $DL^\leftarrow$

In Section 2 we introduced Schelling's segregation game in an informal way. Let us now model it in $DL^\leftarrow$.

4.1 Propositional variables

We need three kinds of propositional variables:

$$\begin{aligned} \text{At}(i, k, l) & \text{ “agent } i \text{ is at location } (k, l) \text{”} \\ \text{R}(i) & \text{ “agent } i \text{ is red”} \\ \text{Done}(i) & \text{ “it was already } i \text{'s turn in the present step”} \end{aligned}$$

where i is an agent in $\mathbb{A} = \{1, \dots, |\mathbb{A}|\}$ and every (k, l) is a location in $[1..N] \times [1..N]$. The following abbreviations will be useful:

$$\begin{aligned} \text{B}(i) & \stackrel{\text{def}}{=} \neg \text{R}(i) \\ \text{NB}^{<1}(k, l) & \stackrel{\text{def}}{=} \bigwedge_i \bigwedge_{k', l' \leq N : |k' - k|, |l' - l| \leq 1} \neg (\text{At}(i, k', l') \wedge \text{B}(i)) \\ \text{NR}^{<1}(k, l) & \stackrel{\text{def}}{=} \bigwedge_i \bigwedge_{k', l' \leq N : |k' - k|, |l' - l| \leq 1} \neg (\text{At}(i, k', l') \wedge \text{R}(i)) \\ \text{Hpy}^{<1}(i, k, l) & \stackrel{\text{def}}{=} (\text{R}(i) \wedge \text{NB}^{<1}(k, l)) \vee (\text{B}(i) \wedge \text{NR}^{<1}(k, l)) \\ \text{Free}(k, l) & \stackrel{\text{def}}{=} \bigwedge_i \neg \text{At}(i, k, l) \\ \text{Segreg}^{<1} & \stackrel{\text{def}}{=} \bigwedge_i \bigvee_{k, l} (\text{At}(i, k, l) \wedge \text{Hpy}^{<1}(i, k, l)) \end{aligned}$$

$\text{NB}^{<1}(k, l)$ can be read “location (k, l) has no blue neighbour” and similarly for $\text{NR}^{<1}(k, l)$. Moreover, i is happy at a given location (k, l) —noted $\text{Hpy}^{<1}(i, k, l)$ — if and only if “at (k, l) , agent i has no neighbour with a different colour”. $\text{Free}(k, l)$ can be read “location (k, l) is free”. Finally, the property of segregation holds —noted $\text{Segreg}^{<1}$ — if and only if “every agent is happy at his place”.

Let us compute the length of these formulas. The length of both $\text{NB}^{<1}(k, l)$ and $\text{NR}^{<1}(k, l)$ is in $O(|\mathbb{A}|)$ (the quantification over the locations (k', l') being void because there are exactly 9 such locations including (k, l) and his neighbours), and is therefore in $O(N^2)$ (because $|\mathbb{A}| < N^2$). The same holds for the length of $\text{Hpy}^{<1}(i, k, l)$. Finally, the length of $\text{Segreg}^{<1}$ is in $O(|\mathbb{A}| \times N^2 \times N^2)$, i.e. it is in $O(N^6)$.

4.2 Describing the agents' moves

Every move of agent i from location (k, l) to location (k', l') can be described by a complex event in our language of events.

$$\text{move}(i, k, l, k', l') \stackrel{\text{def}}{=} \text{At}(i, k, l)? ; \text{Free}(k', l')? ; \text{At}(i, k, l) \leftarrow \perp ; \text{At}(i, k', l') \leftarrow \top$$

Then a move of the simulation is described by a nondeterministic composition of agent moves (plus some turn taking management):

$$\text{move} \stackrel{\text{def}}{=} \bigcup_{i,k,l} \left(\neg \text{Hpy}^{<1}(i, k, l)?; \neg \text{Done}(i)?; \bigcup_{k',l'} \text{move}(i, k, l, k', l'); \text{Done}(i) \leftarrow \top \right); \\ (\top? \cup \neg \bigvee_{i,k,l} (\neg \text{Hpy}^{<1}(i, k, l) \wedge \neg \text{Done}(i))?; \text{Done}(1) \leftarrow \perp; \dots; \text{Done}(|\mathbb{A}|) \leftarrow \perp)$$

Nondeterministic choice corresponds to the scheduler in simulation model implementations. A simulation step consists in $|\mathbb{A}|$ executions of **move**.

Just as $\text{Hpy}^{<1}(i, k, l)$, the length of $\text{move}(i, k, l, k', l')$ is in $O(N^2)$. The length of **move** is in $O(|\mathbb{A}| \times N^4 \times N^2) = O(N^8)$.

4.3 Domain and initial state

In order to properly model the segregation game, we must describe the *domain laws* relating the propositional variables $\text{At}(i, k, l)$, $\text{R}(i)$ and $\text{Done}(i)$. They say that there cannot be a location inhabited by two different agents (χ_1), that an agent cannot be at two different places (χ_2), and that an agent is at least at one place (χ_3):

$$\begin{aligned} \chi_1 &= \bigwedge_{i_1 \neq i_2, k, l} \neg(\text{At}(i_1, k, l) \wedge \text{At}(i_2, k, l)) \\ \chi_2 &= \bigwedge_{i, (k, l) \neq (k', l')} \neg(\text{At}(i, k, l) \wedge \text{At}(i, k', l')) \\ \chi_3 &= \bigwedge_i \bigvee_{k, l} \text{At}(i, k, l) \end{aligned}$$

(In order not to overload notation we left implicit that $1 \leq k, l \leq N$.) Together these formulas make up the domain laws:

$$\text{Laws} = \chi_1 \wedge \chi_2 \wedge \chi_3$$

The length of **Laws** is determined by that of χ_1 , which is in $O(|\mathbb{A}|^2 \times N^2) = O(N^6)$.

The *initial state* is described by

$$\text{Init} = (\bigwedge_{i \in J} \text{R}(i)) \wedge (\bigwedge_{i \notin J} \neg \text{R}(i)) \wedge (\bigwedge_{i \in \mathbb{A}} \neg \text{Done}(i) \wedge \text{At}(i, k_i, l_i))$$

where $J \subseteq \mathbb{A}$ is the set of red agents and (k_i, l_i) is the initial location of agent i . Its length is in $O(|\mathbb{A}|) = O(N^2)$.

4.4 Describing and proving properties

In our language we can express things such as “segregation *will always* occur after n moves” (φ_1), “segregation *may* occur within n moves” (φ_2), “when segregation occurs then none of the agents will move” (φ_3), etc.:

$$\begin{aligned} \varphi_1 &= \forall \text{move}^{=n}. \text{Segreg}^{<1} \\ \varphi_2 &= \exists \text{move}^{\leq n}. \text{Segreg}^{<1} \\ \varphi_3 &= \text{Segreg}^{<1} \rightarrow \forall \text{move}. \perp \end{aligned}$$

Other kinds of properties will be discussed in Section 5.

Given a property described by a formula φ , what we are interested in is to check whether the formula

$$(\text{Laws} \wedge \text{Init}) \rightarrow \varphi$$

is $\text{DL}^\leftarrow$ valid, where φ is one of the above properties.

The difference between `Init` and `Laws` is that while `Init` has just to be true in the initial state, the laws must be true in any update of the current state. In order to ensure that our modelling works properly the first thing to do is to check that the domain laws are preserved by any sequence of events from `move`. To prove this it suffices to prove that

$$\text{Laws} \rightarrow \forall \text{move}. \text{Laws}$$

is $\text{DL}^\leftarrow$ valid.³

It is important to note that the lengths of the formulas `Laws`, $\text{Segreg}^{<1}$, $\text{Hpy}^{<1}(i, k, l)$, `Init` and of the complex event `move` are polynomial in the domain size parameter N . Therefore the length of the formulas $\text{Laws} \rightarrow \forall \text{move}. \text{Laws}$ and $(\text{Laws} \wedge \text{Init}) \rightarrow \varphi_k$ is polynomial in N (precisely, their length is in $O(N^8)$). The validity problem in $\text{DL}^\leftarrow$ being in PSPACE we obtain the following results.

Proposition 3. *The validity of $\text{Laws} \rightarrow \forall \text{move}. \text{Laws}$ and of $(\text{Laws} \wedge \text{Init}) \rightarrow \varphi_k$, for $\varphi_k \in \{\varphi_1, \varphi_2, \varphi_3\}$, can be checked in space polynomial in N .*

All our decision problems being in PSPACE, we can envisage to use existing theorem provers for PSPACE problems to check the above properties, such as provers for modal logic K, for description logic ALC, or for Quantified Boolean Formulas. This requires a polynomial transformation of the formulas to be checked into the language of the respective logic.⁴

4.5 Varying the agents' tolerance

The agents modelled here are extremely intolerant. More tolerant agents can be described as follows:

$$\begin{aligned} \text{NB}^{<2}(k, l) &\stackrel{\text{def}}{=} \bigwedge_{i_1, i_2 : i_1 \neq i_2} \bigwedge_{(k_1, l_1), (k_2, l_2) : |k_1 - k|, |l_1 - l|, |k_2 - k|, |l_2 - l| \leq 1} \\ &\quad \neg(\text{At}(i_1, k_1, l_1) \wedge \text{At}(i_2, k_2, l_2) \wedge \text{B}(i_1) \wedge \text{B}(i_2)) \\ \text{NB}^{<p}(k, l) &\stackrel{\text{def}}{=} \bigwedge_{i_1, \dots, i_p : i_m \neq i_n \text{ if } m \neq n} \bigwedge_{(k_1, l_1), \dots, (k_p, l_p) : |k_m - k|, |l_m - l| \leq 1} \\ &\quad \neg(\text{At}(i_1, k_1, l_1) \wedge \dots \wedge \text{At}(i_p, k_p, l_p) \wedge \text{B}(i_1) \wedge \dots \wedge \text{B}(i_p)) \end{aligned}$$

$\text{NB}^{<p}(k, l)$ is read “location (k, l) has less than p blue neighbours”. $\text{NR}^{<p}(k, l)$ is defined accordingly. $\text{Hpy}^{<1}(i, k, l)$ and $\text{Segreg}^{<1}$ can be generalised to $\text{Hpy}^{<p}(i, k, l)$ and $\text{Segreg}^{<p}$ in the obvious way. One may also stipulate that an agent is happy if the percentage of agents in his neighbourhood with a colour different from his is below some threshold.

The length of $\text{NB}^{<2}$ is in $O(|\mathbb{A}|^2)$ and is therefore in $O(N^4)$. It follows that the length of $\text{Hpy}^{<2}(i, k, l)$ is also in $O(N^4)$, that of $\text{Segreg}^{<2}$ is in $O(N^8)$, and that of `move` is in $O(N^{10})$. Generally, the length of $\text{NB}^{<p}(k, l)$ is in $O(|\mathbb{A}|^p)$; the parameter p being at most 8, it follows that the length of $\text{NB}^{<p}(k, l)$ is in $O(N^{16})$, that of $\text{Hpy}^{<p}(i, k, l)$ is in $O(N^{16})$, too, that of $\text{Segreg}^{<p}$ is in $O(N^{20})$, and that of `move` is in $O(N^{22})$. Overall, such properties can still be checked in polynomial space just as those of Section 4.4.

³ From this it follows by standard principles of modal logic that $\text{Laws} \rightarrow \forall \text{move}. \text{Laws}$ and $\text{Laws} \rightarrow \forall \text{move}^{\leq n}. \text{Laws}$ are $\text{DL}^\leftarrow$ valid. By the induction axioms of PDL one can also prove that $\text{Laws} \rightarrow \forall \text{move}^{<\infty}. \text{Laws}$ is $\text{DL}^\leftarrow$ valid.

⁴ While we know that such a transformation exists because all these problems are in the same complexity class, it remains to find an elegant such transformation.

5 More expressive languages

We now discuss some generalisations of our logic that allow to naturally express other properties one would like to check in simulations.

Let us introduce two new modal operators $\geq k \epsilon$ and $\geq \frac{1}{2} \epsilon$, where $\geq k \epsilon.\varphi$ reads “ φ is true in at least k of the possible updates by ϵ ” and $\geq \frac{1}{2} \epsilon.\varphi$ reads “ φ is true in most of the states after ϵ ”. So the formula $\exists \epsilon.\varphi$ of Section 3 is nothing but $\geq 1 \epsilon.\varphi$.

$$\begin{aligned} V \models \geq k \epsilon.\varphi & \text{ iff } |\{V' : (V, V') \in R_\epsilon \text{ and } V' \models \varphi\}| \geq k \\ V \models \geq \frac{1}{2} \epsilon.\varphi & \text{ iff } |\{V' : VR_\epsilon V' \text{ \& } V' \models \varphi\}| > |\{V' : VR_\epsilon V' \text{ \& } V' \not\models \varphi\}| \end{aligned}$$

This allows to formulate interesting properties of segregation game such as the following.

- “segregation will occur within n moves *at least* k times” (ψ_1);
- “segregation will occur at some point in the future *at least* k times” (ψ_2);
- “segregation will occur within n moves *exactly* k times” (ψ_3);
- “segregation will occur at some point in the future *exactly* k times” (ψ_4);
- “segregation will occur after n moves in most of the cases” (ψ_5).

$$\begin{aligned} \psi_1 &= \geq k \text{ move}^{\leq n}. \text{Segreg}^{< p} \\ \psi_2 &= \geq k \text{ move}^{< \infty}. \text{Segreg}^{< p} \\ \psi_3 &= \geq k \text{ move}^{\leq n}. \text{Segreg}^{< p} \wedge \neg \geq k+1 \text{ move}^{\leq n}. \text{Segreg}^{< p} \\ \psi_4 &= \geq k \text{ move}^{< \infty}. \text{Segreg}^{< p} \wedge \neg \geq k+1 \text{ move}^{< \infty}. \text{Segreg}^{< p} \\ \psi_5 &= > \frac{1}{2} \text{ move}^{\leq n}. \text{Segreg}^{< p} \end{aligned}$$

Model checking requires some more bookkeeping in order to count valuations, but can still be done in polynomial space. We obtain a PSPACE completeness result for the validity checking problem by using Savitch’s theorem in the same way as we did in the proof of Theorem 2.

6 Conclusion

In this paper we have shown how to do social simulation in a dynamic logic with assignments, tests, sequential and nondeterministic composition, and bounded and non-bounded iteration.

Instead of our logic we could also have used other logical approaches to reasoning about actions such as the Situation Calculus [14], the Fluent Calculus [18], or the Event Calculus [16]. However, while these formalisms allow to represent more or less the same things, their mathematical analysis is less developed: while there are some decidability results, there are no complexity results that could be compared to the PSPACE completeness result for our logic.

As we have said in the introduction the kind of properties we want to prove can be viewed as planning problems with a finite horizon. We could therefore have used existing finite horizon planners in order to prove properties of simulations. It has to be noted that planners typically build plans, while we are only interested in proving plan existence. However, it is an interesting research avenue to exploit a possible convergence of the fields of simulation and planning.

References

1. F. Baader, D. Calvanese, D.L. McGuinness, D. Nardi, and P.F. Patel-Schneider, editors. *Description Logic Handbook*. Cambridge University Press, 2003.
2. J. van Benthem, J. van Eijck, and B. Kooi. Logics of communication and change. *Information and Computation*, 204(204):1620–1662, 2006.
3. T. Bylander. The computational complexity of propositional strips planning. *Artificial Intelligence*, 69:165–204, 1994.
4. R. Conte and M. Paolucci. Responsibility for societies of agents. *J. of Artificial Societies and Social Simulation*, 7(4), 2004.
5. F. Dignum, B. Edmonds, and L. Sonenberg. Editorial: The use of logic in agent-based social simulation. *J. of Artificial Societies and Social Simulation*, 7(4), 2004.
6. H. P. van Ditmarsch, W. van der Hoek, and B. Kooi. Dynamic epistemic logic with assignment. In *Proc. of AAMAS’05*, pages 141–148. ACM Press, 2005.
7. B. Edmonds. How formal logic can fail to be useful for modelling or designing MAS. In G. Lindeman, editor, *Proc. of the International Workshop on Regulated Agent-Based Social Systems: Theories and Applications (RASTA’02)*, LNAI, pages 1–15. Springer-Verlag, 2004.
8. M. Fasli. Formal systems and agent-based social simulation equals null? *J. of Artificial Societies and Social Simulation*, 7(4), 2004.
9. M. Fattorosi-Barnaba and F. de Caro. Graded modalities I. *Studia Logica*, 44:197–221, 1985.
10. M. R. Garey and D. S. Johnson. *Computers and Intractability: A Guide to the Theory of NP-Completeness*. W. H. Freeman Co., 1979.
11. W. van der Hoek, D. Walther, and M. Wooldridge. On the logic of cooperation and the transfer of control. *JAIR*, 37:437–477, 2010.
12. I. Horrocks. Using an expressive description logic: Fact or fiction? In *Proc. of KR’98*, pages 636–649, 1998.
13. H. A. Kautz and B. Selman. Planning as satisfiability. In *Proc. of ECAI’92*, pages 359–363, 1992.
14. R. Reiter. *Knowledge in Action: Logical Foundations for Specifying and Implementing Dynamical Systems*. MIT Press, 2001.
15. T. C. Schelling. Dynamic Models of Segregation. *Journal of Mathematical Sociology*, 1:143–186, 1971.
16. M. Shanahan. *Solving the frame problem: a mathematical investigation of the common sense law of inertia*. MIT Press, 1997.
17. P. Taillandier, A. Drogoul, D.A. Vo, and E. Amouroux. GAMA: a simulation platform that integrates geographical information data, agent-based modeling and multi-scale control. In *Proc. of PRIMA’10*, 2010.
18. M. Thielscher. The logic of dynamic systems. In *Proc. 14th Int. Joint Conf. on Artificial Intelligence (IJCAI’95)*, pages 1956–1962, Montreal, Canada, 1995.
19. W. van der Hoek. On the semantics of graded modalities. *J. of Applied Non-Classical Logics*, 2(1), 1992.
20. Jan van Eijck. Making things happen. *Studia Logica*, 66(1):41–58, 2000.
21. U. Wilensky. Netlogo segregation model. Technical report, Center for Connected Learning and Computer-Based Modeling, Northwestern University, Evanston, IL, 1997.
22. U. Wilensky. Netlogo. Technical report, Center for Connected Learning and Computer-Based Modeling, Northwestern University, Evanston, IL, 1999.

Simulating Research Behaviour

Nardine Osman, Jordi Sabater-Mir, and Carles Sierra

Artificial Intelligence Research Institute (IIIA-CSIC), Barcelona, Catalonia, Spain

Abstract. This paper proposes a simulation for research behaviour, focusing on the process of writing papers, submitting them to journals and conferences, reviewing them, and accepting/rejecting them. The simulation is currently used to evaluate the OpinioNet reputation model, which calculates the reputation of researchers and research work based on inferred opinions. The goal is to verify whether the reputation model succeeds in encouraging ‘good’ research behaviour or no, although the simulator is elaborate enough to be used for the analysis of other aspects of paper writing, submission, and review processes.

1 Introduction

The classical way in which scientific publications are produced, evaluated and credited, is being challenged by the use of modern computer science technologies [1, 5]. In particular, software versioning tools and reputation mechanisms make it realistic to think about a publication process where the publications are ‘liquid’, in the sense that they are persistently accessible over Internet and modified along time. Credit is then given to authors based on opinions, reviews, comments, etc. This would produce many beneficial results, for instance, to reduce the current large number of very similar publications (i.e. salami papers) or to organise conferences by just searching for the most prestigious liquid publications satisfying certain keywords. OpinioNet [4] is one reputation model that has been proposed with the intent of encouraging ‘good’ research behaviour. For instance, its equations are designed to give more attention to the quality of a researcher’s work than their quantity.

This paper proposes a simulator that would simulate the course of scientific publications, from writing papers and submitting them to journals and conferences to reviewing them and accepting/rejecting them. The main goal of the simulator is to verify whether the OpinioNet reputation model succeeds in encouraging ‘good’ research behaviour or not, although the simulator is rich enough to be used for the analysis of other aspects of paper writing, submission, and review processes.

The rest of this paper is divided as follows. Section 2 provides the background needed for understanding the followed formal publications model and the OpinioNet reputation model being evaluated. Section 3 introduces the basics of the simulator. The simulator’s algorithm is then presented by Section 4, and its assumptions and hypothesis are presented by Section 5. Section 6 then discusses the preliminary results, before concluding with Section 7.

2 Background

Reputation is widely understood as the group’s opinion about the entity in question. The OpinioNet reputation model [4] is based on the concept that in the case of the lack of explicit opinions, opinions may be deduced from related entities. For instance, in the field of publications, one may assume that if a paper has been accepted by a reputable conference, then the paper should be of a minimum quality. Similarly, a conference becomes reputable if it accepts high quality papers. OpinioNet is essentially based on this notion of opinion propagation in structural graphs, such as the publications graphs (where nodes of this graph may represent knowledge objects, such as conference proceedings and conference papers, and relations stating which is part of what).

OpinioNet understands opinions as probability distributions over an evaluation space, for a particular attribute, and at a moment in time; for example, one can define a set of elements for the evaluation space for *quality* of a node as $\{\text{poor}, \text{good}, \text{v.good}, \text{excellent}\}$. The set of attributes that opinions may address can be, for instance, $\{\text{novelty}, \text{clarity}, \text{significance}, \text{correctness}\}$. OpinioNet then defines the structural graph accordingly: $SG = \langle N, G, O, E, A, T, \mathcal{E}, \mathcal{F} \rangle$, where N is the set of nodes, G is the set of agents that may generate opinions about nodes, O is the set of all opinions, E is the ordered evaluation space for O , A is the set of attributes that opinions may address, T represents calendar time, $\mathcal{E} \subseteq N \times N$ specifies which nodes are part of the structure of which others (i.e. $(n, n') \in \mathcal{E}$ implies n is part of n'), $\mathcal{F} : G \times N \times A \times T \rightarrow O$ is a relation that links a given agent, node, attribute, and time to their corresponding opinion.

A single opinion is then represented as the probability distribution $\mathbb{P}(E|G, N, A, T) \in O$. We note that probability distributions subsume classical approaches and are more informative. Hence, the adoption of this approach by our simulator does not necessarily restrict its application to other scenarios.

3 Simulation Basics

In theory, it is researchers’ behaviour (defined through their profiles) that would influence the creation and evolution of papers, journals, and, eventually, fields of research. However, to keep the simulation simple, our example focuses on the evolution of one specific aspect of a research community—namely, the growth of the community’s contributions—and neglects other aspects that are not deemed crucial for the evaluation of the reputation module—such as the rise and fall of the community itself, its journals, its fields of research, etc. As such, and for the sake of simplifying the simulation, we choose to simulate a single community with a fixed number of researchers researching a given subject; we say it is the researchers’ profiles that control the production and dissemination of single contributions; and we keep the number of journals that could accept/reject these contributions fixed. In other words, we say one ‘top rated’ journal is sufficient to represent the acceptance of a contribution by any ‘top rated’ journal. We argue that since our current interest is in the future of authors’ contributions

(such as papers or book chapters), the number of journals becomes irrelevant: what is crucial is the quality of the journals (if any) that accept the authors' contributions. In what follows, we define the simulator's input and output.

3.1 The Simulator's Input and Output

The simulator requires the following tuple as input: $\langle SG^0, J, U, \mathcal{J}, \mathcal{U}, \mathbf{T} \rangle$, where SG^0 describes the initial state of the system (or the initial SG graph), which should include at least a fixed number of researchers and journals; J describes the set of journal profiles (defined shortly); U describes the set of researcher profiles (defined shortly); $\mathcal{J} \subseteq N \times J$ is a function that maps a journal in N (where N is the set of knowledge objects in SG^0) to a journal profile in J ; $\mathcal{U} \subseteq G \times U$ is a function that maps a researcher in G (where G is the set of researchers in SG^0) to a researcher profile in U ; and $\mathbf{T} \in \mathbb{N}^*$ describes the number of years to be simulated. Then, every simulation year Y results in a modified structural graph SG^Y . The evolution of the SG graph is then presented: $E_{SG} = \{SG^0, \dots, SG^{\mathbf{T}}\}$.

3.2 Journals' Profiles

Journals are categorised through profiles that define their quality and their required number of reviewers. A journal's profile $j \in J$ is defined as the tuple: $j = \langle \mathbb{J}, \mathbf{RN} \rangle$, where similar to the opinions on quality, $\mathbb{J}$ is a probability distribution over the evaluation space E describing the quality of the journal; and $\mathbf{RN}$ describes the number of reviewers needed to review a paper, and it is specified as a Gaussian function over the set of natural numbers $\mathbb{N}$.

The rules for accepting/rejecting contributions depends on the quality of the journal $\mathbb{J}$. For example, very good journals are very strict about the quality of the papers they accept, other lower quality ones are not as strict. Hence, a journal's acceptance threshold $\mathbf{AT}$ may be defined in terms of its quality $\mathbb{J}$. A preliminary definition could be to have $\mathbf{AT} = emd(\mathbb{J}, \mathbb{T})$, where emd is the earth movers distance that calculates the distance (whose range is $[0, 1]$, where 0 represents the minimum distance and 1 represents the maximum possible distance) between two probability distributions [6],¹ and $\mathbb{T} = \{e_n \mapsto 1\}$ (where $\forall e_i \in E \cdot e_n > e_i$) describes the ideal distribution, or the best distribution possible.

3.3 Researchers' Profiles

Similar to journals, researchers' behaviour is also categorised through profiles that define their quality, their productivity, etc. A researcher's profile $u \in U$ is defined as the tuple: $u = \langle \mathbb{Q}, \mathbf{RP}, \mathbf{CN}, \mathbb{C}, \mathbf{CA}, \mathbf{CP}, \mathbf{SC}, \mathbf{V}, \mathbf{SS}, \mathbf{RvP}, \mathbf{RV}, \mathbf{RT} \rangle$, where

¹ If probability distributions are viewed as piles of dirt, then the earth mover's distance measures the minimum *cost* for transforming one pile into the other. This cost is equivalent to the 'amount of dirt' times the distance by which it is moved, or the distance between elements of the ordered evaluation space E .

- **Q** describes the researcher’s research quality, and it is specified as a probability distribution over the evaluation space E (we assume that researchers have a fixed and ‘intrinsic’ quality of research — Section 5 argues the need for this intrinsic value — which is different from reputation values that reflect the view of the community and are calculated by reputation algorithms);
- **RP** describes the research productivity in terms of the produced number of papers per year (since produced research work is usually presented and preserved through papers, whether published or unpublished), and it is specified as a Gaussian function over the set of natural numbers $\mathbb{N}$;
- **CN** describes the researcher’s usual number of coauthors per contribution, and it is specified as a Gaussian function over the set of natural numbers $\mathbb{N}$;
- **C** describes the accepted research quality of coauthors, and it is specified as a probability distribution over the evaluation space E ;
- **CA** describes the accepted affinity level of coauthors (currently, affinity measure describes how close are two researchers’ profiles; however, in future simulations, one may also consider affinity measures that describe how close are two researchers with respect to numerous social relations), and the range of its value is the interval $[0, 1]$, where the value 0 represents minimum affinity and the value 1 represents maximum affinity;
- **CP** describes the level of persistency in sticking with old coauthors, it is defined in terms of the number of past papers that two researchers have coauthored together, and it is specified as a Gaussian function over the set of natural numbers $\mathbb{N}$;
- **SS** describes the submission strategy of the researcher, and the range of its value is the interval $[-1, +1]$, where a value -1 represents an extreme ‘risk-averse’ strategy in which the researcher does not submit a paper to any journal unless its paper is of the highest quality possible, a value of $+1$ represents an extreme ‘risk-seeking’ strategy in which the researcher doesn’t mind submitting a paper to a journal of much higher quality, and the value 0 represents a more neutral approach in which the researcher usually submits its papers to journals of the same quality (of course, values in between represent different levels of risk-averse and risk-seeking strategies);²
- **RvP** describes the researcher’s review productivity in terms of the number of papers the researcher accepts to review per year, and it is specified as a Gaussian function over the set of natural numbers $\mathbb{N}$;
- **RV** describes the review quality in terms of how close the researcher’s reviews are from the true quality of the papers in question, it is defined in terms of the distance from the true quality of the paper in question, and the range of this distance is the interval $[-1, +1]$; and
- **RT** describes the reviewers’ threshold for accepting to review a paper for a given journal, it is defined in terms of the earth mover’s distance between

² We assume that the quality of journals and that of researchers may be compared since they are measured on the same scale. Our assumption is based on the idea that the quality of the researchers, their papers, and the journals that accept those papers are all based on the quality of the research work being carried out and presented.

the reviewer’s quality of research and that of the journal’s, and the range of its value is the interval $[0, 1]$.

We note that although the researchers’ profiles may seem too complex, (1) many ideas have already been overly simplified, as illustrated by Section 5.1, and (2) additional simplifications are straightforward.

4 Simulation Algorithm

While the previous section has introduced the simulator as a black box, this section presents an overview of the simulation algorithm. (For the simulator’s technical details, we refer the interested reader to our technical document [3].) The algorithm’s steps are outlined below.

1. *Generate the groups of coauthors for the given year*

The idea is that each group of coauthors will produce one paper that will then be added to the *SG* graph. In summary, the algorithm selects the authors one by one, giving the authors that intend to write more papers this year (specified by the research productivity **RP** of the researcher) a higher probability of being selected first. Then, for each selected author, the algorithm searches, in an iterated manner, for a suitable group of coauthors, where the ‘suitability’ is based on the restrictions imposed by each researcher through its preferred number of coauthors (**CN**), the accepted quality of coauthors (**C**), the accepted affinity of coauthors (**CA**), and the accepted persistency of coauthors (**CP**). The algorithm iterates until all researchers are assigned to as many coalitions as needed.

2. Then, for each paper resulting from a created group of coauthors, the simulator performs the following:

- (a) *Calculate the intrinsic quality of the paper*

We base our simulation on the idea that papers have a true quality that researchers (and reviewers) often try to guess. Of course, in reality, this value does not exist. However, simulation may assume such values to compare and analyse the performance of researchers. We assume that when researchers from various qualities (where a researcher’s true quality is specified by the parameter **Q**) are grouped together then the resulting paper’s true value would be an aggregation of the researchers’ true value.

- (b) *Choose the journal to submit the paper to*

After a paper is created, it is submitted to some journal. The selection of the journal assumes that researchers tend to have certain submission strategies (specified through the parameter **SS**), and the submission strategy for a given journal is an aggregation of its authors’.³ The final

³ Although, in fields that are known to have an enormous number of authors per paper (for example, it is not uncommon for Physics articles to have a few thousand authors each) and by the law of large numbers, the simple aggregation method of the authors’ **SS** values would fail since it would result in similar values for all papers. In such cases, it might be useful to calculate the resulting submission strategy by adopting it from the leading author.

calculated submission strategy is then used to help select the journal to submit to by enforcing constraints on the distance between the paper’s true quality (calculated by step 2(a) above) and that of the journal’s ($\mathbb{J}$).

- (c) *Choose the reviewers to review the paper*

This action is based on the number of reviews needed $\mathbf{RN}$, the availability of the reviewers (researchers are assumed to review a certain number of papers per year, determined by $\mathbf{RvP}$, and they accept the journals’ requests to review papers based on a first come first serve basis), the quality of the researcher $\mathbb{Q}$, the quality of the journal $\mathbb{J}$, and the reviewer’s threshold for accepting a journal ($\mathbf{RT}$).

- (d) *Generate the reviewers’ opinions (reviews) about the paper in question*

Reviewers’ opinions are based on the intrinsic quality of the paper (calculated by step 2.(a).) and the researcher’s review quality ($\mathbf{RV}$), which determines how close the review would be to the paper’s true value.

- (e) *Accept/Reject the paper by the chosen journal*

This calculation is based on the quality of the journal $\mathbb{J}$, the journal’s acceptance threshold ($\mathbf{AT}$), and the reviewers’ aggregated opinions, where the aggregation take into consideration the reviewers’ reputation at the time. Note that if the paper is accepted, then it is linked (in the *SG* graph) to the journal through the *part of* relation.

- (f) *Reputation measures are calculated by OpinioNet*

After reviews are created and papers are accepted/rejected accordingly, the simulator calls the OpinioNet reputation model to calculate the reputation of papers based on the new reviews and acceptance results, as well as the authors’ reputation based on the reputation of their papers.

3. Repeat the entire process for the following year.

5 Assumptions and Hypotheses

After introducing the proposed simulation algorithm, and before moving on to the experiments and results, this section is intended to clarify our stance by highlighting and discussing the assumptions we make as well as clarifying the claims the simulation algorithm is designed to test.

5.1 Assumptions

As discussed earlier, trying to simulate the real behaviour of researchers requires a thorough study of various aspects, from how people choose their coauthors and how they choose where to submit their work to, to how do journals select reviewers, and how is the quality of a paper related to the research quality of its authors. We argue that the proposed simulation algorithm is sophisticated enough to capture the actions that have an impact on reputation measures, yet it is simplified enough to overlook unnecessary complicated behaviour. As a result, a bunch of assumptions are made, which we discuss below. We note that many of the fixed values that we refer to in our assumptions are in fact either drawn

from a predefined Gaussian function, or some noise is added to them to make our scenarios more realistic.

On the Static Nature of the Research Community. We say both the community and the researchers' behaviour are static: researchers do not join or leave the community; journals do not evolve or die; the field of research is fixed; each paper cites a fixed number of other papers; a researcher's productivity does not change with time; a researcher's quality of research does not evolve with time; a researcher's review productivity does not evolve with time; a researcher's review quality does not evolve with time and his reviews always fall at a fixed distance from that of the true quality of the paper being reviewed; a researcher's submission strategy does not evolve with time; a researcher's acceptable journal quality for reviewing papers is fixed; and journals do not evolve and they always accept papers of the same quality.

These assumptions are introduced to keep the simulation simple. Although, to keep the simulation more realistic, some randomness is introduced when generating the measures specifying a 'fixed' behaviour. We postpone the study of dynamic and evolving behaviour for future work.

On Selecting Coauthors. Selecting the coauthors to collaborate with is usually a complex matter that depends on a variety of issues, such as the subject of study, the practicality of collaboration, and so on. Our proposed simulator, however, is not aimed at studying the dynamics of human relations, their collaboration, and coalition formations, but the production of papers on an annual basis. Hence, for simplicity, the production of papers assumes researchers produce a fixed number of papers per year, and coauthor their papers with a group of other researchers. Again, for simplicity, we assume the strategy of selecting coauthors is fixed and that coauthors are selected based on their quality, their affinity, and their persistence. Of course, different weights may be given to each.

On the True Quality of Researchers and Research Work. Although in reality the true quality of a researcher is neither definite nor accessible, we do assume that researchers have defined true and fixed qualities, since we feel the need to base several other actions on these qualities. For example, we say the true quality of a paper is based on the true quality of its authors. And this is crucial because we say, for instance, that how good a reviewer is depends on how close its opinion is from the true quality of the paper in question.

But how is the true quality of the paper calculated? We argue that when researchers from various qualities are grouped together then the resulting paper's true (quality) value would be an aggregation of the researchers' true (quality) value. However, we also assume that when the dispersion in quality is large (where the Gini coefficient [2] is used to describe dispersion), the aggregation tends to follow a mean nature; and when the dispersion in quality is low, the aggregation tends to follow a more superadditive nature. The equation we propose for calculating a paper's true value $\mathbb{X}$ is then defined as follows:

$$\mathbb{X} = G \cdot \text{mean}(\{\mathbb{Q}_a\}) + (1 - G) \cdot \text{superadd}(\{\mathbb{Q}_a\}) \quad (1)$$

where, $\{Q_a\}$ represents the set of all authors' true value Q , G represents the Gini coefficient, *mean* represents some mean function, and *superadd* represents some superadditive function.⁴⁵

The superadditivity implies that when a group of researchers in the same quality range are combined, then they may be able to produce some work that is a little bit better than what each may produce on their own. However, when very good researchers coauthor papers with very poor ones, then the poor ones could possibly pull the quality of the produced paper below the excellent researcher's standard quality. In other words, the effect that a single researcher has on the true quality of a paper is dependent on its quality of coauthors.

On Selecting the Journal to Submit a paper to. We say researchers have different, and 'fixed', submission strategies. We define submission strategies following the prospect theory classification of strategies into risk seeking, risk averse, and risk neutral ones. The submission strategy of a paper is then an aggregation of its authors'. For example, if the submission strategy is 'risk seeking', then the paper may be submitted to some journal which is of better quality than the paper in question. If the submission strategy is 'risk averse', then the paper cannot be submitted to a journal unless it is of higher quality. Of course, varying levels of these strategies are considered.

On Selecting the Reviewers. Journals usually try to get good reviewers, based on availability. But how do reviewers choose whether to accept/reject reviewing papers for a given journal. We assume that reviewers accept journals based on a first come first serve basis, as long as the journal is of acceptable quality and the reviewer is available to do more reviews.

On Accepting/Rejecting papers by a given Journal. We assume accepting/rejecting a paper is only based on reviewers' opinions, and not on the number of papers submitted, the acceptance rate, etc. This is necessary because we have already assumed the number of journals to be fixed. In other words, we say one journal of a given quality is enough to represent all potential journals of that same quality.

Furthermore, when accepting/rejecting a paper, the journal's editors do not base their decision on the true quality of the paper (since this information is not available), but on the aggregated reviewers' opinions. When aggregating reviewers' opinions, we say that the reliability of a review (or opinion) is based on the researcher's current reputation in his/her community (as calculated by the Opin-ioNet algorithm), rather than how confident it claims to be (which is how the current review process works). We believe this is a stronger reliability measure since the reviewer does not assess himself, but is assessed by the community.

⁴ We argue that it does not matter much which exact mean or superadditive functions we choose, since the effect of that would be minute. In any case, we hope future extensive simulations would clarify which choices are better.

⁵ It is not clear yet whether the proposed approach for calculating a paper's true values would provide better results than a simple variance of authors' quality values. Future extensive simulations could also clarify which choices are better.

On the Fate of Papers. Finally, we say that, for the sake of simplicity, both accepted and rejected papers are forgotten. Neither of them is submitted to other journals in the following years; only new papers are created each year. In practice these new papers would in fact be a modification of (i.e. a new version of) already existing ones. However, we currently postpone the simulation of the *version of* relation that links related papers for future simulations. This assumption is acceptable since the current simulation simply focuses on the number of journals and the quality of the journals that accept them, rather than the *evolution* of papers.

5.2 Hypotheses

The OpinioNet reputation model has been used in an attempt to encourage ‘good’ research behaviour [3]. The proposed simulator aims at verifying whether OpinioNet achieves its goal or not through testing the following hypotheses.

Hypothesis 1 *It is more profitable to produce few high quality papers than several lower quality ones.*

Hypothesis 2 *It is more profitable to follow a ‘risk-neutral’ submission strategy.*

The first hypothesis implies that it would be more profitable, in terms of reputation, to spend more time on producing few high quality research papers than numerous papers of lower quality. This, we believe, lowers the dissemination overhead in researchers’ contributions and encourages researchers to spend more time on high quality research, as opposed to wasting time on repackaging already existing ideas for the sole purpose of increasing reputation.

The second hypothesis implies that it is more profitable (again, in terms of reputation) to submit one’s contributions to journals that lie in the same quality range of the paper. For instance, if authors choose journals that are much better, then they end up wasting the community’s time and resources, and they also waste time before their work is accepted. However, if the authors choose journals that are of much lower quality than the work submitted, then they miss the chance of having this work published in more reputable journals. Naturally, proving/disproving this hypothesis will be influenced by the assumption that papers may only be submitted once. We believe resubmitting usually requires the creation of new versions, which we postpone for future (and more advanced) simulations. Nevertheless, the current simulation may illustrate the effort, time, and potential gain in reputation that could be wasted by preferring one submission strategy over another.

To test the claims above, we define two different simulations, one for each of these claims. In each of these simulations, we keep the value of the relevant parameters in the researchers’ profiles fixed while we vary the other values. For instance, the parameter describing the researcher’s productivity (**RP**) should be fixed when testing the first hypothesis, while the parameter describing the researcher’s submission strategy (**SS**) should be fixed for testing the second hypothesis. The following section discussed the details of our simulation examples and their results.

6 Results and Analysis

For our preliminary simulation, we choose to simulate a small research community composed of a fixed set of 20 researchers. The simulation then runs for 10 time-steps, where each time-step represents one calendar year. In other words, our simulated example represents a fixed community of 20 researchers with varying behaviour and its evolution over 10 years.

At each time-step, papers are added following the constraints of the various profiles. With the addition of each paper, the following measures are calculated: (1) the OpinioNet reputation of the papers affected by this addition, and (2) the OpinioNet reputation of authors affected by this addition. The evolution of these measures along time is then plotted for further analysis, as illustrated shortly.

Evaluating Hypothesis 1. In this experiment, we divide the researchers into 3 groups: (1) those with a high quality research and low productivity level, (2) those with a medium quality research and a medium productivity level, and (3) those with low quality research and a high productivity level. We note that the productivity level represents the number of papers produced per year. All the other values defining the researchers' profiles are kept fixed. For example, all researchers share the same criteria in selecting their co-authors.

The results are presented by Figure 1, which shows that those who focus on the quantity cannot do better than those who focus on the quality of their work. However, if we have two researchers that have the same quality of work, then it is not very clear whether focusing on quantity would help or not. For this reason, we run a second experiment, where all the researchers now share the same quality of work, but have different productivity levels. Four categories of productivity are distinguished: (1) those with very high productivity level per year, (2) those with an average productivity level per year, (3) those with a very low productivity level per year, and (4) those who produce a paper every several years (around 5 years on average). The results of this experiment are presented by Figure 2, which illustrates that as long as a researcher is producing papers relatively 'frequent' enough, then s/he needs not focus on the quantity of their papers. However, there is a limit for this 'frequency'. For example, we show that researchers who produce only one paper every 5 years cannot compete with those who are continuously active.

These results confirm hypothesis 1, which states that *it is more profitable to produce few high quality papers than several low quality ones*; yet, researchers are required to remain active not become dormant for long periods of time.

Evaluating Hypothesis 2 In this experiment, we divide the researchers into 3 groups: (1) those with a 'risk-seeking' submission strategy, (2) those with a 'risk-neutral' submission strategy, and (3) those with 'risk-averse' submission strategy. All the other values defining their profiles are kept fixed. For example, all researchers are of a medium-high quality; all researchers share the same criteria in selecting their co-authors, etc. The results of this experiment are presented by Figure 3, which shows that the researchers that are more picky in their

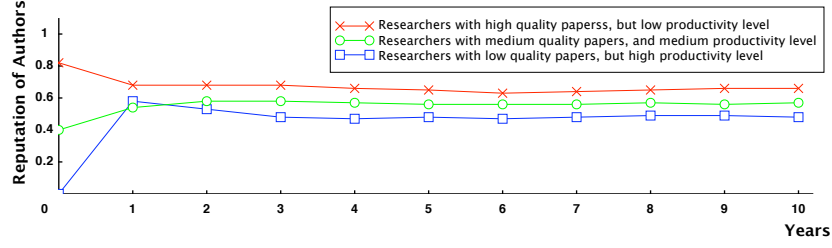


Fig. 1. Reputation of authors: first experiment results evaluating Hypothesis 1

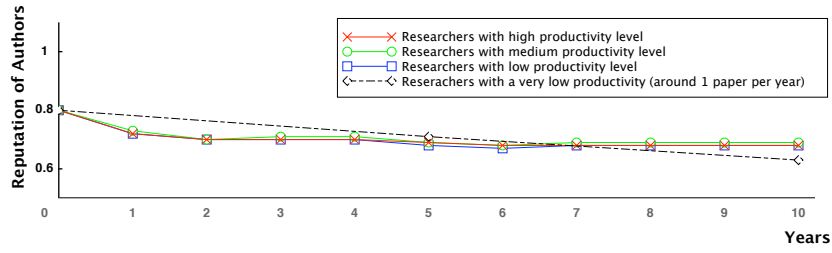


Fig. 2. Reputation of authors: second experiment results evaluating hypothesis 1

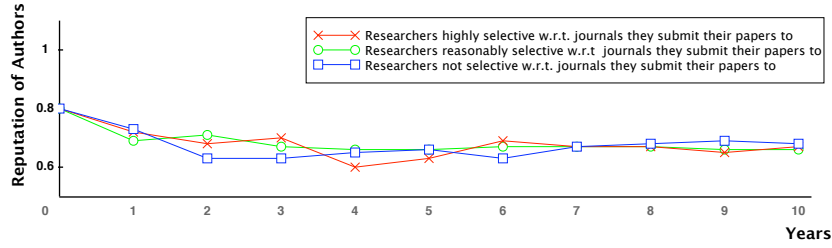


Fig. 3. Reputation of authors: experiment results evaluating hypothesis 2

selection of the journals to submit their papers to have their reputation fluctuating much more than others. This is because they risk their work being rejected by some journals, but when they do get published, their gain in reputation is relatively high.

An intriguing issue to note is that all three groups tend to have a similar reputation after a long period of time (10 years in this case). We believe that this might be related to the quality of the researchers. For example, researchers with a very high quality might be better off choosing a ‘risk-seeking’ submission strategy than a ‘risk-averse’ one. We hope future simulations to clarify this issue.

7 Conclusion

This paper has presented a simulator that simulates the details of the publication process, from writing papers and submitting them to journals, to the review process and the final decision of accepting/rejecting papers. The simulator is rich enough to be used in analysing different aspects of the publication process. However, this paper has focused on using it to verify the OpinioNet reputation model’s success in encouraging ‘good’ behaviour, such as encouraging the focus on the quality of work produced as opposed to its quantity. Future work should help us understand the details and preconditions of hypotheses 1 and 2 better. We also plan to extend our future simulations for testing additional hypotheses such as: (1) Is it more profitable to collaborate with researchers of high research quality? (2) Are more reputable researchers less susceptible to the selected quality of coauthors? (3) Is it more profitable to repackage one’s research work into different versions? And so on.

Acknowledgements

This work has been supported by the LiquidPublications project (funded by the European Commission’s FET programme) and the Agreement Technologies project (funded by CONSOLIDER CSD 2007-0022, INGENIO 2010).

References

1. Fabio Casati, Fausto Giunchiglia, and Maurizio Marchese. Publish and perish: why the current publication and review model is killing research and wasting your money. *Ubiquity*, 2007:3:1–3:1, January 2007.
2. Corrado Gini. *Variability and Mutability (Italian: Variabilità e mutabilità)*. Clarendon Press, 1912. Reprinted in *Memorie di metodologica statistica* (Ed. Pizetti E, Salvemini, T). Rome: Libreria Eredi Virgilio Veschi (1955).
3. Nardine Osman, Jordi Sabater-Mir, Carles Sierra, Adrián Perreau de Pinninck Bas, Muhammad Imran, Maurizio Marchese, and Azzurra Ragone. Credit attribution for liquid publications. Deliverable D4.1, LiquidPublications Project, June 2010.
4. Nardine Osman, Carles Sierra, and Jordi Sabater-Mir. Propagation of opinions in structural graphs. In Helder Coelho, Rudi Studer, and Michael Wooldridge, editors, *Proceeding of the 19th European Conference on Artificial Intelligence (ECAI’10)*, volume 215 of *Frontiers in Artificial Intelligence and Applications*, pages 595–600, Amsterdam, The Netherlands, The Netherlands, 2010. IOS Press.
5. Nardine Osman, Carles Sierra, Jordi Sabater-Mir, Joseph R. Wakeling, Judith Simon, Gloria Origgi, and Roberto Casati. LiquidPublications and its technical and legal challenges. In Danièle Bourcier, Pompeu Casanovas, Mélanie Dulong de Rosnay, and Catharina Maracke, editors, *Intelligent Multimedia: Managing Creative Works in a Digital World*, volume 8, pages 321–336. EPAP, Florence, 2010.
6. Shmuel Peleg, Michael Werman, and Hillel Rom. A unified approach to the change of resolution: Space and gray-level. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 11(7):739–742, 1989.

Agent Simulation of Peer Review: The PR-1 Model

Francisco Grimaldo¹, Mario Paolucci², Rosaria Conte²

¹ Departament d'Informàtica
Universitat de València
Av. Vicent Andrés Estellés, s/n, Burjassot, Spain, 46100
francisco.grimaldo@uv.es

² Italian National Research Council (CNR)
Institute of Cognitive Sciences and Technologies (ISTC)
Viale Marx 15, Roma, Italy, RM 00137
mario.paolucci@istc.cnr.it, rosaria.conte@istc.cnr.it

Abstract. Peer review lies at the core of current scientific research. It is composed of a set of social norms, practices and processes that connect the abstract scientific method with the society of people that apply the method. As a social construct, peer review should be understood by building theory-informed models and comparing them with data collection. Both these activities are evolving in the era of automated computation and communication: new modeling tools and large bodies of data become available to the interested scholar. In this paper, starting from abstract principles, we develop and present a model of the peer review process. We also propose a working implementation of a subset of the general model, developed with Jason, a framework that implements the Belief-Desire-Intention (BDI) model for multi agent systems. After running a set of simulations, varying the initial distribution of reviewer skill, we compare the aggregates that our simplified model produces with recent findings, showing how for some parameter choice the model can generate data in qualitative agreement with measures.

1 Introduction

Science is both a method - a logically coherent set of norms and processes - and a social activity, in which people and organizations endeavour to apply the method. One of the most important elements of the social structure of science is peer review, the process that scrutinizes scientific contributions before they are made available to the community.

As with any social process, peer review should be the object of scientific investigation, and should be evaluated with respect to a set of parameters. Common sense would suggest, at least, considerations of fairness and efficiency. In addition, two specific dimensions very relevant to research are innovation promotion

A preliminary version of this paper has been presented at EUMAS 2010 with the title “A proposal for agent simulation of peer review.”

and fraud detection. Science evolves by revolutions [6], and peer review should be evaluated with respect to its reaction to novelty. Is the current system of peer review supporting radical innovation, or is it impeding it?

Fraud detection, especially for politically relevant matters as medicine and health, is also extremely important; its actual effectiveness at ensuring quality has yet to be fully investigated. In [7], the review process is found to include a strong “lottery” component, independent of editor and referee integrity. While the multiple review approach to a decision between two options is supported by Condorcet’s jury theorem, if we move beyond simple accept/reject decisions by considering scoring and ranking, we find several kinds of potential failures that are not waived by the theorem.

These questions are particularly relevant right now, because, on the one hand, peer review is ready to take advantage of the new information publishing approach created by Web 2.0 and beyond. On the other hand, we perceive a diffuse dissatisfaction of scientists towards the current mechanisms of peer review. This is sometimes testified just anecdotally; list of famous papers that were initially rejected and striking fraudulent cases abound. Leaning on examples is an approach that we do not support because it is known to induce bias [12]. However, some recent papers have shown some numerical evidence on the failures of peer review [4].

In fact, peer review is just a specific case of mutual scoring. Following [8, 9], it is a reciprocal and symmetric type of evaluation which includes narrow access and transparency to the target (at least this is how it is designed in the case of teamwork, see the example of scientific research evaluation). Peer review is the standard that journals and granting agencies use to ensure the scientific quality of their publications and funded projects.

The question that follows is then - can we improve on this process? We are not going to fall for the technology trap, and just suggest that by updating peer review to the Web X.0 filtering, tagging, crowdsourcing, and reputation management practices [9], every problem will disappear - in fact, change could make the problems worse; consider for example the well known averaging effect of searching and crowd filtering [3].

Instead, we propose to create a model (or better, a plurality of models) of peer review, that takes into account recent theoretical developments in recommender systems and reputation theories, and test “in silico” the proposed innovations. In this work, we draw an overview of how we foresee such a model, and we present a first, partial implementation of it. The literature of simulation models about peer review is scarce; the only other approach we have found is [11], where the authors focus on an optimizing view of the reviewer for his or her own advantage.

The rest of the paper is organized as follows: the next section outlines a general model of peer review as well as a restricted model focusing on the roles of the reviewer and the conference. Section 3 explains how the latter has been implemented as a Multi-Agent System (MAS) over Jason [2]. In section 4 we show the aggregates that our simplified model produces when varying the distribution

of reviewers ability. Finally, in section 5 we state the conclusions of this work and discuss about future lines of research.

2 Description of the proposed model

In this section, we draw the outline of a model of peer review (PR-M in the following) and of its subset that we have implemented. We use agent-based simulation as a modelling technique [1]. With respect to statistical techniques employed for example in [4] or [7], the agent-based or individual-based approach allows us to model the process explicitly. In addition, it helps focusing on agents, their interaction, and possibly also their special roles - consider for example the proposal in [7] of increasing pre-screening of editors or editorial boards. Such a change is based on trust in the fair performance of a few individuals who take up the editors role. Thus, these individuals deserve detailed modeling, that could allow us to reason on their goals and motivations [5].

In this model, we want to catch the whole social process of review, and not just the workings of the single selection process. We will try to simulate the whole lifecycle of peer review, that will allow for example - in the complete model - to reason about role superposition between author and reviewer. This approach distinguishes our effort from that of other authors like [4].

2.1 PR-M

The key entities in our system are: the *paper*, as the basic unit of evaluation; the *author* of the paper; and the *reviewer*, which participate in a program committee of a specific *conference*. We define them in the following paragraphs.

Paper. Here, we do not focus on research but on its evaluation. We assume that the actual value of a paper - that we take as the basic research brick - is difficult to ascertain. Thus, while we give to each paper an actual value, we speculate that the value is only accessible through a procedure that implies the possibility of mistakes.

As a consequence, value is hidden by noise and evaluating papers is modeled as a difficult task - though, noise can obviously be canceled by repeated independent evaluations. In our model, we give papers an intrinsic fixed value. But there is another, different value that can be calculated and that changes in time: the number of citations that the paper receives.

The value of a paper as the number of its citation should, in an ideal case, reflect its actual value. In PR-M, we plan to implement a citation system so that approved papers can be cited by other papers, thus creating a network of citations. The decision process will be carried on by the simulated author. With both an intrinsic value and a citation count, after an initial bootstrapping phase, we could check the correlation between these two. The larger the correlation, the better the whole system of peer review is performing.

Author. Authors create papers and submit them to the conferences. With the citation network, the author will also decide on what papers are to be included in the bibliography. We plan to develop a probabilistic choice where a paper will have a higher chance to be cited depending on a list of factors including paper presence in a conference where the author is in the PC, or has submitted a paper; being co-authored by the author himself; and being a highly cited paper, thus mirroring the positive feedback mechanism that operates in research. Authors could have individual preferences on the weights. By varying the distribution of the intrinsic value of the papers submitted as well as the author preferences, the PR-M model will allow us to analyze the evolution of the quality of the papers published by each conference.

Reviewer. Reviewers can be part of the program committee (PC) of any number of conferences. In every simulation cycle, representing one year or conference edition, they evaluate a certain number of papers for each conference that enlists them in the PC.

The PR-M model characterizes reviewers by a probability value, named reviewer skill (s), that represents the chance they actually understand the paper they are reviewing.

The distribution of s values is the primary cause of reviewing noise. We will experiment with several distributions, including a uniform distribution of s values across reviewers (which we consider a low level of noise in evaluations) and other, left-skewed distributions where a low level of reviewing skill is more frequent.

Conference. As with the paper, we use the term conference in a general sense; it covers also, for example, the journal selection process. The authors' decision about what conference to send their works to is crucial, since the number of papers received is a measure of success for the conference, and their quality will determine the conference's quality. Can the review-conference system ensure quality in the face of very strong noise, variable reviewers skill, thanks to some selection process of PC composition that leans on the simplest measurable quantity - disagreement? The number of evaluations a paper receives are just a few - three being a typical case. Thus, the conference is where all the process comes together - are three reviews enough to cancel noise? For what distributions of papers and reviewers skill?

2.2 PR-1

In this paper, we only present a restricted implementation of the full model. This restricted model, that we call PR-1, contains a subset of the features in PR-M, focusing on the roles of the reviewer and the conference only. Thus, the authors and the papers are not included in the following PR-1 definition.

PR-1 represents the peer review problem by a tuple $\langle R, C \rangle$, where R is the set of *reviewers* participating in the PC of a set of *conferences* C .

Each reviewer $r \in R$ has an associated skill value $s \in [0, 1]$. The result of reviewing is accurate with probability s , and completely random with probability $(1 - s)$. To test different distributions in the unit segment, we use the beta distribution. Depending on its two parameters (see figure 1), this distribution can easily express very diverse shapes such as: a uniform skill distribution ($\alpha = 1, \beta = 1$); a set of moderately low skill reviewers ($\alpha = 2, \beta = 4$), and a mix of very good and very bad reviewers ($\alpha = 0.4, \beta = 0.4$).

Conferences $c \in C$ are represented by the tuple:

$$c = \langle np, rp, pr, R_c, pa, ac, I, d, e \rangle$$

Each conference receives a number of papers np every year, and employs a subset of reviewers $R_c \subseteq R$ to prepare rp reviews for each paper. The size of the PC ($|R_c|$) depends on the number of reviews done per PC member pr .

Papers have an associated value representing their intrinsic value, and receive a review value from each reviewer. The intrinsic values follow a uniform distribution over a N -values ordered scale, interpretable as the standard from strong reject to strong accept scores. Conferences accept the best pa papers whose average review value is greater than the acceptance value ac . That is, pa determines the size of the conference measured in terms of the number of papers accepted.

After the reviewing process, the conference updates the images $i \in I$ of each reviewer in R , according to the disagreement with the other reviewers of the same paper. This disagreement is calculated for each paper as the difference between the review value given by the reviewer and the average review value for that paper. When this difference gets higher than a disagreement threshold d , the reviewer disagreement count grows by one; values are recorder in an image representation of the form $i = \langle r, nd, nr \rangle$, where r is the reviewer, nd is the accumulated number of disagreements and nr is the total number of reviews carried out. These images are then used to discard the e reviewers with a higher ratio nd/nr and select e new reviewers from R . This way, conferences perform a selection process which selects reviewers who provide similar evaluations. Given our choice for reviewers' mistakes (if they don't understand the paper, the evaluation is random), this mechanism should also select good reviewers.

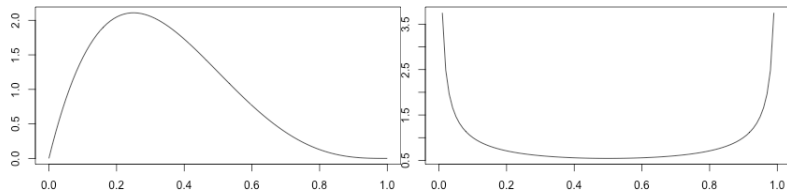


Fig. 1. Beta distributions used in the paper. From left to right, values for (α, β) : $(2, 4)$ corresponding to low skill reviewers, $(0.4, 0.4)$ corresponding to a mix of very good and very bad reviewers. The uniform distribution, also used in the paper, is not shown.

3 Implementation details

The PR-1 model has been programmed as a MAS over Jason [2], which allows the definition of BDI agents using an extended version of AgentSpeak(L) [10]. This MAS represents both conferences and reviewers as agents interacting in a common environment.

The reviews carried out by the pool of reviewers can be simply programmed in AgentSpeak(L) as shown in Table 1. Here, we use the belief `skill` to set the skill value associated with each reviewer. As already mentioned, we can change the distribution of these values through the (α, β) parameters of a beta distribution (lines 1–3). Each time the reviewer has to review a paper, the `+?review` test plan is executed (lines 6–11). Then, the review is accurate with probability S , and completely random with probability $(1 - S)$.

Table 1. `reviewer.asl` file defining the reviewer’s behavior.

```

1 skill(tools.beta(1,1)).           // Uniform distribution
2 // skill(tools.beta(2,4)).        // Low skill reviewers
3 // skill(tools.beta(0.4,0.4)).    // Polarized reviewer skill
4
5 // Plan to review papers
6 +?review(IdPaper, Value, Review) : skill(S) & paper_scale_values(N)
7   <- if (math.random < S)
8     {
9       Review = Value
10    } else {
11      Review = math.floor(math.random(N)) + 1
12    }

```

Conferences can be configured through a set of beliefs in the `conference.asl` file. Table 2 shows the ontology of beliefs used to set parameters such as: the number of papers received (`n_received_papers`), how many of them can be accepted (`max_papers_accepted`) or the number of discordant reviewers that get substituted per year (`n_reviewers_exchanged`). Additionally, a set of image beliefs will be managed by each conference in order to represent the images of the reviewers in the pool. In addition to the previous beliefs, the `conference.asl` file contains the set of plans dealing with the goals involved in the peer review system. Table 3 shows some snippets of the plans in this file. For instance, the plan `!celebrateConference` (lines 1–4) first launches the subgoal related to the reviewing process (`!reviewProcess`). For each paper received, a number of RxP reviews are collected (line 11). Then, the conference accepts the best PA papers (lines 30–34) amongst those exceeding the acceptance value AC (lines 13–17). The image of the reviewers in the PC is updated according to the disagreements with the other reviewers of the same paper (lines 19–27). These new images will be used to satisfy the goal of updating the members of the PC (`!updateReviewers`) in line 3.

Table 2. The ontology by the conference agents.

Belief formula	Description
n.received.papers (NP)	NP is the amount of papers received by the conference.
reviews.x.paper (RP)	RP is the number of reviews done for each paper.
papers.x.reviewer (PR)	PR is the number of papers reviewed by each reviewer.
max.papers.accepted (PA)	PA is the maximum number of papers the conference accepts.
paper.scale.values (N)	N is the scale of values for the papers.
accept.value (AC)	AC is minimum value for a paper to be accepted.
image (R, ND, NR)	R is the number of the reviewer, ND is the accumulated number of disagreements, and NR is the number of reviews done by reviewer R .
disagreement.threshold (D)	D is the disagreement threshold to punish reviewers.
n.reviewers.exchanged (E)	E is the number of reviewers exchanged each year.

4 Results

As a proof of concept, in this paper we show what happens in our simplified model (PR-1) if we change the distribution of reviewers' abilities. Thus, we experiment with different initial probability distributions for the only characteristic of reviewers' skill, that is, the probability that a reviewer gets his/her paper right. We consider three cases, that is, uniform ability, low average skill, and polarized skill¹. The shape of the beta distributions that we apply are shown in Fig. 1.

For this first set of experiments, we have ten conferences (which are essentially the same) receiving 100 submissions each (np), drawn from a uniform distribution. Papers are assigned an intrinsic value in a 10-values ordered scale, interpretable as the standard from strong reject to enthusiastically accept scores. We have fixed $pa = 100$ and $ac = 5.5$, so that all papers whose average review value is greater than 5 are accepted. We have set $rp = pr = 3$, i.e., the same number of reviews per paper and per PC member. Thus, a conference will need as many reviewers as it receives papers. That is, conferences will employ 100 reviewers each from a pool of 500 reviewers. There is no limit to PC memberships for an individual reviewer. Ideally, the same group of 100 reviewers could constitute the PC of all ten conferences. Finally, we use a disagreement threshold of 4 (d) and a 10% reviewer turnover rate ($e = 10$). Substitute reviewers are selected randomly in the pool.

4.1 Measured quantities

For each set of experiments, we measure several quantities, that we present, in their time evolution, in the following figures. The results are presented with five number summary (the central line marks the median, then the successive quartiles), collecting together the data of the different conferences (that are equivalent in PR-1) and in a window of five consecutive years.

¹ High average skill is not considered because the uniform distribution already yields a high quality selection process.

Table 3. Plan snippets from the conference.asl file.

```

1  +!celebrateConference(Year)
2    <- !reviewProcess;
3      !updateReviewers;
4      !!celebrateConference(Year + 1).
5
6  +! reviewProcess : n_received_papers(RP) & reviews_x_paper(RxP) &
7    accept_value(AC) & max_papers_accepted(PA) &
8    disagreement_threshold(D) & ...
9    <- for ( .range(PaperId,1,RP) ) {
10      PaperValue = math.floor(math.random(MaxValue)) + 1;
11      for ( .range(1,1, RxP) ) { /* Ask for reviews ... */ }
12      // Evaluate the paper
13      .findall(Review, review(PaperId,_,Review), Reviews);
14      AvgReview = math.average(Reviews);
15      if ( AvgReview > AC ) {
16        +accepted_paper(PaperId, PaperValue, AvgReview);
17      }
18      // Update the image of the reviewers
19      for ( review(PaperId, R, Review) ) {
20        ?image(R, ND, NR);
21        .abolish(image(R,_,_));
22        if ( math.abs(AvgReview - Review) > D ) {
23          +image(R, ND+1, NR+1);
24        } else {
25          +image(R, ND, NR+1);
26        }
27      }
28    }
29    // Limit the number of accepted papers
30    while ( .count(accepted_paper(_,_,_)) > PA ) {
31      .findall(acc_paper(R, PId), acc_paper(PId,_,R), AcceptedPapers);
32      .min(AcceptedPapers, acc_paper(_,PaperIdMin));
33      .abolish(accepted_paper(PaperIdMin,_,_));
34    }

```

The average accepted quality is the primary measure of success for the selection system. Paper quality, if the review process works perfectly, should select 20 top score papers, 20 with quality 9, and 10 of quality 8, leading to an ideal score of 9.2. The worst possible case (papers are accepted completely at random), as a reference value, would simply be the mean of scores from one to ten, amounting to 5.5.

In parallel to the paper selection process, based on disagreement measures between reviewers, program committees are reorganized with the aim to select the best reviewers. Thus, another quantity we measure is the average quality of reviewers that are part of program committees, under different initial conditions for their distribution. In principle, better reviewers should select better papers.

We also show the number of good papers (i.e., with an intrinsic value greater than 5.5) rejected, and the number of bad papers (i.e., with an intrinsic value less than 5.5) accepted. While the previous quantities can be seen as measures of efficiency, these two can be thought of as measures of fairness.

In fact, good papers rejected and bad papers accepted are especially important because of the high-stakes nature of investment that researchers do on each

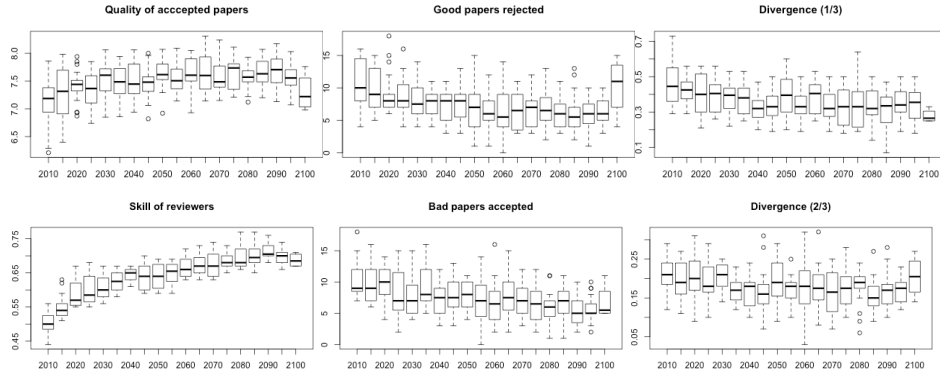


Fig. 2. Results (shown as five-number summary) for a beta distribution with parameters $(1.0, 1.0)$, that is, a uniform distribution, averaged over ten conferences and in periods of five years. *First column*, above, average quality of accepted papers; below, quality of reviewers. Both observable quantities improve substantially in time. *Second column*, above, good papers rejected, below, bad papers accepted, both showing a small improvement in time. *Third column*, divergence values calculated at $1/3$ and $2/3$, both decreasing in time.

paper - on the one hand, an “out-of-the-blue” rejection can seriously impact career, especially in small research groups; on the other hand, the publication of bogus papers creates a stigma on journals and conferences. Finally, another interesting measure of success for a conference review process had been defined in [4] as *divergence*: the normalized distance between the ordering of the accepted papers, and the ordering induced by another quality measure.

In [4], divergence was calculated with real data of an anonymised “large conference”, comparing review results against paper citation rates registered five years later. We perform a similar calculation, not against citation rates but against our idealized paper quality. The distance used is calculated simply by the (normalized) number of elements ranked in the top ($1/3$ or $2/3$) by the review process that are not in the top ($1/3$ or $2/3$) in the ideal quality ordering. The result for the large conference, that the authors of [4] claim to be disappointingly comparable to random sorting, is a value of 0.63 at $1/3$ and 0.32 at $2/3$.

Note how this ordering concerns only the set of accepted papers; good rejected papers or bad accepted ones do not enter this calculation. This value can be considered as another measure of efficiency of the system: the lower it is, the more efficient the peer review.

4.2 Uniform ability

Here we show the results obtained from a reviewer skill distribution with parameters $(1,1)$ - a uniform distribution.

From figure 2, we can see how the quality of accepted papers starts already over the average. The process improves in time for both the paper quality and

reviewer skill; however, only the second has a significant effect. The convergence process seems to manage selecting good reviewers, but this happens without a substantial quality improvement. Mistakes in paper evaluations show only a slight decrease. Finally, divergence from the optimal acceptance ordering remains constant - perhaps after a slight improvement in the first years. At about 0.28 and 0.2, it remains far better than the levels 0.63 and 0.32 reported in [4].

4.3 Low average skill

Apparently, our simulated reviewers perform better than the real ones. What if we decrease their average skill, for example drawing them from a beta distribution with parameters (2.0, 4.0), shown in figure 1 (left)? The results are presented in figure 3. With such a bad average reviewer skill, the quality of accepted papers results lower than in the previous case, and the agreement process yields no or little improvement in time - except in reviewers skill, whose increase however does not seem enough to improve the quality of accepted papers. There just aren't enough good reviewers around to make the process work. Good papers rejected and bad papers accepted abound, making up for more than half the body of accepted papers; divergence, ending at 0.6 and 0.25, seems directly comparable to the values in [4].

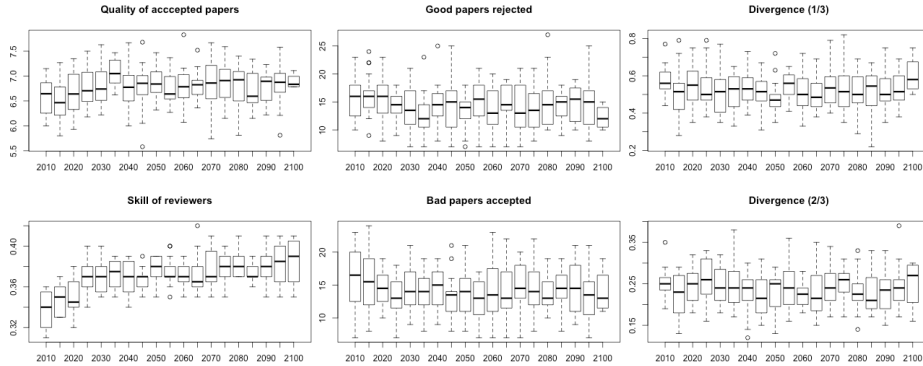


Fig. 3. Results (shown as five-number summary) for a beta distribution with parameters (2.0, 4.0), averaged over ten conferences and in periods of five years. *First column*, above, average quality of accepted papers; below, quality of reviewers. There is no substantial improvement, apart from an increase in reviewers quality. *Second column*, above, good papers rejected, below, bad papers accepted, both stable in time. *Third column*, divergence values calculated at 1/3 and 2/3.

4.4 Polarized skill

So far we have shown a relatively good selection process, starting with reviewers with uniform distribution, and a relatively bad one, where most reviewers are of

low skill. With yet another shape of the skill distribution, we want to measure how effective the agreement process is in selecting good reviewers. To this purpose, we choose an initial distribution with a double peak - in this experiment, as can be seen from figure 1 (right), most reviewers are very bad or very good. We surely have more than enough good ones for a nearly perfect review process - but will the system be able to select them? Figure 4 shows this is indeed the case. This time, the success of the reviewer selection process takes the average paper quality up with it, obtaining better results than in the uniform case. There are nearly no bad papers accepted, nor good papers rejected towards the end. Divergence is similarly affected, leveling at 0.2 and 0.12 at the end.

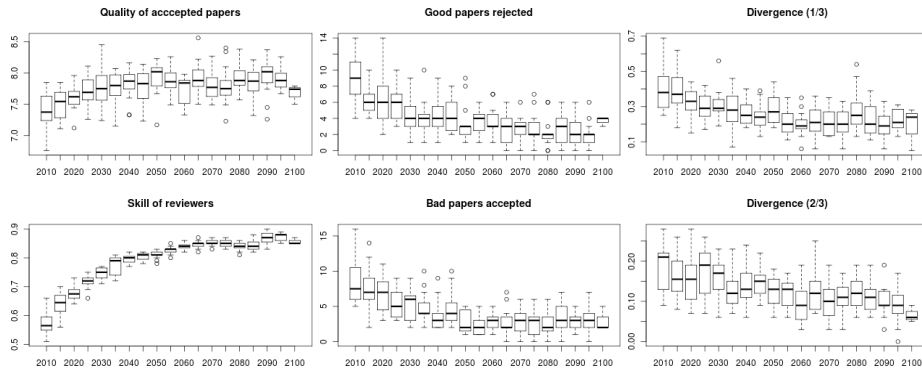


Fig. 4. Results (shown as five-number summary) for a beta distribution with parameters $(0.4, 0.4)$, averaged over ten conferences and in periods of five years. *First column*, above, average quality of accepted papers; below, quality of reviewers. Both observable quantities improve substantially in time. *Second column*, above, good papers rejected, below, bad papers accepted. Both show a marked decrease in time. *Third column*, divergence values calculated at $1/3$ and $2/3$, both decreasing in time.

5 Discussion and Future work

This paper is a first step towards a model of peer review devoted to study and to enhance the way of evaluating scientific research. We have sketched the main elements involved as well as the relations amongst them. A first restricted version of the full model, that we call PR-1, has been implemented as a MAS over Jason. The results show how a conference review process based on disagreement control between reviewers can i) improve in time both the quality of accepted papers and the reviewer skill of PC members; ii) reduce the number of good papers rejected and bad papers accepted; and iii) lower the divergence between the ordering of the accepted papers and an ideal quality ordering. Reviewer selection improves on both the efficiency and the fairness of the review process. The results, for

what regards a measure of divergence between reviews and actual quality of the paper, are shown to be qualitatively comparable with the observed data in [4].

Quite a large number of issues still remain open for future work. Limiting PC memberships for individual reviewers and considering role superposition between author and reviewer is one of the next steps. Furthermore, subsequent versions of the PR model should include the active role of the authors when deciding which conference to send their works to, as it can vary the distribution of the papers submitted to a conference.

Acknowledgements

The first author has been jointly supported by the Spanish MEC and the European Commission FEDER funds, under grants Consolider-Ingenio 2010 CSD2006-00046, TIN2009-14475-C04-04 and JC2010-0062. We are indebted with the reviewers for many improvements to this version of the paper.

References

1. E. Bonabeau. Agent-based modeling: methods and techniques for simulating human systems. *Proceedings of the National Academy of Sciences of the United States of America*, 99 Suppl 3(Suppl 3):7280–7287, May 2002.
2. R. H. Bordini and J. F. Hübner. Jason. Available at <http://jason.sourceforge.net/>, March 2007.
3. T. Brabazon. The google effect: Googling, blogging, wikis and the flattening of expertise. *Libri*, 56:157–167, 2006.
4. F. Casati, M. Marchese, A. Ragone, and M. Turrini. Is peer review any good? a quantitative analysis of peer review. Technical report, Ingegneria e Scienza dell’Informazione, University of Trento, 2009.
5. R. Conte and C. Castelfranchi. *Cognitive Social Action*. London: UCL Press, 1995.
6. T. S. Kuhn. *The Structure of Scientific Revolutions*. University Of Chicago Press, 3rd edition, December 1996.
7. B. D. Neff and J. D. Olden. Is peer review a game of chance? *BioScience*, 56(4):333–340, April 2006.
8. M. Paolucci, T. Balke, R. Conte, T. Eymann, and S. Marmo. Review of internet user-oriented reputation applications and application layer networks. *Social Science Research Network Working Paper Series*, September 2009.
9. M. Paolucci, S. Picascia, and S. Marmo. Electronic reputation systems. In *Handbook of Research on Web 2.0, 3.0, and X.0*, chapter chapter 23, pages 411–429. IGI Global, 2010.
10. A. S. Rao. AgentSpeak(L): BDI agents speak out in a logical computable language. In S. Verlag, editor, *Proc. of MAAMAW’96*, number 1038 in LNAI, pages 42–55, 1996.
11. S. Thurner and R. Hanel. Peer-review in a world with rational scientists: Toward selection of the average. Aug. 2010.
12. J. Wainberg, T. Kida, and J. F. Smith. Stories vs. statistics: The impact of anecdotal data on accounting decision making. *Social Science Research Network Working Paper Series*, March 2010.

An Agent-Based Proxemic Model for Pedestrian and Group Dynamics: Motivations and First Experiments

Lorenza Manenti^{1,2}, Sara Manzoni^{1,2}, Giuseppe Vizzari^{1,2}, Kazumichi Ohtsuka³, and Kenichiro Shimura³

¹ Complex Systems and Artificial Intelligence (CSAI) research center
Department of Computer Science, Systems and Communication (DISCo)
University of Milan - Bicocca, Viale Sarca 336/14, 20126 Milano, Italy
{manenti, manzoni, vizzari}@disco.unimib.it

² Crystal Project, Centre of Research Excellence in Hajj and Omrah (Hajjcore)
Umm Al-Qura University, Makkah, Saudi Arabia.

³ Research Center for Advanced Science & Technology, The University of Tokyo, Japan
tukacyf@mail.ecc.u-tokyo.ac.jp, shimura@tokai.t.u-tokyo.ac.jp

Abstract. The simulation of pedestrian dynamics is a consolidated area of application for agent-based models: successful case studies can be found in the literature and off-the-shelf simulators are commonly employed by end-users, decision makers and consultancy companies. These models, however, generally consider individuals, their interactions with the environment and among themselves, but they generally neglect (or treat in a simplistic way) aspects like (i) the impact of cultural heterogeneity among individuals and (ii) the effects of the presence of groups and particular relationships among pedestrians. This work is aimed, on one hand, at introducing some fundamental anthropological considerations on which most pedestrian models are based, and in particular Edward T. Hall's work on proxemics. On the other hand, the paper describes an agent-based model encapsulating in the pedestrian's behavioural model effects representing both proxemics and a simplified account of influences related to the presence of groups in the crowd. The model is tested in a simple scenario to evaluate the implications of some modeling choices and the presence of groups in the simulated scenario. Results are discussed and compared to experimental observations and to data available in the literature.

Key words: crowd modeling and simulation, agent-based models

1 Introduction

There are several features of crowds of pedestrians suggesting that they can be considered as complex entities: the mix of competition for the space shared by pedestrians and the collaboration due to the (not necessarily explicit but generally shared) social norms, the dependency of individual choices on the past actions of other individuals and on the current perceived state of the system, the possibility to detect self-organization and emergent phenomena they are all indicators of the intrinsic complexity of a crowd. The relevance of human behaviour, and especially of the movements of pedestrians, in built

environment in normal and extraordinary situations, and its implications for the activities of architects, designers and urban planners are apparent (see, e.g., [1] and [2]), especially considering dramatic episodes such as terrorist attacks, riots and fires, but also due to the growing issues in facing the organization and management of public events (ceremonies, races, carnivals, concerts, parties/social gatherings, and so on) and in designing naturally crowded places (e.g. stations, arenas, airports). Computational models for the simulation of crowds are thus growingly investigated in the scientific context, and these efforts led to the realization of commercial off-the-shelf simulators often adopted by firms and decision makers⁴. Even if research on this topic is still quite lively and far from a complete understanding of the complex phenomena related to crowds of pedestrians in the environment, models and simulators have shown their usefulness in supporting architectural designers and urban planners in their decisions by creating the possibility to envision the behaviour/movement of crowds of pedestrians in specific designs/environments, to elaborate what-if scenarios and evaluate their decisions with reference to specific metrics and criteria.

A Multi-Agent Systems (MAS) approach to the modeling and simulation of complex systems has been applied in very different contexts, from the study of social systems [3], to biology (see, e.g., [4]), and it is considered one of the most successful types of applications of agent-based computing [5], even if this approach is still relatively young, compared, for instance, to analytical equation-based modeling. The MAS approach has also been adopted in the pedestrian and crowd modeling context, especially due to the expressiveness of the approach, that is particularly suited to the definition of models in which autonomous and possibly heterogeneous agents can be defined, situated in an environment, provided with the possibility to perceive it, decide and try to carry out the most appropriate line of action, possibly interacting with other agents as well as the environment itself. The approach can lead to the definition of models that are richer and more expressive than other approaches that were traditionally adopted in the modeling of pedestrians, that respectively consider pedestrians as particles subject to forces (see, e.g., [6]), in physical approaches, or particular states of cells in which the environment is subdivided, in CA approaches (see, e.g., [7, 8]).

The main aim of this work is to present the motivations, fundamental research questions and directions, and some preliminary results of an agent-based modeling and simulation approach to the multidisciplinary investigation of the complex dynamics that characterize aggregations of pedestrians and crowds. This work is set in the context of the Crystals project⁵, a joint research effort between the Complex Systems and Artificial Intelligence research center of the University of Milano–Bicocca, the Centre of Research Excellence in Hajj and Omrah and the Research Center for Advanced Science and Technology of the University of Tokyo. The main focus of the project is on the adoption of an agent-based pedestrian and crowd modeling approach to investigate meaningful relationships between the contributions of anthropology, cultural characteristics and existing results on the research on crowd dynamics, and how the presence

⁴ see, e.g., Legion Ltd. (<http://www.legion.com>), Crowd Dynamics Ltd. (<http://www.crowddynamics.com/>), Savannah Simulations AG (<http://www.savannah-simulations.ch>).

⁵ <http://www.csai.disco.unimib.it/CSAI/CRYSTALS/>

of heterogeneous groups influence emergent dynamics in the context of the Hajj and Omrah. The last point is in fact an open topic in the context of pedestrian modeling and simulation approaches: the implications of particular relationships among pedestrians in a crowd are generally not considered or treated in a very simplistic way by current approaches. In the specific context of the Hajj, the yearly pilgrimage to Mecca that involves over 2 millions of people coming from over 150 countries, the presence of groups (possibly characterized by an internal structure) and the cultural differences among pedestrians represent two fundamental features of the reference scenario. Studying implications of these basic features is the main aim of the Crystals project.

The paper breaks down as follows: the following section describes some basic anthropological and sociological theories that were selected to describe the phenomenologies that will be considered in the agent-based model definition. Section 3 briefly describes a model encompassing basic anthropological rules for the interpretation of mutual distances by agents and basic rules for the cohesion of groups of pedestrians, while Section 4 summarizes the results of the application of this model in a simple simulation scenario. Conclusions and future developments will end the paper.

2 Interdisciplinary Approach

Pedestrian and crowd modeling research context regards events in which a large number of people may be gathered or bound to move in a limited area; this can lead to serious safety and security issues for the participants and the organizers. The understanding of the dynamics of large groups of people is very important in the design and management of any type of public events. In addition to safety and security concerns also the comfort of event participants is another aim of the organizers and managers of crowd related events. Large people gatherings in public spaces (like pop-rock concerts or religious rites participation) represent scenarios in which crowd dynamics can be quite complex due to different factors (the large number and heterogeneity of participants, their interactions, their relationship with the performing artists and also exogenous factors like dangerous situations and any kind of different stimuli present in the environment [9]). The traditional and current trend in social sciences studying crowds is still characterized by a non-dominant behavioral theory on individuals and crowds dynamics, although it is recognized that a behavioural theory is needed to improve the current state of the art in pedestrian and crowd modeling and simulation [10].

2.1 Proxemics

The term *proxemics* was first introduced by Hall with respect to the study of set of measurable distances between people as they interact [11]. In his studies, Hall carried out analysis of different situations in order to recognize behavioral patterns. These patterns are based on people's culture as they appear at different levels of awareness. In [12] Hall proposed a system for the notation of proxemic behavior in order to collect data and information on people sharing a common space. Hall defined proxemic behavior and four types of perceived distances: *intimate distance* for embracing, touching or whispering;

personal distance for interactions among good friends or family members; *social distance* for interactions among acquaintances; *public distance* used for public speaking. Perceived distances depend on some additional elements which characterize relationships and interactions between people: posture and sex identifiers, sociofugal-sociopetal (SFP) axis, kinesthetic factor, touching code, visual code, thermal code, olfactory code and voice loudness.

Proxemic behavior includes different aspects which could be useful and interesting to integrate in crowd and pedestrian dynamics simulation. In particular, the most significant of these aspects being the existence of two kinds of distance: *physical* distance and *perceived* distance. While the first depends on physical position associated to each person, the latter depends on proxemic behavior based on culture and social rules.

It must be noted that some recent research effort was aimed at evaluating the impact of proxemics and cultural differences on the fundamental diagram [13], a typical way of evaluating both real crowding situations and simulation results.

2.2 Crowds: Canetti's Theory

Elias Canetti's work [14] proposes a classification and an ontological description of the crowd; it represents the result of 40 years of empirical observations and studies from psychological and anthropological viewpoints. Elias Canetti can be considered as belonging to the tradition of social studies that refer to the crowd as an entity dominated by uniform moods and feelings. We preferred this work among others dealing with crowds due to its clear semantics and explicit reference to concepts of loss of individuality, crowd uniformity, spatio-temporal dynamics and *discharge* as a triggering entity generating the crowd, that could be fruitfully represented by computationally modeling approaches like MAS.

The normal pedestrian behaviour, according to Canetti, is based upon what can be called the *fear to be touched* principle:

"There is nothing man fears more than the touch of the unknown. He wants to see what is reaching towards him, and to be able to recognize or at least classify it."

"All the distance which men place around themselves are dictated by this fear."

A discharge is a particular event, a situation, a specific context in which this principle is not valid anymore, since pedestrians are willing to accept being very close (within touch distance). Canetti provided an extensive categorization of the conditions, situations in which this happens and he also described the features of these situations and of the resulting types of crowds. Finally, Canetti also provides the concept of *crowd crystal*, a particular set of pedestrians which are part of a group willing to preserve its unity, despite crowd dynamics. Canetti's theory (and precisely the fear to be touched principle) is apparently compatible with Hall's proxemics, but it also provides additional concepts that are useful to describe phenomena that take place in several relevant crowding phenomena, especially from the Hajj perspective.

Recent developments aimed at formalizing, embedding and employing Canetti's crowd theory into computer systems (for instance supporting crowd profiling and modeling) can be found in the literature [15, 16] and they represent a useful contribution to

the present work. Two recent works represent a relevant effort towards the modeling of groups, respectively in particle-based [17] and in CA-based [18] approaches: in both cases, groups are modeled by means of additional contributions to the overall pedestrian behaviour representing the tendency to stay close to other group members.

3 An Agent-Based Proxemic Model

This section will describe a first step towards an agent-based model encompassing abstractions and mechanisms accounting based on fundamental considerations about proxemics and basic group behaviour in pedestrians. We first defined a very general and simple model for agents, their environment and interaction, then we realized a proof-of-concept prototype to have an immediate idea of the implications of our modeling choices.

The simulated environment represents a simplified real built environment, a corridor with two exits (North and South); later different experiments will be described with corridors of different size (10m wide and 20 m long as well as 5m wide and 10 m long). We represented this environment as a simple euclidean bi-dimensional space, that is discrete (meaning that coordinates are integer numbers) but not “discretized” (as in a CA). Pedestrians, in other words, are characterized by a position that is a pair $\langle x, y \rangle$ that does not denote a cell but rather admissible coordinates in an euclidean space. Movement, the fundamental agent’s action, is represented as a displacement in this space, i.e. a vector. The approach is essentially based on the Boids model [19], in which however rules have been modified to represent the phenomenologies described by the basic theories and contributions on pedestrian movement instead of flocks. Boundaries can also be defined: in the example Eastern and Western borders cannot be crossed and the movement of pedestrians is limited by the pedestrian position update function, which is an environmental responsibility. Every agent $a \in A$ (where A is the set of agents representing pedestrians of the modeled scenario) is characterized by a position pos_a represent by a pair of coordinates $\langle x_a, y_a \rangle$. Agent’s action is thus represented by a vector $\overline{m}_a = \langle \delta_{x_a}, \delta_{y_a} \rangle$ where $|\overline{m}_a| = \sqrt{\delta_{x_a}^2 + \delta_{y_a}^2} < M$ where M is a parameter depending on the specific scenario representing the maximum displacement per time unit.

More complex environments could be modeled, for instance by means of a set of relevant objects in the scene, like points of interest but also obstacles. These objects could be perceived by agents according to their position and perceptual capabilities, and they could thus have implications on their movement. Objects can (but they do not necessarily must be) in fact be considered as attractive or repulsive by them. The effect of the perception of objects and other pedestrians, however, is part of agents’ behavioural specifications. For this specific application, however, the perceptive capability of an agent a are simply defined as the set of other pedestrians that are present at the time of the perception in a circular portion of space or radius r_p centered at the current coordinates of agent a . In particular, each agent $a \in A$ is provided with a perception distance per_a ; the set of perceived agents is defined as $P_a = p_1, \dots, p_i$ where $d(a, i) = \sqrt{(x_a - x_i)^2 + (y_a - y_i)^2} \leq per_a$.

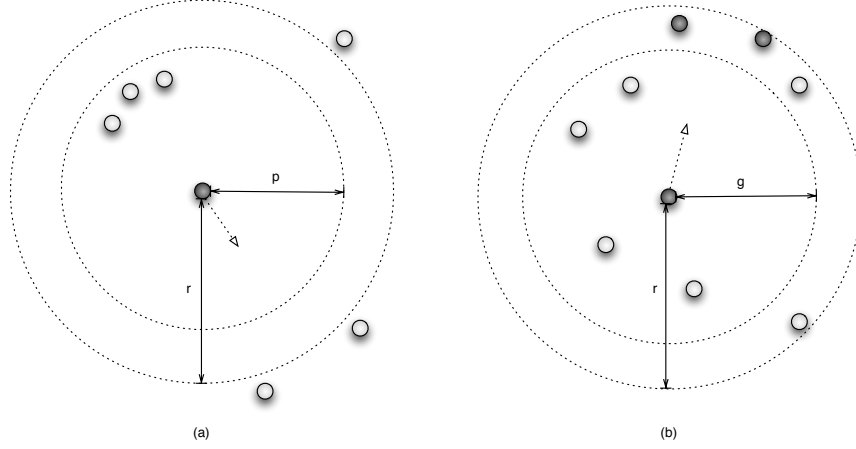


Fig. 1. Basic behavioural rules: a basic proxemic rule drives an agent to move away from other agents that entered/are present in his/her own personal space (delimited by the proxemic distance p) (a), whereas a member of a group will pursue members of his/her group that have moved/are located beyond a certain distance (g) but within his/her perception radius (r) (b).

Pedestrians are modeled as agents situated in an environment, each occupying about 40 cm^2 , characterized by a state representing individual properties. Their behaviour has a goal driven component, a preferred direction; in this specific example it does not change over time and according to agent's position in space (agents want to get out of the corridor from one of the exits, either North or South), but it generally changes according to the position of the agent, generating a path of movement from its starting point to its own destination. The preferred direction is thus generally the result of a stochastic function possibly depending on time and current position of the agent. The goal driven component of the agent behavioural specification, however, is just one of the different elements of the agent architectures that must include elements properly capturing elements related to general proxemic tendencies and group influence (at least), and we also added a small random contribution to the overall movement of pedestrians, as suggested by [20]. The actual layering of the modules contributing to the overall is object of current and future work. In the scenario, agents' goal driven behavioural component is instead rather simple: agents heading North (respectively South) have a deliberate contribution to their overall movement $\bar{m}_a^g = \langle 0, M \rangle$ (respectively $\bar{m}_a^g = \langle 0, -M \rangle$).

We realized a simulation scenario in a rapid prototyping framework⁶ and we employed it to test the simple behavioural model that will be described in the following. In the realized simulator, the environment is responsible for updating the position of agents, actually triggering their action choice in a sequential way, in order to ensure fairness among agents. In particular, we set the turn duration to 100 ms and the maximum covered distance in one turn is 15 cm (i.e. the maximum velocity for a pedestrian is 1.5 m/s).

⁶ Nodebox – <http://www.nodebox.org>

3.1 Basic Proxemic Rules

Every pedestrian is characterized by a culturally defined proxemic distance p ; this value is in general related to the specific culture characterizing the individual, so the overall system is designed to be potentially heterogenous. In a normal situation, the pedestrian moves (according to his/her preferred direction) maintaining the minimum distance from the others above this threshold (rule *P1*). More precisely, for a given agent a this rule defines that the proxemic contribution to the overall agent movement $\overline{m}_a^p = 0$ if $\forall b \in P_a : d(a, b) \geq p$.

However, due to the overall system dynamics, the minimum distance between one pedestrian and another can drop below p . In this case, given a pedestrian a , we have that $\exists b \in P_a : d(a, b) < p$; the proxemic contribution to the overall movement of a will try to restore this condition (rule *P2*) (please notice that pedestrians might have different thresholds, so b might not be in a situation so that his/her *P2* rule is activated). In particular, given $p_1, \dots, p_k \in P_a : d(a, p_i) < p$ for $1 \leq i \leq k$, given c the centroid of $pos_{p_1}, \dots, pos_{p_k}$, the proxemic contribution to the overall agent movement $\overline{m}_a^p = -k_p \cdot c - pos_a$, where k_a is a parameter determining the intensity of the proxemic influence on the overall behaviour.

These basic considerations, also schematized in Figure 1-A, lead to the definition of rules support a basic proxemic behaviour for pedestrian agents; these agents are not characterized by any particular relationship binding them, with the exception of a shared goal, i.e. they are not a group but rather an unstructured set of pedestrians.

3.2 Group Dynamic Rules

We extended the behavioural specification of agents by means of an additional contribution representing the tendency of group members to stay close to each other. First of all, every pedestrian may be thus part of a group, that is, a set of pedestrians that mutually recognize their belonging to the same group and that are willing to preserve the group unity. This is clearly a very simplified, heterarchical notion of group, and in particular it does not account for hierarchical relationships in groups (e.g. leader and followers), but we wanted to start defining basic rules for the simplest form of group.

Every pedestrian is thus also characterized by a culturally defined proxemic distance g determining the way the pedestrian interprets the minimum distance from any other group member. In particular, in a normal situation a pedestrian moves (according to his/her preferred direction and also considering the basic proxemic rules) keeping the maximum distance from the other members of the group below g (Rule *G1*). More precisely, for a given agent a , member of a group G , this rule defines that the group dynamic contribution to the overall agent movement $\overline{m}_a^g = 0$ if $\forall b \in (P_a \cap (G - \{a\})) : d(a, b) < g$.

However, due to the overall system dynamics, the maximum distance between one pedestrian and other members of his group can exceed g . In this case, he/she will try to restore this condition by moving towards the group members he/she is able to perceive (rule *G2*). In particular, given $p_1, \dots, p_k \in (P_a \cap (G - \{a\})) : d(a, p_i) \geq g$ for $1 \leq i \leq k$, given c the centroid of $pos_{p_1}, \dots, pos_{p_k}$, the proxemic contribution to the

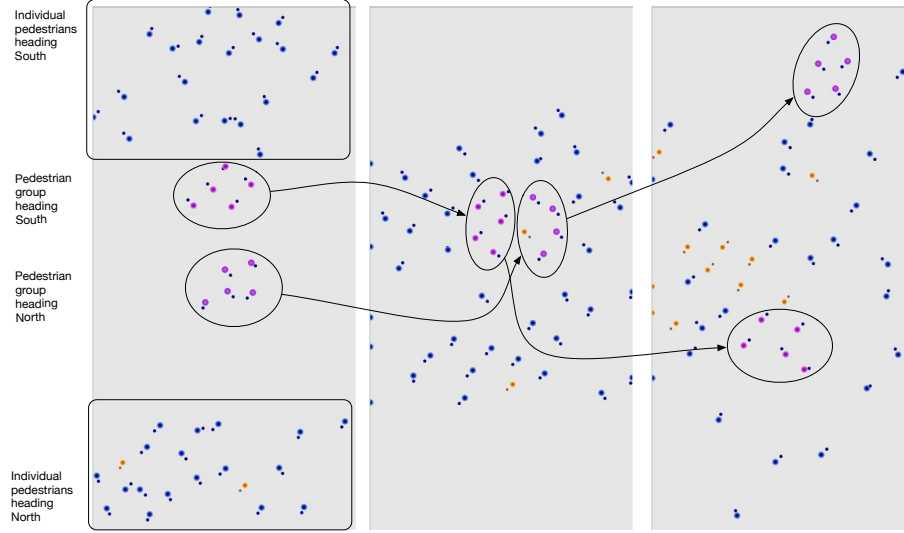


Fig. 2. Screenshots of the prototype of the simulation system.

overall agent movement $\overline{m}_a^g = k_g \cdot \overline{c - pos_a}$, where k_a is a parameter determining the intensity of the group dynamic influence on the overall behaviour.

This basic idea of group influence on pedestrian dynamics, also schematized in Figure 1-B, lead to the extension of the basic proxemic behaviour for pedestrian agents of the previous example. We tested the newly defined rules in a similar scenario but including groups of pedestrians. In particular, two scenarios were analyzed. In the first one, we simply substituted 4 individual pedestrians in the previous scenario with a group of 4 pedestrians. The group was able to preserve its unity in all the tests we conducted, but the average travel time for the group members actually increased. Individuals, in other words, trade some of their potential speed to preserve the unity of the group. In a different scenario, we included 10 pedestrians and a group of 4 pedestrians heading North, 10 pedestrians and a group of 4 pedestrians heading South. In this circumstances, the two groups sometimes face and they are generally able to find a way to form two lanes, actually avoiding each other. However, the overall travel time for group members actually increases in many of the simulations we conducted.

In Figure 2 two screenshots the of the prototype of the simulation system that was briefly introduced here. Individual agents, those that are not part of a group, are depicted in blue, but those for which rule P2 is activated (they are afraid to be touched) turn to orange, to highlight the invasion of their personal space. Members of groups are depicted in violet and pink. The two screenshots show how two groups directly facing each other must manage to “turn around” each other to preserve their unity but at the same time advance towards their destination.



Fig. 3. Experiments on facing groups: several experiments were conducted on real pedestrian dynamics, some of which also considered the presence of groups of pedestrians, that were instructed on the fact that they had to behave as friends or relatives while moving during the experiment.

4 Experimental results

4.1 Proxemic Distance Evaluation

We conducted several experiments with the above described model and simulator, to evaluate the plausibility of the overall system dynamics achieved with such simple basic rules and to calibrate the parameters to fit actual data available from the literature or acquired in the experiments. In particular, we first of all focused on the influence of the proxemic distance p on the overall system dynamics. We started considering Hall's personal distance as a starting point for this model parameter. Hall reported ranges for the various proxemic distances, considering a close phase and a far phase for all the different perceived distances (described in Section 2.1). In particular, we considered both an average value for the far phase of the personal distance (1m) and a low end value (75 cm) that is actually the border between the far and the close phases of the personal distance range. In general, the higher value allowed to achieve relatively results in scenarios characterized by a low density of pedestrians in the environment. For densities close and above one pedestrian per square meter, the lower value allowed achieving a smoother flow, more consistent with the results available in the literature.

A summary of the achieved results is shown in Figure 4: the graphs represent the fundamental diagrams [9] of the data achieved in the simulation of a 10 m long and 5 m wide corridor. For these experiments we considered that the influence of the different components of pedestrian behaviour (i.e. the weights of their contribution to the overall movement vector), that is goal attraction, proxemic repulsion and group cohesion, is equal for the first two components while the third is less significant (about one third of the previous ones). This setting supported a good balance between flow smoothness, collision avoidance and group cohesion in a preliminary face validation phase.

We varied the number of agents altering the density of pedestrians in the environment; to keep constant the number of pedestrians in the corridor, the two ends were joined (i.e. pedestrians exiting from one end were actually re-entering the corridor from the other). For each run only complete pedestrian trips were considered (i.e. the first

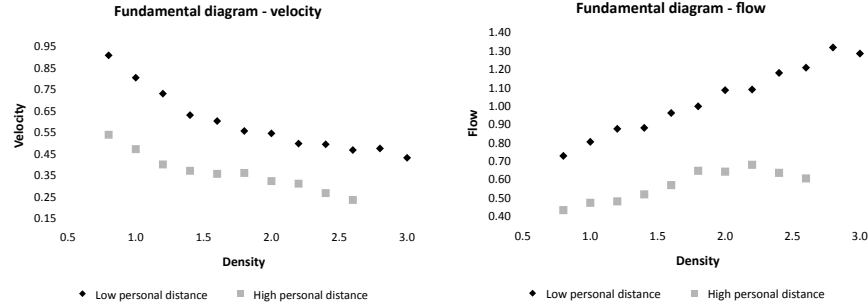


Fig. 4. Fundamental diagrams for the 10m long and 5m wide corridor scenario. The two data series respectively refer to different values for the proxemic distance, respectively the low end (75cm) and the average value (1m) of personal distance.

pedestrian exit event was discarded because related to a partial crossing of the corridor) and in high density scenarios a significant number of starting turns were also discarded to avoid transient starting conditions. The results of the simulations employing the low personal distance are consistent with empirical observations discussed in [9].

In parallel to the modeling effort, a set of experiments were conducted (in June 2010) to back-up with observed data some intuitions on the implications of the presence of groups in specific scenarios; two photos of one of the experiments are shown in Figure 3. In particular, this experiments is characterized by two sets of pedestrians moving in opposite directions in a constrained portion of space. In the set of pedestrians, in some of the experiments, some individuals were instructed to behave as friends or relatives, trying to stay close to each other in the movement towards their goal. It must be noted that this kind of situation is simple yet relevant for the understanding of some general principle on pedestrian movement and on the implications of the presence of groups in a crowd. The analysis of a first set of experiments were not conclusive on the effects of the presence of groups in the two facing sets of pedestrians (in some experiments the overall travel time was lower when groups were present, in other occasion it was longer), but the simulation results achieved with a low personal distance were consistently more in tune with the observed data than those achieved with a high personal distance.

4.2 Groups and Individuals

We also analyzed the implications of the presence of groups in the environment. The generated data, as well as the empirical observations, still do not lead to conclusive results. In simulations carried out in low density scenarios, however, the average speed of group members is consistently lower than the one of single individuals. It must be considered that, when compared to basic individuals, their overall movement has an additional component that sometimes contrasts the tendency to move towards the goal, to stay close to other group members. In high density scenarios, instead, the average speed of group members is generally higher than that of single individuals. This is

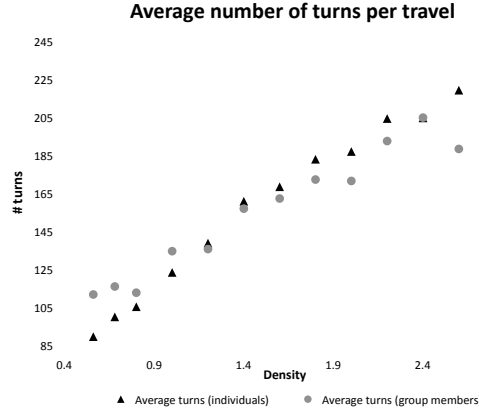


Fig. 5. Average number of turns per travel, individuals compared to group members.

probably due to the fact that the presence of the group has a greater influence on the possibility of other individuals to move, generating for instance a higher possibility of members on the back of the group to follow the “leaders”. Figure 5 compares the average number of turns per complete travel time of individuals and group members in the same conditions of the experiments described in the previous section.

5 Conclusions and Future Works

The paper has presented the research setting in which an innovative agent-based pedestrian crowd modeling and simulation effort is set. Preliminary results of the first stage of the modeling phase were described. Future works are aimed, on one hand, at consolidating the preliminary results of this first scenarios, but also extending the range of simulated scenarios characterized by relatively simple spatial structures for the environment (e.g. bends, junctions). On the other hand, we want to better formalize the agent behavioural model and its overall architecture, but we also plan to extend the notion of group, in order to capture phenomenologies that are particularly relevant in the context of Hajj (e.g. hierarchical groups, but also hierarchies of groups). Finally, we are working at the integration of these models into an existing open source framework for 3D computing (Blender⁷), also to be able to embed these models and simulations in real portions of the built environment defined with traditional CAD tools.

Acknowledgments

This work is a result of the Crystal Project, funded by the Centre of Research Excellence in Hajj and Omrah (Hajjcore), Umm Al-Qura University, Makkah, Saudi Arabia.

⁷ <http://www.blender.org/>

References

1. Batty, M.: Agent based pedestrian modeling (editorial). *Environment and Planning B: Planning and Design* **28** (2001) 321–326
2. Willis, A., Gjersoe, N., Havard, C., Kerridge, J., Kukla, R.: Human Movement Behaviour in Urban Spaces: Implications for the Design and Modelling of Effective Pedestrian Environments. *Environment and Planning B* **31**(6) (2004) 805–828
3. Axtell, R.: Why Agents? On the Varied Motivations for Agent Computing in the Social Sciences. Center on Social and Economic Dynamics Working Paper **17** (2000)
4. Christley, S., Zhu, X., Newman, S.A., Alber, M.S.: Multiscale Agent-based Simulation for Chondrogenic Pattern Formation *in vitro*. *Cybernetics and Systems* **38**(7) (2007) 707–727
5. Luck, M., McBurney, P., Sheory, O., Willmott, S., eds.: *Agent Technology: Computing as Interaction*. University of Southampton (2005)
6. Helbing, D., Molnár, P.: Social force model for pedestrian dynamics. *Phys. Rev. E* **51**(5) (May 1995) 4282–4286
7. Schadschneider, A., Kirchner, A., Nishinari, K.: CA Approach to Collective Phenomena in Pedestrian Dynamics. In Bandini, S., Chopard, B., Tomassini, M., eds.: *Cellular Automata, 5th International Conference on Cellular Automata for Research and Industry, ACRI 2002*. Volume 2493 of *Lecture Notes in Computer Science*, Springer (2002) 239–248
8. Blue, V.J., Adler, J.L.: Cellular Automata Microsimulation of Bi-directional Pedestrian Flows. *Transportation Research Record* **1678** (2000) 135–141
9. Schadschneider, A., Klingsch, W., Klüpfel, H., Kretz, T., Rogsch, C., Seyfried, A.: Evacuation Dynamics: Empirical Results, Modeling and Applications. In Meyers, R.A., ed.: *Encyclopedia of Complexity and Systems Science*. Springer (2009) 3142–3176
10. Kuligowski, E.D., Gwynne, S.M.V.: The Need for Behavioral Theory in Evacuation Modeling. In: *Pedestrian and Evacuation Dynamics 2008*. Springer (2010) 721–732
11. Hall, E.T.: *The Hidden Dimension*. Anchor Books (1966)
12. Hall, E.T.: A System for the Notation of Proxemic behavior. *American Anthropologist* **65**(5) (1963) 1003–1026
13. Chattaraj, U., Seyfried, A., Chakroborty, P.: Comparison of Pedestrian Fundamental Diagram Across Cultures. *Advances in Complex Systems* **12**(3) (2009) 393–405
14. Canetti, E.: *Crowds and power*. Farrar, Straus and Giroux (1984)
15. Bandini, S., Manzoni, S., Redaelli, S.: Towards an Ontology for Crowds Description: a Proposal Based on Description Logic. In Umeo, H., Morishita, S., Nishinari, K., Komatsuzaki, T., Bandini, S., eds.: *ACRI. Volume 5191 of Lecture Notes in Computer Science*, Springer (2008) 538–541
16. Bandini, S., Manenti, L., Manzoni, S., Sartori, F.: A Knowledge-based Approach to Crowd Classification. In: *Proceedings of the The 5th International Conference on Pedestrian and Evacuation Dynamics, March 8-10, 2010*, (2010)
17. Moussaïd, M., Perozo, N., Garnier, S., Helbing, D., Theraulaz, G.: The Walking Behaviour of Pedestrian Social Groups and its Impact on Crowd Dynamics. *PLoS ONE* **5**(4) (04 2010) e10047
18. Sarmady, S., Haron, F., Talib, A.Z.H.: Modeling Groups of Pedestrians in Least Effort Crowd Movements Using Cellular Automata. In Al-Dabass, D., Triweko, R., Susanto, S., Abraham, A., eds.: *Asia International Conference on Modelling and Simulation, IEEE Computer Society* (2009) 520–525
19. Reynolds, C.W.: Flocks, Herds and Schools: a Distributed Behavioral Model. In: *SIGGRAPH '87: Proceedings of the 14th annual conference on Computer graphics and interactive techniques*, New York, NY, USA, ACM (1987) 25–34
20. Batty, M.: Agent-based pedestrian modelling. In: *Advanced Spatial Analysis: The CASA Book of GIS*. Esri Press (2003) 81–106

Validation of Agent-Based Simulation through Human Computation : An Example of Crowd Simulation

Xing Pengfei, Michael Lees, Hu Nan, and Vaisagh Viswanathan T.

School of Computer Engineering,
Nanyang Technological University,
Singapore, 639798
{pfxing,mhlees*}@ntu.edu.sg,
{huna0002, vaisagh1}@e.ntu.edu.sg

Abstract. Agent-based modeling as a methodology for understanding natural phenomena is becoming increasingly popular in many disciplines of scientific research. Validation is still a significant problem for agent-based modelers and while various validation methodologies have been proposed, none have been widely adopted. Data plays a key role in the validation of any simulation system, typically large amounts of observable real world data are necessary to compare with model outputs. However, the complex nature of the studied natural systems will often make data collection difficult. This is certainly true for crowd and egress simulation, where data is limited and difficult to collect. In this paper we propose a new technique for validation of agent-based models, particularly those which relate to human behavior. This methodology adopts ideas from the field of Human Computation as a means of collecting large amounts of contextual behavioral data. The key principle is to use games as a means of framing behavioral questions to try and capture people's natural and instinctive decisions. We outline some key design challenges for such games and present one example game in the form of Escape. Escape is an egress based game where people are tasked to escape from rooms inhabited by other people. We show some preliminary studies which highlight some interesting applications of the game in addressing validation of behavioral based crowd and egress simulation.

Keywords: Simulation, Validation, Games, Human Computation

1 Introduction

The term Complex System [9] has been used to describe systems, which are a formation of interconnected components that exhibit emergent complex behavior. The field of Agent-Based Modeling and Simulation (ABMS) is a popular field of modeling and simulation that is ideally suited to simulate such systems. An agent-based model is defined by a collection of interacting autonomous agents

with individually defined properties and behaviors. Through development of simple rules and the complex temporal interactions of the many agents, more complex macro level system properties will emerge. From a modeling perspective this is an attractive method for reasoning about complex systems since many forms of complex system the individual behaviors and characteristics are well understood, or easy to describe. Given this understanding it is clear to see why ABMS has seen a dramatic increase in its application in a wide variety of disciplines including: Sociology, Biology, and Economics. However, the major motivation for using ABMS to understand complex systems has an undesirable consequence when ascertaining the validity of models. One central motivation for ABMS is that the behavior of the system emerges in some unknown and unexpected way. This unpredictable nature of the system makes agent-based models very challenging to validate and current approaches and techniques are still under-developed with many philosophical and methodical questions remaining unanswered [14].

Validation is important for simulation; without some form of validation and verification there can only be limited confidence in the accuracy of the model. As explained, agent-based models pose unique challenges with regard to validation and these challenges become even more significant when trying to model natural systems which involve human behavioral aspects, such as crowd or egress simulation. For these types of model, challenges occur in many aspects of validation: collection of appropriate data, the non-determinism of human free will, etc. Establishing hard validation for such systems is very difficult [10, 14]; it is perhaps better to state that one gains confidence in the simulation by experimentation and comparison with observable data. Despite these problems associated with validity, simulations which consider human behavior still have their place, and are widely applied in areas of critical importance (e.g., epidemics [4], finance [12]). Due to the stochastic nature of these systems, both simulation replication and an abundance of related data are critical. Through repetition, results will have an associated confidence, which can help modellers understand the likelihood of the obtained result.

It is possible to achieve validation in a number of ways, and the suitability of approaches is often dependent on the motivation of modelling [3, 17]. Epstein [3] pointed out that the goal for models is not always prediction, but rather explanation. Epstein believes that in social science, models can be used to help guide data collection. In this way models can be meaningful in illuminating core dynamics. However, for ensuring predictive power, validation and data are important; modellers must be able to show that their model is capable of producing some known output data for a given set of known input data. Without this validation process it is hard to make any claims about the predictive power of a model. As well as validation, the data available from a physical system affects the experimental frame, the model resolution and model instantiation. This is what makes egress or crowd simulation so hard to validate and justify; data for emergency evacuation scenarios is very limited. The data that is available, is often at an aggregate level, e.g., number of people through exit, density of people, etc. More sophisticated techniques, such as video tracking, are still limited in

terms of accuracy and throughput of data collection. The use of engineered experiments also has significant limitations, recreating the stressful conditions of a real evacuation is very difficult under controlled conditions. Safety constrains the experimental conditions such that participants must be given prior knowledge of the conditions and therefore will not feel the same pressure or fear. Another key challenge is the non-deterministic nature of people’s behavior during egress. Therefore, to obtain confidence in the data, and hence validation, one must obtain multiple instances of the same data sets. Again, the difficulty of obtaining such data makes this validation an exigent task.

In this paper we propose a novel method for validation of agent-based evacuation and crowd simulation using concepts from the field of Human Computation [21]. We introduce our iPhone/iPad evacuation game which is intended to procure data and information, which can then be utilized in agent-based crowd and evacuation simulations we are currently developing [11]. The game is still in the early stages of development and as such the data collected is intended to be illustrative. The key contribution of this paper is therefore to introduce the concept of validation through human computation and to illustrate the usefulness and feasibility of such an approach. The remainder of this paper is organized as follows, section two introduces the simulation application area of crowd simulation and discusses what type of data is necessary for different simulation approaches and how it can be applied. Section three summarizes the field of Human Computation and provides a survey of significant and related work in the area. Section four continues with a detailed description of the evacuation game and its design. Section five outlines the three test cases used in this paper and presents some initial findings and the implications for crowd simulation models. The paper concludes with a summary and some possible directions for future work.

2 Related Work

Human crowds and their egress can and have been modelled in a number of ways; existing systems apply techniques such as flow modeling , particle systems, cellular automata and agent-based behavioral models. Each modeling approach has its own benefits along with its own drawbacks. A flow-based approach [19] works well under certain conditions where a homogeneous, optimal behavior is appropriate. This is typically in very high density locations where individual choice has little impact on the overall crowd motion or egress. The force-based particle systems [8] assume individuals (or particles) in the crowd exert forces on each other. Motion of individuals is then calculated by resolving the interacting forces on each individual agent. This approach allows for greater heterogeneity in the crowd, but interactions between individual must be described in terms of opposing forces, which can make it hard to quantify complex interactions and decisions. Cellular automata models [24] offer a great deal of heterogeneity, for example models for exit choice can consider exit size and exit density [5]. Many models focus on the physical aspects of crowd motion, such as density,

exit throughput and congestion. While these issues are clearly essential when analyzing egress or crowd motion, there are other, non-physical factors, which can greatly affect egress effectiveness. In [11] a cognitive agent framework was developed to provide tools for modeling human crowds. This framework incorporated many non-physical issues, for example: emotion, group formation, and gender and age variation. Other existing egress and crowd models have considered different cognitive aspects, in particular social norm [15], panic [15] and emotion [11]. While the arguments for this naturalistic approach seem convincing, the real challenge comes about in validating the approach and verifying its implementation.

Validation is typically application specific, every natural system will present unique challenges for validation. The ABMS community has identified specific issues associated with validating agent-based models. Some have correctly argued the need for a systematic approach and methodology for validating the models which are developed [10]. Others have discussed traditional (empirical) and alternative notions of validity and their implications in the context of agent-based modeling [14, 25]. From the perspective of behavioral simulation some have suggested expert based facial validation as one feasible approach [6]. This popular approach involves generating observable output from the simulation and asking a domain expert to assess the likelihood of the derived output from the given input. Sterman [18] proposed direct experimental validation as the most appropriate way to validate behavioral simulations. His arguments and approach closely mirror those which we propose in this paper. Sterman designed contextual games and scenarios in which to place subjects. These real life “role-play” scenarios, or interactive surveys, would be used to verify and test rules proposed within the simulation. The game would be controlled to present the subject with the precise scenarios in which the rules are intended to apply. The response can then be compared with the contextual rules, as specified in the model. While we agree with this approach to validating behavior based models, the approach does not scale well when the model requires the response of many participants. One of the distinct advantages which the human computation approach affords the experimenter, is the ability to conduct large scale interactive surveys. Having a ubiquitous game also allows data to be collected from many different countries with vastly different cultural and social behavior.

From the perspective of egress or crowd simulation there have been a number of attempts to validate various types of models. Statistical validation approaches generally look at macroscopic or aggregate properties of the real crowd and try to recreate the same properties within the model [1]. Obtaining data from a real crowd is also difficult, this can be achieved through manual counting, simple mechanical counting or through the use of video analysis [2]. Such approaches are valuable, but validation is still very context dependent and a model with various individual characteristics would need some way of knowing the approximate initial crowd composition in order to guarantee the model and system correspond. Facial validation is also one popular approach for validating virtual crowds, [16] uses a form of interactive survey whereby participants were placed

in a virtual world and surrounded by a computer generated crowd. Participants were then questioned about their experience and in particular their impressions of the artificial crowd.

3 Human Computation

Human computation is a method for utilizing humans as computational elements. In much the same way as other distributed computing platforms utilize idle clock cycles of a CPU (e.g., BOINC), human computation attempts to utilize idle human cycles through games. Typically the games will be carefully designed in order to make a repetitive mundane task interesting, which in turn will encourage people to play the game and unknowingly perform the task. Many researchers have collected fascinating statistics as to the number of man-hours spent on gaming around the world. As one example Jane McGonigal [13] has suggested the average American 21 year old will have spent 10000 hours of their lives playing games, which is approximately the same number of hours that same person will have spent within the education system. On May 23rd 2010 the Google main page was changed to include a version of the classic PacMan game, some estimations are that 4.82 million man hours of productivity were lost in a single day (approximately US\$120 million). The core motivation of Human Computation is to attempt to exploit this resource by designing games with some productive or positive side-effect.

As well as the massive amount of computational resources offered, humans are also capable of types of computation which machines cannot perform. One example of this is the ability to assign meaning to images through meta tagging, see (ESP Game and Peek-a-boom [22, 23]). The principle can be further refined to state that humans can be used to generate data by playing games, which can then be used by machines. In the ESP game for example, players are asked to textually identify contents of images and will succeed once they match the same word with their playing partner (who is randomly assigned via the internet based game). This way the humans, by playing the game, are generating image meta data (i.e., stating an image contains a car, house, grass, etc.) which can then be used by Google image search (or similar search engine). Current machine vision techniques are not capable of assigning such semantics to image contents, and so the alternative to such games would be to ask thousands of people to manually describe the contents of the billions of images on the internet. In this paper we adapt this concept to the validation of behavioral models, that is we propose developing games which will be used to collect instinctive behavioral responses, which can then be used to quantitatively validate behavioral models.

4 The Game : Escape

When designing games with a specific purpose, as is the case with human computation, there are a number of important considerations which must be accounted for [20]. From the perspective of simulation validation, there are further unique

design criteria which must be incorporated in the game design. During the design process we have considered the following principles:

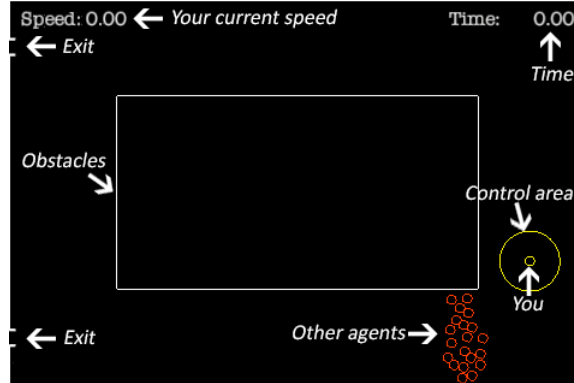
1. *Accuracy*: The game mechanics and artificially controlled players need to be designed carefully so as to closely reflect reality. We don't claim or strive for completely correct behavior for the artificial crowd in the game. What is important is that the player perceives the artificial crowd as human-like and therefore, makes the same decision in the game as they would in the equivalent real world scenario.
2. *Experimental Control*: In order to be able to draw conclusions from each of the scenarios, it is important that the experimenter is able to easily control conditions. If the experimenter wants to study the effects of one particular control on the players decision, the levels should be designed carefully to factor out as many other possible variables which may affect decisions.
3. *Ubiquity*: One fundamental benefit of the human computation approach for validation is the ability to collect large amounts of data. Therefore it is critical that the game receives the correct level of exposure and is played by as many people as possible.
4. *Playability*: Obviously it is important that the game is enjoyable to play so players will continue to play the game. This principle in fact contains all the standard game design fundamentals which have been developed over the past four decades of computer game development.

None of the above are trivial tasks and the development of a successful computer game is typically a slow and iterative process. The current incarnation of *Escape* has attempted to address the first three design principles of accuracy, experimental control and pervasiveness. The experiments conducted for this paper attempt to somewhat verify the accuracy of the game mechanics. It is important to note that the above principles will often lead to conflicting requirements for the game, which necessitate a number of trade-offs. For example, The choice of the iPhone as the platform was mainly for reasons of exposure and the ubiquity of the system. We decided against a more immersive, powerful 3D platform as this would be both harder to develop and would have lesser exposure, which in turn would result in fewer samples.

Our game, *Escape*, has been developed as an iPhone application using an open source game engine Cocos2D-iPhone. The game is still in the developmental stage and while all the game mechanics and level designer have been fully developed, the game still lacks a polished graphical front-end and storyline. We currently have two separate applications available: the *Escape* game itself and a convenient level editor for experimenters to quickly develop new scenarios with specific objectives in mind. The game itself consists of a single human agent controlled by the player and some number of competing AI controlled agents. Obstacles and walls are configured, as well as specific target locations which trigger level completion. In the level editor, level designers can simply click, drag and drop to design the shapes and positions of obstacles and goals (and sub-goals) and positions of agents.

The level designer is used in four distinct stages, firstly obstacles are drawn as a polygon from a set of user specified points. The designer must then specify a series of roadmaps, which are necessary for the RVO2 [7] motion planning system we adopt. In the third stage the designer specifies the AI controlled agents initial location and desired goals, and in the final stage the player agent is specified along with its goal. Once complete, the level specification is output to an XML file which can be loaded into the game or loaded back into the level editor if modification is required.

Fig. 1. The Escape game.



The game itself (see Figure 1) is essentially a time-stepped simulation where the AI agents are performing sense-think-act cycles within each time step. Obviously the hardware on even modern iPhones is somewhat limited, so the time step must be chosen appropriately to achieve the correct level of interactivity and playability - this is currently set to be 0.25 seconds. The AI controlled agents use the RVO2 motion planning system which has previously been applied to crowd simulation. The agents are $0.4m$ in diameter and have a series of specified waypoints, a preferred velocity for each agent is calculated using the current location of the AI agent and its current sub-goal. This preferred velocity is passed into the RVO2 calculation and an actual velocity is calculated, which will try to avoid any oncoming collisions with minimal deviation from the specified preferred velocity. The human player controls the player agent by drawing (i.e., dragging with a finger) a preferred velocity, the actual motion of the player agent is also governed by the RVO2 mechanism¹. Once a level is completed and the user agent escapes, a file is written to the phone which records the positions of all agents at every time step during that level. These files are then processed and analysed to obtain the desired results from the game.

¹ In theory any motion planning system could be employed within the game, we chose RVO2 in particular as it produces high quality motion with high computational efficiency

5 Test Scenarios

In this paper we present preliminary results from two test scenarios we feel are important for validating behavioral crowd simulation models. These tests have been conducted on a relatively small scale, with 25 participants. This preliminary study is intended to achieve a number of different objectives. Firstly, the results will offer an initial insight into the behavior of people in the chosen scenarios. Secondly, they should highlight any issues with the current game mechanics and allow the game to be improved prior to general release. Finally, the initial tests should enable refinement of the scenarios and test cases to ensure we are measuring the correct behavioral decisions.

Each participant is asked to play a total of 12 different level instances (6 of each type) with varying environmental configurations. The sample set are predominantly members of Parallel and Distributed Computing Centre (PDCC) within the school of Computer Engineering. Other participants have been sourced through word of mouth. All tests are conducted on an iPad running IOS 4.2.1. The game begins with an information entry screen where a participant is asked to enter their name, gender, age and place of birth. This information is stored within each output file generated during the game. Once the participant has entered their details, they are presented with a short game description outlining the purpose of the game and their objectives.

5.1 Density and Distance

One behavior which is common to many egress models is the trade-off made by people when opting for less crowded but longer routes. Obviously this trade-off is different for different people and any correct behavioral model should account for this, the difficulty is in understanding the average tradeoff made by an average person, or by a specific type of person. The scenario used here presents the player with two route choices R_1 and R_2 (See Figure 2(a))the experimenter can vary the relative distance and density of both routes, thereby obtaining the average behavior of a number of players for a given configuration.

Figure 2(a) shows the environmental configuration for this experiment. The player is presented with two exit choices involving a longer route (R_1) and a shorter route (R_2). Six different initial conditions of this scenario are presented to the player during the experiments. The number of agents occupying space on route R_2 is varied from zero to fifty in increments of ten. The computer controlled agents will always choose exit E_2 and are created within the area A (See Figure 2(a)). Given the initial configuration (and a fixed random seed) for each scenario it would be possible to calculate the optimal route choice in terms of escape time. However, one of the critical aspects of this experiment is that humans do not necessarily make optimal decisions in such scenarios, our experiments will hopefully offer some insight into exactly when and why humans choose the non-optimal route.

Figure 2(b) presents the summarized results from this test case. The results indicate even with a relatively small sample size we are achieving a definite tran-

sition in behavior. Once the route R_2 is occupied with more than 10 agents the majority of people will opt for the longer, but empty route R_1 . Further experiments could clarify the exact point at which the likelihood of choosing either of the two routes becomes the same. It would also be interesting to understand how the choice varies with a greater distinction in route length.

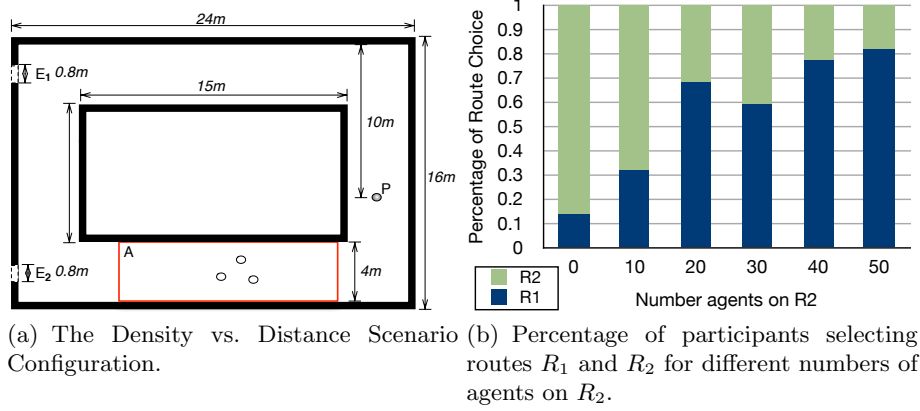


Fig. 2. Density and Distance Scenario

5.2 Exit Size

The second case study examines the effect that exit size has when an evacuee chooses their exit. The literature contains many examples of behavioral rules which use exit size, exit throughput and exit density (e.g., [5]). However, the assumption that a human can quickly and accurately calculate exit throughput and density is perhaps incorrect. This game scenario is intended to investigate the hypothesis that in fact exit size alone can be the determining factor when making an exit choice.

Figure 3(a) depicts the environment used in this case study, the player is placed in the centre of the room with two exits at equal distance. The environment is populated with 50 other AI controlled agents placed uniformly in the circular area indicated by A in Figure 3(a). The AI agents will attempt to leave the environment through either of the two exits. There are again six different instances of this case study (presented as six different levels), with the relative number of agents choosing each exit varying in each instance. The larger exit is chosen to be twice the size of the smaller exit and the probability of an AI agent choosing the small exit (E_1) varied between 0, 0.2, 0.4, 0.6, 0.8 and 1.0.

Figure 3(b) presents the summarized results from this test case. As expected, it seems that in most cases the participants will opt for the larger of the two exits regardless of the relative proportion of AI agents choosing the small exit. However, a slight increase in those players choosing E_2 can be observed as the number of AI agents choosing E_2 decreases. For the case when all agents choose exit E_1 the results show a surprising change in behavior. From experimental

observation we realize this is due to the level specification. In this case the player is pushed toward the smaller exit by the other agents who are all moving to E_1 . The player typically attempts to move towards the larger exit, but due to pushing, eventually submits to the will of the crowd and chooses the smaller exit. This result emphasizes the need for careful scenario testing to ensure they are correctly designed to present a valid set of choices to the player.

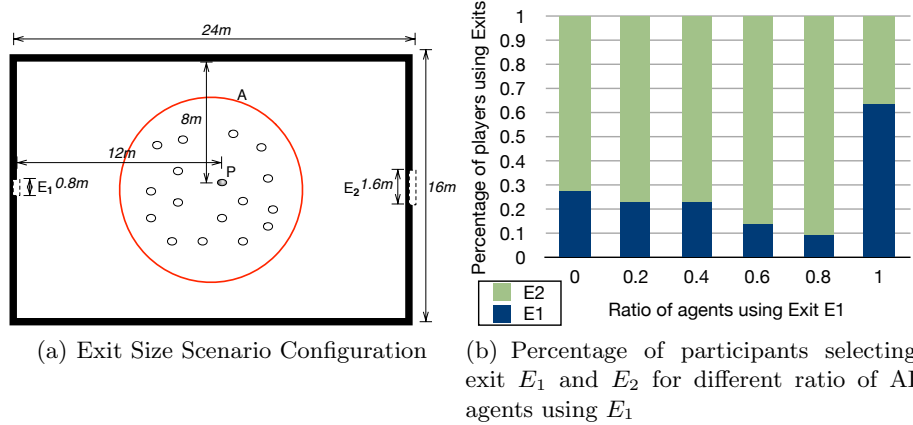


Fig. 3. Exit Size Scenario

6 Discussion and Future Work

In this paper we have proposed a new approach to validating agent-based models which uses principles from the field of Human Computation. This approach is well suited to validating human behavior models which are known to present unique challenges in terms of validation. We do not see this methodology as a single technique for achieving validation for all forms of agent-based models or behavioral models. Instead we view it as a complementary approach to existing methods and another technique modelers can utilize when building confidence in their models. Importantly, our technique has the key advantage of being able to generate large sets of data from a large sample set. The approach also offers the advantages outlined by Sterman [18], in that using carefully designed scenarios and games provides a more contextual and interactive form for phrasing the question. This paper also presented the game Escape as an example of using human computation as a way to validate behavioral crowd simulation. Our assumption is that the simple 2D interface is capable of capturing the same reactive decisions that people make in reality. This may be a strong assumption and is something which we plan to investigate in future work. One can argue that this is equivalent to other forms of facial validation where experts are asked to observe 2D animation output from the crowd simulation. Initial results indicate the feasibility of the approach, but also highlight the challenges in designing

appropriate scenarios. We believe that eventually the resultant focused and statistical data can be applied in the development, parameterization and validation of decision models and route algorithms for crowd simulation.

There are still a number of issues we hope to address in our continued work. Firstly, Escape will be completed and a game story will be constructed around the current basic game mechanics. Once the graphics and story have been completed, we plan for further testing before placing the game on general release. We still have to clarify issues of data collection from phones prior to release and methods for distributing new game levels and scenarios. Game controls may also be developed to allow for other scenario types. For example, in all current scenarios the player has complete information regarding the environment, we plan to develop scenarios with limited sensor range to test other possible decision cases. One important question is the accuracy of the game and the correctness of the behavioral data that is collected. It is important that the game itself is validated against real world data. We are currently looking at the use of video or RFID tracking with real world experiments as two possible methods for validating Escape. Finally, we have not shown explicitly how to use this data, or the game, for validating an agent-based model. In future work we plan to functionally validate a group-based perception model we are developing.

7 Acknowledgments

This research has been funded by the NTU SU Grant M58020019

References

1. B. Banerjee and L. Kraemer. Validation of agent based crowd egress simulation. In *Proceedings of the 9th International Conference on Autonomous Agents and Multi-agent Systems: volume 1 - Volume 1*, AAMAS '10, pages 1551–1552, Richland, SC, 2010. International Foundation for Autonomous Agents and Multiagent Systems.
2. J. Berrou, J. Beecham, P. Quaglia, M. Kagarlis, and A. Gerodimos. Calibration and validation of the legion simulation model using empirical data. In N. Walda, P. Gattermann, H. Knoflach, and M. Schreckenberg, editors, *Pedestrian and Evacuation Dynamics 2005*, pages 167–181. Springer Berlin Heidelberg, 2007.
3. J. M. Epstein. Why model? *Journal of Artificial Societies and Social Simulation*, 11(4):12, 2008.
4. N. M. Ferguson, D. A. T. Cummings, C. Fraser, J. C. Cajka, P. C. Cooley, and D. S. Burke. Strategies for mitigating an influenza pandemic. *Nature*, 442(7101):448–452, April 2006.
5. A. Ferscha and K. Zia. On the efficiency of lifebelt based crowd evacuation. In *Proceedings of the 2009 13th IEEE/ACM International Symposium on Distributed Simulation and Real Time Applications*, DS-RT '09, pages 13–20, Washington, DC, USA, 2009. IEEE Computer Society.
6. A. J. Gonzalez. Validation of human behavioral models. In *Proceedings of the Twelfth International Florida Artificial Intelligence Research Society Conference*, pages 489–493. AAAI Press, 1999.

7. S. J. Guy, J. Chhugani, S. Curtis, P. Dubey, M. Lin, and D. Manocha. Pedestrians: a least-effort approach to crowd simulation. In *Proceedings of the 2010 ACM SIGGRAPH/Eurographics Symposium on Computer Animation*, SCA '10, pages 119–128, Aire-la-Ville, Switzerland, Switzerland, 2010. Eurographics Association.
8. D. Helbing and P. Molnár. Social force model for pedestrian dynamics. *Physical Review E*, 51(5):4282–4286, May 1995.
9. N. R. Jennings. An agent-based approach for building complex software systems. *Commun. ACM*, 44:35–41, April 2001.
10. F. Klügl. A validation methodology for agent-based simulations. In *Proceedings of the 2008 ACM symposium on Applied computing*, SAC '08, pages 39–43, New York, NY, USA, 2008. ACM.
11. L. Luo, S. Zhou, W. Cai, M. Y.-H. Low, and M. Lees. Toward a generic framework for modeling human behaviors in crowd simulation. In *IAT*, pages 275–278. IEEE, 2009.
12. T. Lux and M. Marchesi. Scaling and criticality in a stochastic multi-agent model of a financial market. *Nature*, 397(6719):498–500, February 1999.
13. J. McGonigal. Gaming can make a better world. TED online, February 2010. http://www.ted.com/talks/jane_mcgonigal_gaming_can_make_a_better_world.html.
14. S. Moss. Alternative approaches to the empirical validation of agent-based models. *Journal of Artificial Societies and Social Simulation*, 11(1):5, 2008.
15. X. Pan, C. Han, K. Dauber, and K. Law. A multi-agent based framework for the simulation of human and social behaviors during emergency evacuations. *AI & Society*, 22:113–132, 2007. 10.1007/s00146-007-0126-1.
16. N. Pelechano, C. Stocker, J. M. Allbeck, and N. I. Badler. Being a part of the crowd: towards validating vr crowds using presence. In *AAMAS (1)'08*, pages 136–142, 2008.
17. P. Sloot, P. Coveney, and J. Dongarra. Editorial board. *Journal of Computational Science*, 1(4):CO2 – CO2, 2010.
18. J. Sterman. Testing behavioral simulation models by direct experiment. Working papers 1752-86., Massachusetts Institute of Technology (MIT), Sloan School of Management, 1986.
19. A. Treuille, S. Cooper, and Z. Popović. Continuum crowds. In *SIGGRAPH '06: ACM SIGGRAPH 2006 Papers*, pages 1160–1168, New York, NY, USA, 2006. ACM.
20. L. von Ahn. Games with a purpose. *Computer*, 39:92–94, 2006.
21. L. von Ahn. Human computation. In *Proceedings of the 4th international conference on Knowledge capture*, K-CAP '07, pages 5–6, New York, NY, USA, 2007. ACM.
22. L. von Ahn and L. Dabbish. Labeling images with a computer game. In *Proceedings of the SIGCHI conference on Human factors in computing systems*, CHI '04, pages 319–326, New York, NY, USA, 2004. ACM.
23. L. von Ahn, R. Liu, and M. Blum. Peekaboom: a game for locating objects in images. In *Proceedings of the SIGCHI conference on Human Factors in computing systems*, CHI '06, pages 55–64, New York, NY, USA, 2006. ACM.
24. Y. Weifeng and T. K. Hai. A novel algorithm of simulating multi-velocity evacuation based on cellular automata modeling and tenability condition. *Physica A: Statistical Mechanics and its Applications*, 379(1):250 – 262, 2007.
25. P. Windrum, G. Fagiolo, and A. Moneta. Empirical validation of agent-based models: Alternatives and prospects. *Journal of Artificial Societies and Social Simulation*, 10(2):8, 2007.

Observation of large-scale multi-agent based simulations

Gildas Morvan^{1,2}, Alexandre Veremme^{1,3}, and Daniel Dupont^{1,3}

¹ Univ. Lille Nord de France, 1bis rue Georges Lefèvre 59044 Lille cedex, France

² LGI2A, U. Artois, Technoparc Futura 62400 Béthune, France. email: `first name.surname@univ-artois.fr`

³ HEI, 13 rue de Toul 59046 Lille Cedex, France. email: `first name.surname@hei.fr`

Abstract. The computational cost of large-scale multi-agent based simulations (MABS) can be extremely important, especially if simulations have to be monitored for validation purposes. In this paper, two methods, based on self-observation and statistical survey theory, are introduced in order to optimize the computation of observations in MABS. An empirical comparison of the computational cost of these methods is performed on a toy problem.

Keywords: large-scale multi-agent based simulations, observation methods, scalability

1 Introduction

Theoretical and practical advances in the field of multi-agent based simulations (MABS) allow modelers to simulate very complex systems to solve real world problems. However the analysis and validation of simulations remain engineering problems that do not have "turnkey" solutions. Thus, MABS users conducting such tasks face two main issues:

1. define validation metrics for the simulation,
2. compute efficiently the metrics.

The first issue is generally solved by constructing a set of *ad-hoc* qualitative or quantitative rules on simulation properties. To evaluate these rules, it is then mandatory to observe the corresponding simulation properties and thus to consider the second issue. A distinctive characteristic of MABS is that global simulation properties are not necessary directly observable: they may need to be computed from local agent properties. Fortunately, most of modern MABS platforms come with observation frameworks and toolboxes. Basically, three types of observation methods are generally available [9]:

1. interactive observation: users select observed properties during simulations, *e.g.*, using a point and click interface,

2. brute-force direct observation: simulation agents sharing a given property are monitored; agent properties are then aggregated by a so-called *observer* agent that computes the observation (fig. 1),
3. indirect observation: the observed property is inferred from the observable consequences of agent actions, *e.g.*, in the environment.

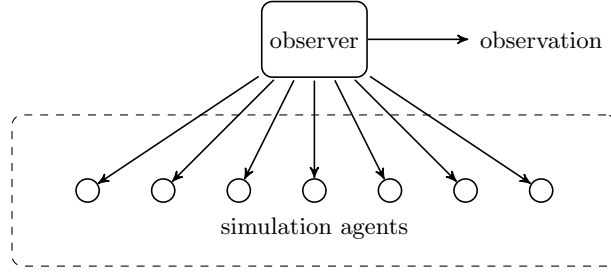


Fig. 1. Brute-force direct observation method

While the first method is clearly unadapted to the observation of large-scale or batch simulations, the second has an important computational cost. This issue is illustrated with a simple case study inspired by "StupidModel" [8]: N agents move randomly in a two dimensional environment $\mathcal{E}$ discretized into $100 \cdot 100$ square cells with Moore neighbourhood during 1000 steps. An area $\mathcal{Z} \subseteq \mathcal{E}$ is defined. The number of agents Z in the area $\mathcal{Z}$ is observed at each simulation step. This simulation is implemented on the MadKit/TurtleKit platform⁴ [7]. Figure 2 shows the CPU times needed to compute unobserved and observed simulations on a Dell Precision 650 workstation⁵, using indirect and brute-force direct observation methods, for the given expected value $E(Z) = N/5$, as a function of the number of simulation agents N .

This work is based on existing implementations (by MadKit/TurtleKit and MASON) of the direct observation method. Thus, an empirical computational complexity metric, *i.e.*, the CPU time needed to compute simulations, is used. This metric, denoted $\mathcal{C}$, depends, in our case study, on the number of simulation agents, N , and on the expected value of the cardinal of the subset of simulation agents computed by *filter*, $E(Z)$.

These results show that, in this case, indirect observation has a minor impact on the computational cost of the simulations. Kaminka *et al.* also note that this method is not intrusive: simulation agents do not have to be modified or

⁴ All the simulations and observation methods described in this paper have also been implemented on the MASON platform [6], leading to similar results.

⁵ CPU: 2×3.06 GHz Intel XeonTM, RAM: 4×1 GB. Full specification: http://www.dell.com/downloads/emea/products/precn/precn_650_uk.pdf. All the results presented in this paper have been computed by this machine.

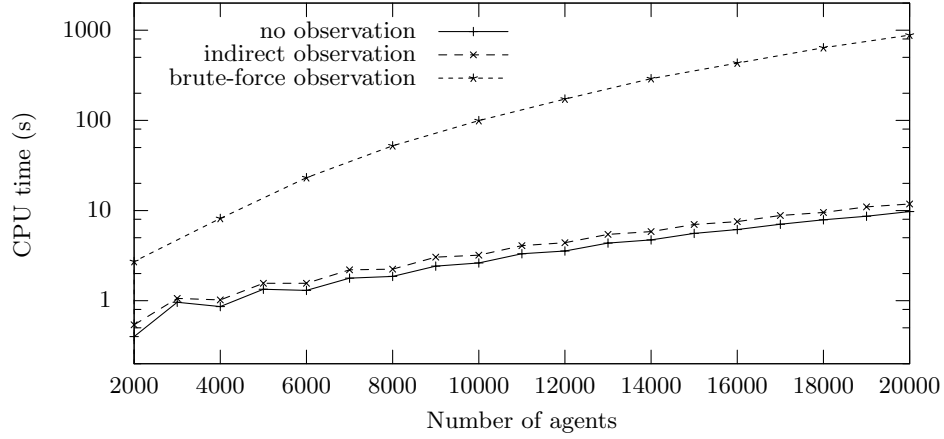


Fig. 2. CPU times (log scale) needed to compute observed and unobserved simulations for $E(Z) = N/5$

accessed during the observation process [5]. However, Wilkins *et al.* underline that the applicability of this method is limited: the observed property might not be inferred [12]; it is often true in complex, *i.e.*, in most of the real world, cases. Thus, direct observation often remains the only available option.

The brute-force direct observation method is, for this particular and very simple toy-problem, linear in the number of simulation agents⁶, *i.e.*, $\mathcal{C}(obs) \propto N$, while the model is exponential, *i.e.*, $\mathcal{C}(model) \propto \alpha^N, \alpha > 1$. However, complex simulations generally involve non-linear observation problems [11]. Thus, improving direct observation appears to be a good lead to improve the efficiency of large-scale complex MABS.

In this paper, two non-brute-force direct observation methods, based on self-observation and statistical survey theory are introduced. An empirical comparison of the computational cost of these methods is performed and discussed on the presented case study.

2 Filtrated direct observation of MABS

Basically, there are two ways to compute a direct observation:

1. a set of agents $\mathcal{A}$ (generally all the simulation agents that share the properties that have to be observed), statically defined, is probed by an observer agent,
2. a subset $\mathcal{A}'$ of $\mathcal{A}$, computed at runtime, is probed (fig. 3).

⁶ Authors would like to thank anonymous reviewers for raising this issue.

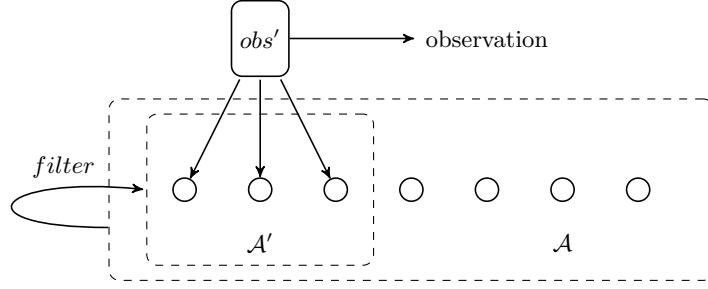


Fig. 3. Filtrated direct observation

Formally, we consider an observation function

$$obs : 2^{\mathcal{A}} \rightarrow \mathcal{I}, \quad (1)$$

where $\mathcal{A}$ is a set of agents and $\mathcal{I}$ represents the values that can be observed. We define $\mathcal{A}' \subseteq \mathcal{A}$ as the minimal subset of agents, dynamically defined by a set of constraints (*e.g.*, in the case study, a unique constraint related to the position of the agent), needed to compute obs correctly. In other words, $\mathcal{A}'$ is a set of agents such as

$$\begin{aligned} & obs'(\mathcal{A}') = obs(\mathcal{A}) \text{ and} \\ & \nexists \mathcal{A}'' \subset \mathcal{A}' \mid obs'(\mathcal{A}'') = obs(\mathcal{A}), \end{aligned} \quad (2)$$

where obs' is a "simplified" observation function, *i.e.*, that can be computed faster than obs because it is specific to $\mathcal{A}'$:

$$\forall \mathcal{A}'' \neq \mathcal{A}' \subseteq \mathcal{A}, \quad obs'(\mathcal{A}') = obs(\mathcal{A}'). \quad (3)$$

Thus, in the case study, the observation function $obs(\mathcal{A})$ observes the position of each agent in $\mathcal{A}$ to count only the ones that are situated in $\mathcal{Z}$, while an observation function $obs'(\mathcal{A}')$ would only returns $|\mathcal{A}'|$ because the agents of $\mathcal{A}'$ are by definition situated in $\mathcal{Z}$.

To use obs' , it is necessary to consider a filtering function $filter$ able to identify the subset $\mathcal{A}'$ (in the case study, this subset only contains the agents situated in $\mathcal{Z}$):

$$\begin{aligned} & filter : 2^{\mathcal{A}} \rightarrow 2^{\mathcal{A}} \mid \forall \mathcal{A}', \mathcal{A}'' \subseteq \mathcal{A}, \\ & \text{if } filter(\mathcal{A}') = \mathcal{A}'', \text{ then } \mathcal{A}'' \subseteq \mathcal{A}'. \end{aligned} \quad (4)$$

The goal is to define and implement a filtering function, such as the cost of the observation computation is reduced, *i.e.*,

$$\mathcal{C}(obs'(filter(\mathcal{A}))) < \mathcal{C}(obs(\mathcal{A})). \quad (5)$$

In the following section, two different implementations of this idea are presented.

3 Implementation of filtrated direct observation methods

3.1 Self-observation

The core idea of this method is to implement the filtering function in the simulation agents themselves. Then, using an organizational structure, denoted *group*, allows to identify the set of agents that has to be observed. Thus, a group is defined as *the set of agents that contains the sufficient and necessary information to compute an observation*. In other words, a group defines, for a given observation, the minimal set of agents that is mandatory to compute it. Agents observe themselves to determine if they have to join, leave or stay in a group. A filtering function is defined by a set of rules R , that specifies the conditions under which an agent has to be observed, evaluated at each simulation step (fig. 4).

Thus, in our case study, we consider a group G , that contains the agents situated in $\mathcal{Z}$. The following set of rules, defined here in natural language, is associated to each agent:

- if the agent is in $\mathcal{Z}$ and does not belong to G , then the agent joins G ,
- if the agent is not in $\mathcal{Z}$ and belongs to G , then the agent leaves G .

The observation system (obs') only probes the agents of G . Figure 5 presents the CPU time difference between simulations observed with self-observation and brute-force methods as a function of the number of agents in the simulation, N , and the mean rate of observed agents, $E(\mathcal{Z})/N$. The dashed line represents the isoline 0, *i.e.*, the conditions for which there is no difference between the two methods. Thus, the area below this line maps the cases for which self-observation is faster.

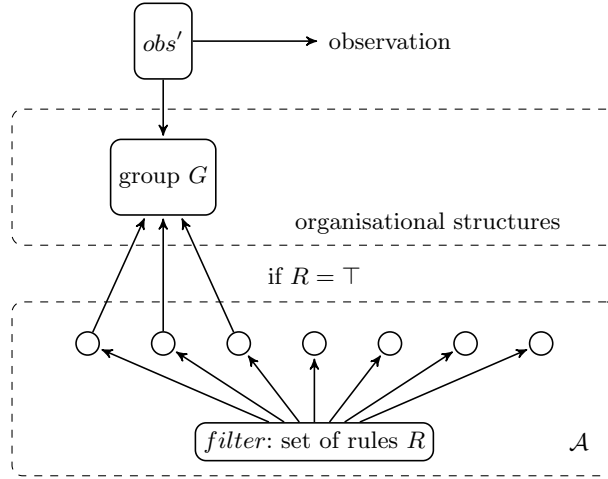


Fig. 4. The self-observation based method

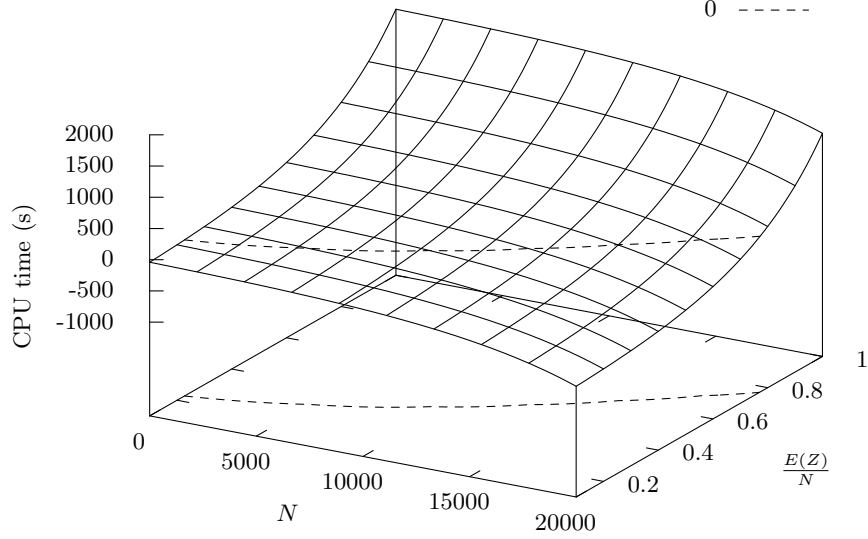


Fig. 5. CPU time difference between simulations observed with self-observation and brute-force methods as a function of the number of agents in the simulation, N , and the mean rate of observed agents, $E(Z)/N$ (response surface estimation)

3.2 Statistical survey

If the equation 2 is rewritten as follows:

$$\begin{aligned} & obs(\mathcal{A}') \simeq obs(\mathcal{A}) \text{ and} \\ & \nexists \mathcal{A}'' \subset \mathcal{A}' \mid obs(\mathcal{A}'') \simeq obs(\mathcal{A}), \end{aligned} \quad (6)$$

i.e., if imprecise observations are authorized, it becomes possible to filter the observed population on a statistical basis. As a result, we do not consider a specific observation function anymore as the set of agents returned by the filtering function is not necessary the *set of agents that contains the sufficient and necessary information to compute the observation*.

Statistical survey theory⁷ provides a formal ground to determine optimal sampling method and size of observed population sample.

Let once again consider our case study; we denote n the size of the observed population, randomly sampled at each simulation step. An estimator of Z , denoted $\hat{Z}$ is constructed from this sample. Many estimator definitions can be found in the literature. In this case, as the population of simulation agents is homogeneous with respect to $E(Z)$ (all the agents have the same probability to

⁷ Proofs of statistical survey theory results presented in this paper will not be given. Interested readers may refer to [1] for an exhaustive presentation of sampling designs, estimator construction and variability estimation methods.

be in $\mathcal{Z}$), n is determined with the Horvitz-Thompson estimator [4] :

$$n^{-1} = \frac{d^2}{4S^2} + \frac{1}{N}, \quad (7)$$

where d is the maximal absolute error accepted for the observation and

$$S^2 \simeq \left(1 - \frac{E(Z)}{N}\right) \cdot \frac{E(Z)}{N}. \quad (8)$$

Impact on the computational cost is shown in figure 6. The semantic is the same than figure 5: the left area maps the cases for which the statistical survey based method is faster than brute-force method.

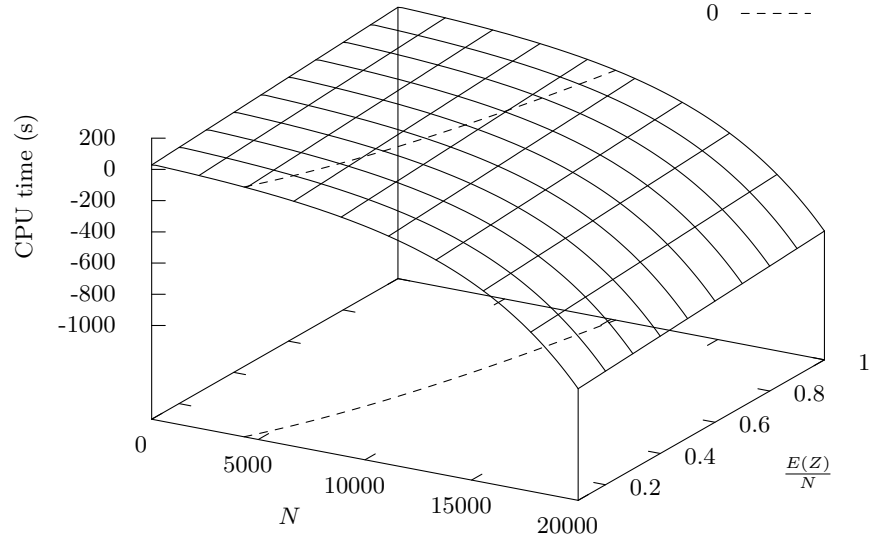


Fig. 6. CPU time difference between simulations observed with statistical survey ($d = 0.08$) and brute-force methods as a function of the number of agents in the simulation, N , and the mean rate of observed agents, $E(Z)/N$ (response surface estimation)

3.3 Discussion

Figure 7 sums up the previous results qualitatively: conditions for which it is preferable to use one method over another are identified. These results are specific to our case study and its implementation; however, they highlight that the choice of an observation method is not trivial and that the performance of the different available methods should be analyzed on a set of simulations before using the model in a production context.

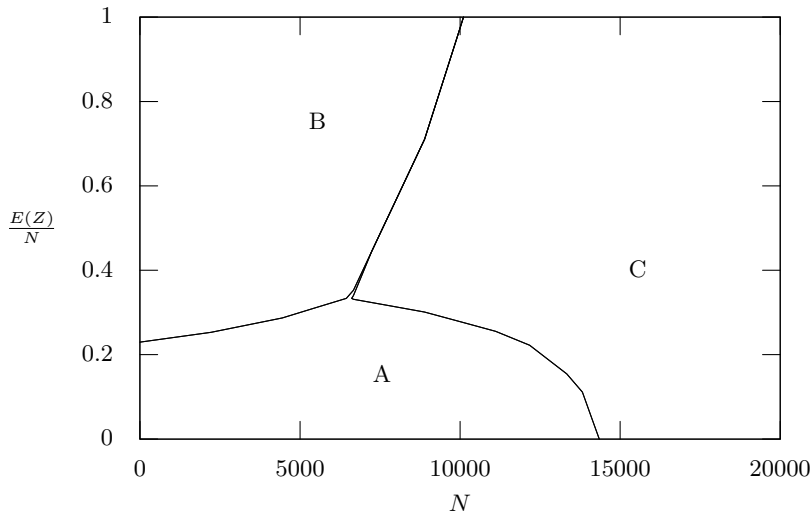


Fig. 7. Map of the fastest observation methods (response surface estimation); A: self-observation, B: brute-force, C: statistical survey ($d = 0.008$)

In a given context, knowing the map of the fastest observation methods allows to dynamically adapt the observation method to use the most efficient one. Impact of dynamic adaptation of the observation method on CPU time is presented in the context of the first example (cf. fig. 2) in figure 8.

4 Conclusion and perspectives

Observation methods presented in this paper allow, under specific conditions identified on a simple case study, to reduce significantly the computational cost of MABS composed of numerous agents. However, filtering is not the only option. Considering that imprecise observations are acceptable, while a precision level is guaranteed, the optimal observation frequency could be determined from the observed property variation. Roughly, the more the variation, the more the observation frequency. However, MABS are used to simulate complex systems with nonlinear dynamics. Dynamic adaptation of observation frequency could be an interesting lead to reduce MABS computational cost.

Moreover, in the statistical survey based method (cf. section 3.2), we consider a simple random sampling method, assuming the population is homogeneous. Real world MABS often involve heterogeneous agents for which the distribution of observed individual properties is not uniform. Clever sampling methods, *e.g.*, a stratified random sampling approach, should then be used. In very complex cases, a machine learning system should be implemented to analyze the impact of sampling method properties on the observation quality and computational

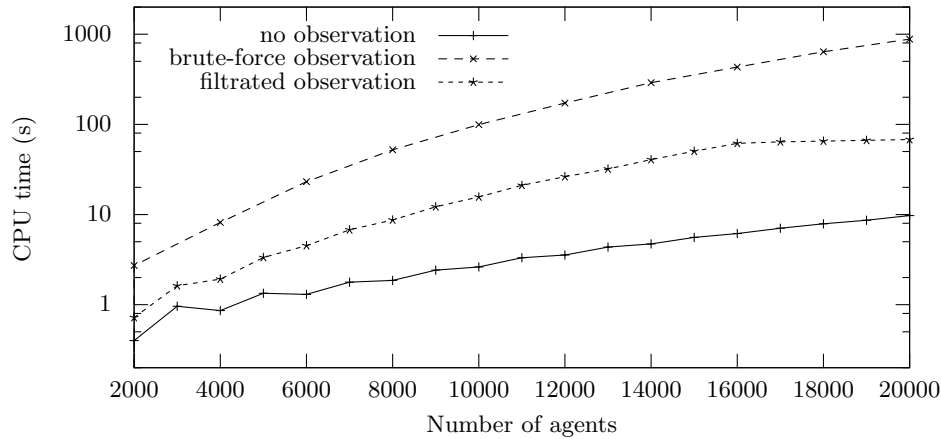


Fig. 8. CPU times (log scale) needed to compute observed and unobserved simulations for $E(Z) = N/5$, filtrated observation method being dynamically adapted to the context

cost, and determine the optimal ones. Similarly, the organizational model used to implement the self-observation based method is very simple: the only organizational structure that is defined is the "group". Using a more comprehensive one, *e.g.*, AGR [2], would allow to consider very complex and fine observations.

From a methodological point of view, authors experimented that setting up an observation method, generally improves the design of simulation validity metrics. Indeed, it forces simulation designers and users to explicitly define local and global observed properties and their sufficient and necessary conditions of observability, and then the validity constraints over them.

While this paper focuses on reducing the complexity of observation, many published works concentrated on agent interactions by dynamically scaling up and down simulated entities or using more structured interaction artifacts [3, 10]. Together, these approaches should lead to the conception of highly efficient large-scale MABS simulators.

References

1. Bethlehem, J.: Applied Survey Methods: A Statistical Perspective. Wiley (2009)
2. Ferber, J., Gutknecht, O.: A meta-model for the analysis and design of organizations in multi-agent systems. In: Proceedings of the Third International Conference on Multi-Agent Systems (ICMAS98). pp. 128–135 (1998)
3. Gaud, N., Galland, S., Gechter, F., Hilaire, V., Koukam, A.: Holonic multilevel simulation of complex systems : Application to real-time pedestrians simulation in virtual urban environment. Simulation Modelling Practice and Theory 16, 1659–1676 (2008)

4. Horvitz, D., Thompson, D.: A generalization of sampling without replacement from a finite universe. *Journal of the American Statistical Association* 47, 663–685 (1952)
5. Kaminka, G., Pynadath, D., Tambe, M.: Monitoring teams by overhearing: A multi-agent planrecognition approach. *Journal of Artificial Intelligence* 17, 83–135 (2002)
6. Luke, S., Cioffi-Revilla, C., Panait, L., Sullivan, K., Balan, G.: Mason: A multiagent simulation environment. *Simulation* 81(7), 517–527 (2005)
7. Michel, F., Beurier, G., Ferber, J.: The turtlekit simulation platform: Application to complex systems. In: *Proceedings of the First International Conference on Signal-Image Technology and Internet Based Systems*. pp. 122–127 (2005)
8. Railsback, S., Lytinen, S., Grimm, V.: Stupidmodel and extensions: A template and teaching tool for agent-based modeling platforms. <http://condor.depaul.edu/slytinen/abm/StupidModelFormulation.pdf> (2005)
9. Railsback, S., Lytinen, S., Jackson, S.: Agent-based simulation platforms: Review and development recommendations. *Simulation* 82(9), 609–623 (2006)
10. Razavi, S., Gaud, N., Mozayani, N., Koukam, A.: Multi-agent based simulations using fast multipole method: application to large scale simulations of flocking dynamical systems. *Artificial Intelligence Review* 35(1), 53–72 (2011)
11. Veremme, A., Lefevre, E., Morvan, G., Jolly, D.: Application of the belief function theory to validate multi-agent based simulations. In: *First international workshop on the theory of belief functions* (2010)
12. Wilkins, D., Lee, T., Berry, P.: Interactive execution monitoring of agent teams. *Journal of Artificial Intelligence Research* 18, 217–261 (2003)

Modeling Emotional Contagion

Jason Tsai¹, Emma Bowring², Stacy Marsella³, and Milind Tambe¹

¹ University of Southern California, Los Angeles, CA 90089
{jasontts, tambe}@usc.edu

² University of the Pacific, Stockton, CA 95211
ebowring@pacific.edu

³ USC Institute for Creative Technologies, Playa Vista, CA 90094
marsella@ict.usc.edu

Abstract. In social psychology, emotional contagion describes the widely observed phenomenon of one person’s emotions being influenced by surrounding people’s emotions. While the overall effect is agreed upon, the underlying mechanism of the spread of emotions has seen little quantification and application to computational agents.

In this paper, we explore computational models of emotional contagion by implementing two models (Bosse et al., Durupinar et al.) and augmenting them to better model real world observations. Our additions include examining the impact of physical proximity and authority figures. We show that these additions provide substantial improvements to the qualitative trends of emotion spreading, more in line with expectations than either of the two previous models. We also evaluate their impact on evacuation safety in an evacuation simulation, ESCAPES, showing substantial differences in predicted safety based on the contagion model.

1 Introduction

Emotional contagion has been shown to arise in a wide range of scenarios in everyday life. Its effects are felt in homes everyday when comedic shows employ laugh tracks to elicit stronger emotional responses from audiences. Less often, but with far more severe implications, it is also felt during the spread of fear and anxiety that surrounds any crowd-based disaster. With the growing interest in emotional modeling in agents, the contagion of these emotions can no longer be marginalized when modeling crowds. Recent work has sought to quantify the qualitative findings of social psychology into useable models with varying degrees of success. Bosse et al. (VU University) introduced one of these in 2009 [3] that used an interaction-based model derived directly from social psychology theories of emotional contagion wherein members of the simulation converged towards a weighted average of each emotional type. Durupinar et al. [6] used an epidemiological-style threshold-based model wherein successive interactions with emotionally ‘infected’ people raises the chance of infection.

While both of these models showed predictions in line with some qualitative findings in social psychology studies, they are inherently very different models of the same phenomenon. Although this type of detail may not be important for understanding the contagion of joy with the use of laugh tracks, it may offer substantially different predictions on the outcome of emotionally-charged crowd simulations. It is in this context

that we explore the modeling of emotional contagion. In particular, we use an evacuation simulation called ESCAPES, described more in Section 3, as the test bed for different models of fear contagion.

Despite their promise, both models come short when applied to an evacuation simulation such as ESCAPES. First, the models do not explore the impact of proximity on the effect of contagion. The VU model provides a parameter (channel strength) that can allow for this manipulation, but never provides guidance on how it should be done. Second, the authors have not introduced guidelines for designing ‘special’ agents such as authority figures that might have stronger resistance to fear contagion and be trained to reduce fear in other agents. Again the VU model provides parameters that allow for this manipulation (receiver openness, sender expressiveness), but have not explored this in their work thus far. The Durupinar model provides no mechanism for either feature.

We propose to augment the VU and Durupinar models with proximity-based effects and authority figure calming and examine their performance in the context of ESCAPES. Through extensive experimental results, we show that without proximity’s effect on contagion, neither model produces qualitatively believable results. After incorporating authority calming effects into each model, we show that the spread of emotions through the population again changes drastically. Finally, as a second-order effect, we also show that the evacuation simulation’s outcome predictions vary substantially, motivating the need for an accurate model of contagion and authority effects.

2 Related Work

Seminal works in social psychology first began the discussion around emotional contagion. In particular, Hatfield et al. [7] first codified the observed phenomena that were just beginning to receive researcher attention. Follow-up work by the co-authors as well as in related fields such as Barsade et al. [1] in managerial sciences continued to detail the effects of the phenomenon in new domains. Recently, there have been attempts to begin quantifying emotional contagion and explore cross-cultural variations in attributes that effect emotional contagion [5, 9].

From a computational perspective, the previously mentioned work of Bosse et al. (VU model) and Durupinar et al. are two of the most recent models of emotional contagion upon which a few follow-up works have been based [2, 8].

3 ESCAPES

Although not the focal point, the ESCAPES evacuation simulation [10] serves as the test bed for our models of emotional contagion, so we describe it briefly here. ESCAPES focuses on the features identified by experts that particularly effect airport evacuations, including first time visitors’ incomplete knowledge of the area, the presence of families, and the presence and effects of authority figures [4]. We also model fear and model its impact on behavior by increasing the speed of more fearful agents.

Although the second-order impacts of changing the emotional contagion model, such as evacuation rates and safety, are dependent upon the specific simulation implementation, we use ESCAPES as a test bed to illustrate obvious deficiencies in the base

models that would occur in any simulation with a spatial component. We detail these in Sections 8 and 9 and also highlight second-order impacts that the different models have upon the evacuation as a whole.

For all the experiments discussed in Section 8 and 9, the same scenario was used (spatial layout can be seen in Figure 3) and 30 trials were run for each setting. It features 2 large spaces, each with an exit, connected by hallways which are lined with smaller spaces that represent shops. 15 seconds into the simulation, an event occurs at the center of the scenario, inciting fear (0.75 for nearby agents, 0.1 for further agents) and a need to evacuate that is communicated by authority figures to pedestrians. The scenario features 100 normal pedestrians, including 10 families of 4 each, as well as 10 authority figures that patrol the scenario.

4 VU Model

Introduced in 2009 by Bosse et al. [3] and built upon in [2, 8], this model is an independent interaction-based model. The initial version that we use here moves people towards a weighted-average of the group’s emotional levels. Since subsequent works do not address the needs of our evacuation simulation, we begin with a discussion of the original model and its attributes. In Sections 6 and 7 we will further explore this base model to examine the impact of proximity and a particular model of authority figures.

We briefly mention the primary components of interest in the VU model here. In particular, emotional contagion is modeled using 5 parameters for every pair of people that may interact: level of sender’s emotion q_S , level of receiver’s emotion q_R , sender’s expressiveness ϵ_S , receiver’s openness δ_R , and the channel strength between S and R α_{SR} . All values are numbers in the interval $[0, 1]$. The parameters are derived from the theory put forth in [1], giving the model a theoretical foundation.

At each time step, each agent calculates the average emotional transfer from all relevant agents. Specifically, from a sender S to a receiver R , the strength of the emotion received would be $\gamma_{SR} = \epsilon_S \cdot \alpha_{SR} \cdot \delta_R$. Logically, stronger channel, stronger sender expressiveness, and stronger receiver openness all lead to stronger emotional transfer to the receiver. [3] details the mathematical formulation, but, qualitatively, the fear level of an agent converges towards a weighted average of the group’s fear level. The speed at which this convergence occurs as well as the weighting depend on the parameter settings for the channel strength, expressiveness, and openness for each agent.

5 Durupinar Model

As opposed to independent interaction models, Durupinar et al. used a threshold model based on epidemiological models of disease contagion. While many types of epidemiological models exist, Durupinar implemented a simple version with only *susceptible* and *infected* states (as opposed to recovered, inoculated, etc. states). The model’s applicability to emotional contagion was not discussed in the initial publication, but its use assumes similarity between disease spread and emotion spread.

Each agent begins with a randomized threshold drawn from a pre-determined log-normal distribution. At each time step, for each agent, a random agent is chosen from

the relevant population group and if the agent is infected, will generate a random dose drawn from a pre-determined log-normal distribution and pass it to the original agent. If the agent is not infected, then a dose of 0.0 is generated. Each agent maintains a running history of the last K doses received. If the cumulative total of all doses in the agent's history exceeds his threshold, the agent enters the infected state. This causes the emotion level to be set to 1.0 with an exponential decay towards 0.0, at which point the agent re-enters the susceptible state. The random dose and threshold are generated from log-normal distributions with user-specified averages and standard deviations and K is a static global variable.

6 Proximity

When used in a simulation that includes physical space, an immediate deficiency arises in both models - the lack of specification of proximity's role in contagion. In any such simulation, proximity must enter the equation in some form. The VU model provides the channel strength parameter, which, if varied properly, can incorporate proximity effects into the contagion. Despite this, the authors did not provide guidance or exploration of possible implementations, thus we explore one in this work. The Durupinar publication did not provide experimental results specifically pertaining to the contagion and do not have a variable parameter such as in the VU model.

In this work, we implement a fixed neighborhood of effect for all agents within both models. Only agents within the specified distance are used in the model's calculations for contagion. In the VU model, this is equivalent to setting all channel strengths to 0.0 for pairs of agents that are too distant from each other and setting the remaining channel strengths to their pre-set levels otherwise. In the Durupinar model, this was a direct augmentation to the contagion effect, where we restrict the population used to neighboring agents only.

7 Authority Figure Effects

The second important modification that we require is specific to our scenario of evacuations, but is an example of the more general need for a contagion model to allow 'special contagion' agents. As noted in recent research [4], the role of authority figures is extremely important in evacuations not only for the information they provide but also for the calming effect they bring to anxious or fearful crowds. Neither model inherently discusses the implementation of agents with unique contagion attributes.

The VU model's individual-specific parameters allows for simple settings that would logically correspond with an authority figure (or to other special agents), but the authors did not explore possible impacts or implementations. While many implementations are possible, we choose to set all authority figures' openness parameters to 0.0, simulating the effect of proper training preventing authority figures from being susceptible to others' influences on their emotions. In combination with the proximity effect, this encourages agents near authority figures to calm their emotions towards 0.0 at each time step, producing the desired effect.

The Durupinar model specifies only population-wide, randomized parameter settings, necessitating model-level augmentations to include unique authority effects. Again, many implementations are possible. We choose to reproduce the resistance of authority figures to fear contagion by removing their contagion module entirely. In addition, to reproduce the calming effect that the VU model can naturally produce with the openness and expressiveness parameters, we introduce two changes to the base Durupinar model. First, we halve the level of fear in the agents surrounding authority figures at each time step. Second, we introduce an inoculated state that agents enter upon contact with an authority figure. They remain in this state for a fixed period of time that is reset as long as they remain in the presence of an authority figure. The second addition prevents the situation where a group of fearful agents at different points in the decay process simply pass fear back and forth to each other despite the presence of an authority figure. With these two augmentations, we are able to reproduce the authority calming effect noted by [4].

8 VU Experiments

We first explore the implications of varying the parameters in the base VU model when applied to the scenario described in Section 3. Then we show results pertaining to the rate and strength of emotion spread under the different versions of the model. Finally, we briefly touch on the implications on the predictions of safety under each of the versions of the model as they appear in the ESCAPES simulation engine.

8.1 Sensitivity Analysis

The parameters of interest in the VU model were the channel strengths, individual expressiveness settings, and individual openness settings. Given that we had a whole population of agents, we elected to use randomly drawn values for each of these based on a normal distribution. We explored variations of the averages and standard deviations used, but surprisingly, none yield substantial changes in the outcome of the simulation from both a contagion perspective (i.e., how the fear spread) and a safety analysis (i.e., how safe the evacuation was). The only exception was, unsurprisingly, when the receiver openness parameter varied tightly around a very low mean, leaving many agents with 0.0 openness. This caused the majority of agents to remain at their initial fear level, sometimes raising all agents' fear levels, which was vastly different from the convergence behavior seen in the other settings.

Figure 1 plots the percentage of people with low fear (≤ 0.1) on the y -axis and the time step on the x -axis. Figure 1a shows the results for variations in average channel strength whereas Figure 1b shows the same results for variations in average receiver openness. In both cases, the parameter being explored varied from 0.1 to 0.9 in increments of 0.2 while keeping a fixed standard deviation of 0.1 and the other two parameters were fixed with an average of 0.5 and a standard deviation of 0.1. As expected, when an event first occurs, those near it become fearful, hence the initial dip. However, due to the global convergence of fear levels and the fact that the vast majority of agents have 0.0 fear and do not know of the event, fear levels quickly decrease back to ≤ 0.1

levels. The tightness of the lines implies that the trend is robust to variations in the average channel strength. The same trend can be seen when the average openness is varied, with the exception of the previously mentioned situation. Similar tightness of lines was observed in other parameter variations.

We also conducted experiments exploring the second-order effects on safety, as measured by the ESCAPES system. In particular we examined the evacuation rates of pedestrians as well as the number of collisions experienced on average. Neither set of results showed significant variation through the parameter space, indicating the results' robustness to parameter variation.

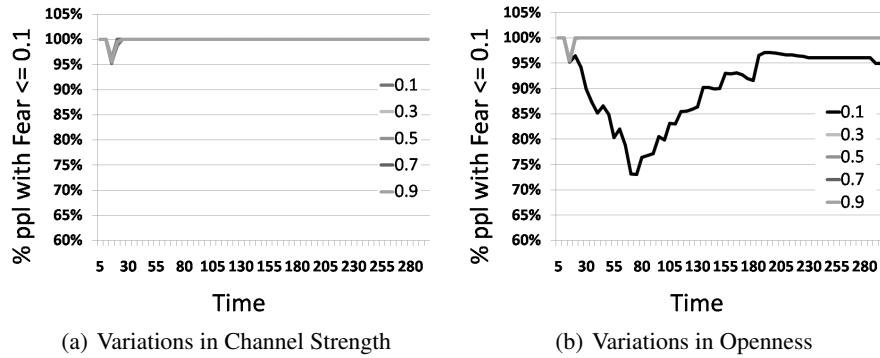


Fig. 1. Percentage of low-fear agents

8.2 Contagion Analysis

Now we discuss the effect on contagion as we include proximity effects and authority figure calming in the base model. Given the relative indifference of the model to parameter variations, we elect to use median values of 0.5 for the average of all parameters and fix the standard deviations at 0.1 for the results shown in this section.

Figure 2 shows the contagion trends of agents in the simulation under the three different models: original base model with ‘worldwide’ neighborhood, a model with a limited neighborhood of contagion, and a model with the limited neighborhood in conjunction with the authority figure modification. Each graph shows the percentage of agents remaining in the simulation that possess the labeled level of fear: ≤ 0.1 and ≥ 0.75 on the y -axis and time steps on the x -axis.

As can be seen, the trends are drastically different in each case. In particular, the base model always sees an extremely steep decrease in fear levels as the majority of agents do not know of the event and possess 0.0 fear, lowering the convergence target to near 0.0. When proximity is introduced, the fear levels reduce slowly, as high fear agents must pass close to low fear agents for this to occur. This leaves a large percentage of agents hovering above 0.1 through the entire duration of the simulation (recall

that there is no decay of emotion in the VU model). When authority figure calming is introduced, a middle ground between ‘worldwide’ and ‘proximity’ is achieved as authority figures are able to constantly reduce nearby agents’ fear levels despite them not encountering new 0-fear agents.

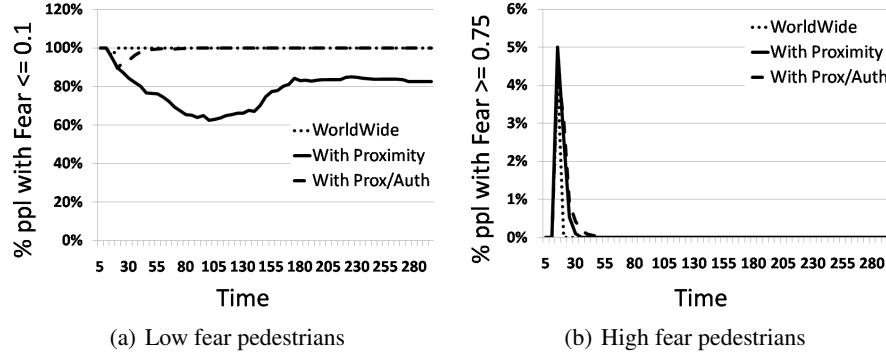


Fig. 2. Comparison of model with additions

Perhaps the strongest indication of the impact that our proximity addition has on realism comes from a series of illustrative snapshots in time of the locations of fearful agents in the scenario. In Figure 3, agents with fear greater than 0.1 are shown as red dots, agents with less than 0.1 but non-zero fear are shown as white dots. Figure 3b shows the location of fearful agents at time step 16 in the base model without proximity. Figure 3c shows the same snapshot for the model with proximity and without authorities. As can be seen, in the first few seconds following the event, agents throughout the scenario instantly become slightly fearful as they converge towards the fearful agents’ emotional level. When proximity is incorporated, however, a much more realistic spread can be seen with nearby agents becoming fearful.

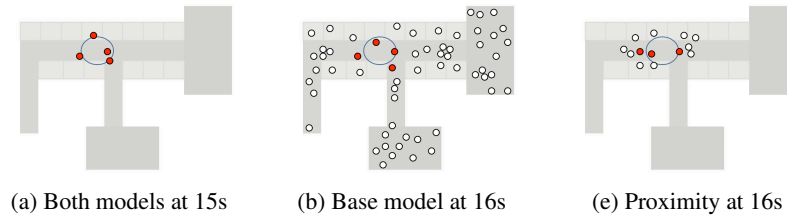


Fig. 3. Effect of proximity in VU model.

8.3 Safety Analysis

Now we evaluate the impact of the models of contagion on the actual evacuations as measured in the ESCAPES system. In particular we show the evacuation rates and average number of collisions of pedestrians in the simulation. Clearly, faster evacuation rates and lower number of collisions indicate better evacuations.

Figure 4a shows the percentage of pedestrians remaining in the simulation on the y -axis and the time step on the x -axis. Figure 4b shows the average number of collisions accumulated by people remaining in the simulation on the y -axis and the time step on the x -axis. The evacuation rate remains unchanged but there are noticeable differences in the number of collisions. In the ‘worldwide’ model and the model with both proximity and authorities, the number of collisions slopes up substantially slower than the model with proximity only as a result of the slower pace of people. The number of authorities was very high in these simulations, creating a situation similar to the ‘worldwide’ model with very little fear in the population. Thus, although the augmentations do not impact evacuation time, the prediction of safety as measured by the number of collisions is strongly affected.

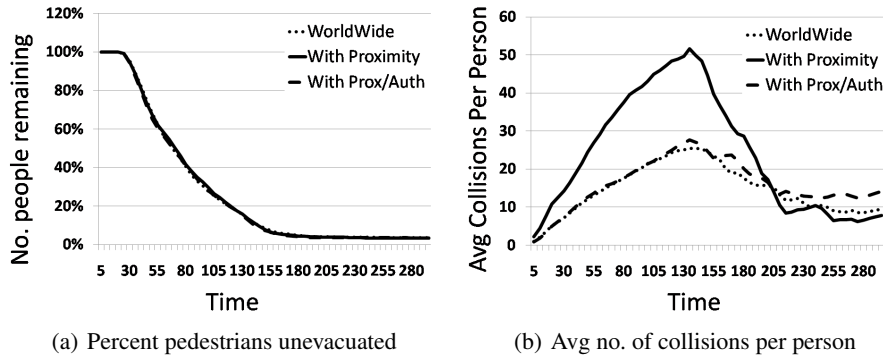


Fig. 4. Comparison of safety between models

9 Durupinar Experiments

Just as for the VU model, we begin with a sensitivity analysis of the Durupinar model. We then evaluate the implications of the augmentations on the way emotions spread in the simulation. Finally we discuss the implications for safety as they manifest in the ESCAPES simulation.

9.1 Sensitivity Analysis

Sensitivity analysis of the Durupinar model is considerably more complex than the VU model, because although the number of key parameters remain the same, they are

interdependent. Lower thresholds, higher dose strengths, or longer dose histories (K) would lead to more agents that are fearful because they would accumulate necessary doses faster. Clearly the relative values are what are important. Thus, we begin with fixed relative values and vary the parameters to identify key sensitivities. In particular, we begin with a baseline of K of 4, dose average of 2, dose standard deviation of 0.5, threshold average of 7, and threshold standard deviation of 2.

Unsurprisingly, altering any one of the parameters' averages OR standard deviations individually drastically alters the magnitude of the contagion effect, but not the overall trends. The exceptions are at extremely low values for K or dose distribution average and at extremely high values for threshold distribution average, when very few agents become fearful at all due to insufficient doses, dose sizes, or extraordinarily high thresholds. Figure 5a shows the percentage of low-fear pedestrians (≤ 0.001) on the y -axis and time steps on the x -axis, with each line representing a different setting of K . Figure 5b shows the same, but with each line showing different settings of the threshold distribution's standard deviation. We use 0.001 instead of 0.1 as before because the decaying aspect of the Durupinar model quickly causes people to fall below 0.1 fear, making 0.001 comparable to 0.1 in the VU experiments. As can be seen, the qualitative trends remain the same over the tested parameter-spaces, with the aforementioned exception. This implies that the model remains robust to parameter changes with respect to the contagion trends that emerge.

We again explored the second-order impacts of parameter variations on the safety of the evacuation by measuring the evacuation rates and average number of collisions of pedestrians in the simulation. As in the VU experiments, we again found no significant variation as the parameters varied across the non-trivial parameter space.

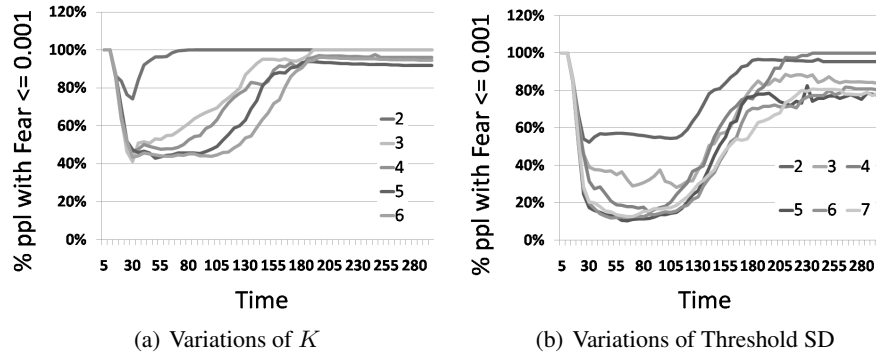


Fig. 5. Percent low-fear pedestrians

9.2 Contagion Analysis

Now we examine the effects that the augmentations have on actual contagion in the simulation. Unlike in the VU model, larger populations cause an *increase* in the overall

level of fear because of the infection model. Figure 6 shows the percentage of the remaining population with ≤ 0.001 fear and ≥ 0.75 fear on the y -axis and time step on the x -axis. The results use the baseline parameter settings mentioned in Section 9.1, but the tightness of trends is consistent through the non-trivial parameter space.

As expected, in the ‘worldwide’ model, a larger percentage of the population becomes fearful than in the other two cases, as shown by the lower point reached by the line. This is due to the fact that the entire population can potentially be infected. The model with proximity has a similar dip, although less pronounced since the susceptible population available consists only of neighboring agents. As shown by the less steep increase towards the tail, the number of fearful people tapers off more slowly in the proximity case than in the worldwide case because new susceptible people are encountered over time. The high fear graph, Figure 6b, shows no real surprises with the exception of the abrupt spikes, which is due to the fact that so few people have high fear at any given point in time and fear levels are set immediately to 1.0 upon infection.

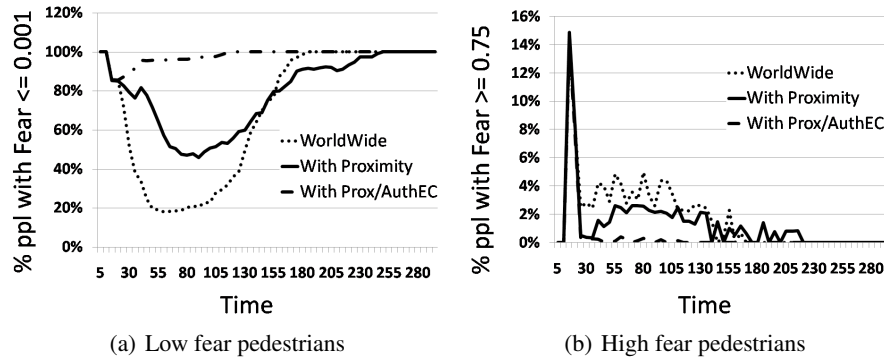


Fig. 6. Comparison of model with additions

As with the VU model, the effect of proximity on the Durupinar model provides far more realism in the contagion of fear than does the base model. This time we show agents with fear greater than 0.001, but the same dramatic increase in realism remains.

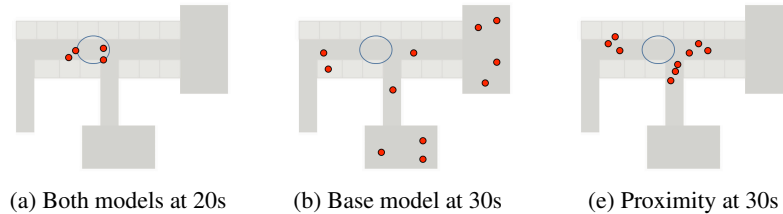


Fig. 7. Effect of proximity in Durupinar model.

9.3 Safety Analysis

Now we evaluate the impacts of the model augmentations on the actual predictions of evacuations as measured in the ESCAPES system. Again, we show the evacuation rates and average number of collisions of pedestrians in the simulation.

Figure 8a shows the percentage of people that remain unevacuated in the simulation on the y -axis and time step on the x -axis. As can be seen, the evacuation rate remains relatively unchanged. Figure 8b, however, shows very noticeable differences between the models in the number of collisions caused on average. In the ‘worldwide’ case, as with VU, the number of collisions slopes up substantially faster than the other two models for the same reasons. Next, the model with only proximity follows the same overall trend, but with a lower peak due to the fewer number of fearful people for the majority of the simulation. Finally, the model with both proximity and authority effects shows the lowest peak due to the lower number of infected people in addition to the authority calming effect slowing the pace of pedestrian travel and, therefore, making it easier for agents to avoid collisions with each other. Thus, although the model augmentations do not appear to impact overall evacuation time, the prediction of safety as measured by the number of collisions is again strongly affected.

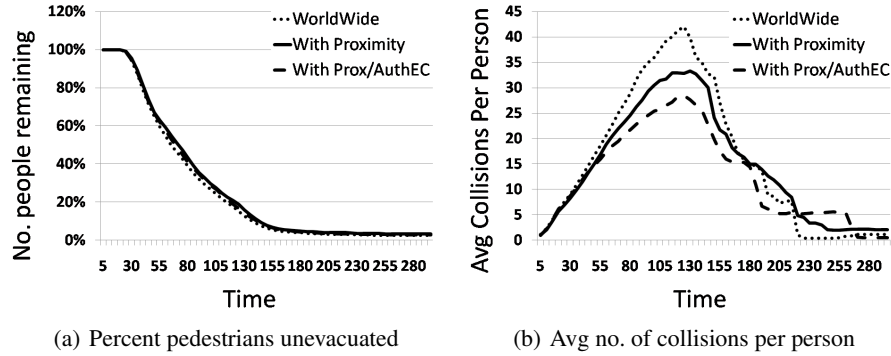


Fig. 8. Comparison of safety between models

10 Conclusions

Although both Durupinar and VU attempt to model emotional contagion, the underlying mechanisms differ drastically, with the VU model using an independent interaction-based approach and the Durupinar model using a threshold framework inherited from epidemiological studies. In the tests conducted, the VU model seemed to reproduce the contagion phenomenon with higher fidelity.

The Durupinar model possesses a number of inherent flaws due to its origins in epidemiological modeling. Its lack of a representation of ‘strength’ of the emotion means

that agents with more fear have the same impact as agents with only slight fear. Qualitatively, this is inconsistent with observations in social psychology [7]. Furthermore, the Durupinar model possesses no mechanism for ‘reverse’ contagion where a fearful agent might be impacted by the *lack* of fear of other agents. This means that a handful of fearful agents entering a room with 100 0-fear agents will not lose their own fear any faster than if they were alone. In fact, if their fear decays slowly enough and the infection sampling is done quickly enough, they will inevitably infect the entire crowd with their fear. While this *may* occur, the Durupinar model unrealistically implies that duration of exposure even to an extreme minority will *inevitably* lead to escalation.

The VU model, however, is not without its shortcomings. In particular, the lack of a decay function for the emotions means agents will *never* lose fear unless they encounter lower fear agents. The base VU model implemented here also never exhibits escalation because it enforces convergence to the weighted average. However, follow-up work has attempted to address this in [2]. Finally, the proximity and authority figure implementations used here, although an improvement over the base model, are but one of the possible ways that they can be done and further exploration is necessary to determine the most accurate, theoretically-based methods.

References

1. S. G. Barsade and D. E. Gibson. Group Emotion: A View from Top and Bottom. In D. Gruenfeld, E. Mannix, , and M. Neale, editors, *Research on Managing on Groups and Teams*, pages 81–102. JAI Press, 1998.
2. T. Bosse, R. Duell, Z. A. Memon, J. Treur, and C. N. V. D. Wal. A Multi-Agent Model for Emotion Contagion Spirals Integrated within a Supporting Ambient Agent Model. In *PRIMA-09*.
3. T. Bosse, R. Duell, Z. A. Memon, J. Treur, and C. N. V. D. Wal. A Multi-Agent Model for Mutual Absorption of Emotions. In *ECMS-09*, 2009.
4. J. Diamond, M. McVay, and M. W. Zavala. Quick, Safe, Secure: Addressing Human Behavior During Evacuations at LAX. Master’s thesis, UCLA Department of Public Policy, June 2010.
5. W. Doherty. The Emotional Contagion Scale: A Measure of Individual Differences. *Journal of Nonverbal Behavior*, 21(2), 1997.
6. F. Durupinar. *From Audiences to Mobs: Crowd Simulation with Psychological Factors*. PhD dissertation, Bilkent University, Dept. Comp. Eng, 2010.
7. E. Hatfield, J. T. Cacioppo, and R. L. Rapson. *Emotional Contagion*. Cambridge University Press, 1994.
8. M. Hoogendoorn, J. Treur, C. v. d. Wal, and A. v. Wissen. An Agent-Based Model for the Interplay of Information and Emotion in Social Diffusion. In *IAT-10*.
9. L.-O. Lundqvist. Factor Structure of the Greek Version of the Emotional Contagion Scale and its Measurement Invariance Across Gender and Cultural Groups. *Journal of Individual Differences*, 29(3):121–129, 2008.
10. J. Tsai, G. Kaminka, S. Epstein, A. Zilka, I. Rika, X. Wang, A. Ogden, M. Brown, N. Friedman, M. Taylor, E. Bowring, S. Marsella, M. Tambe, and A. Sheel. ESCAPES: Evacuation Simulation with Children, Authorities, Parents, Emotions, and Social Comparison. In *AAMAS-11*.

Between Agents and Mean Fields

H. Van Dyke Parunak

3520 Green Court, Suite 250
Ann Arbor, MI 48105 USA
van.parunak@jacobs.com

Abstract. Some agent-based models use analogs of insect pheromones for coordination. We situate these techniques in the spectrum of modeling tools. Analysis and simulation show that pheromone models are intermediate between classical agent-based models and mean-field models, inspired by statistical physics. This position is not fixed, but can be adjusted by pheromone parameters (notably, the propagation factor), providing new design options for ABMs.

Keywords: Pheromones, mean-field models, agent-based models

1 Introduction

Many problems to which agent-based models (ABMs) are applied can also be modeled using mean-field models, typically as differential equations describing the time evolution of system variables. Each type of model has advantages and disadvantages.

One approach to ABM imitates insect pheromones. This paper’s central claim is that pheromone-based coordination is intermediate between classical multi-agent systems and mean-field models. Pheromones provide “lumpy-field models”: each agent generates a non-uniform field whose mode is related to the agent’s location.

We extend a simple model of population dynamics [13] to contrast agents and mean-field models. We analyze this model in both mean-field and conventional agent configurations, to develop intuitions as to the mediating position of pheromone based models. Then we confirm and refine these intuitions using simulations of the system.

Section 2 compares agent-based vs. mean-field models. Section 3 situates pheromone-based coordination this approach in the agent-equation space. An experiment in Section 4 validates and illustrates the claims of Section 3, with special attention to the mapping between computational and physical models of propagation. Section 5 discusses implications of the research, and Section 6 concludes.

2 Agents and Mean Fields

Agent-based modeling is an alternative to equation-based modeling (EBM) [11, 14], in which differential or difference equations capture the evolution of variables. A salient difference [11] is that agent-based modeling directs the modeler’s attention to the individual *entities* in the domain, while equation-based models focus on *variables*. These variables may be extensive (e.g., the total population or agent density), or in-

tensive (e.g., parameters of individual agents). Intensive variables may reflect individual agents, or (as averages) the system as a whole.

EBM favors extensive variables, or averages over intensive ones, which permit parsimonious closed-form equations. The perspective of such a model is centralized, since such variables require global information. In contrast, the natural tendency in ABM is to define agent behaviors in terms of observables accessible to the individual agent, favoring a local rather than a global viewpoint. The evolution of system-level observables does emerge from an ABM, but the modeler is less likely to use these observables to drive the model's dynamics than in an EBM.

Statistical physics offers an analog to this contrast between local and global information. Systems of many interacting particles are analytically intractable. Physicists consider a single typical particle, and estimate the influence of all the others on it as an average influence. The resulting model is called a "mean-field model" (MFM), because it replaces many individual entities by averages over them.

The canonical example is the Ising model of ferromagnetism [3]. An n -dimensional lattice has z sites indexed with i , each with a spin S_i of either +1 or -1. These spins are subject to two influences. The first is an external field of strength h . The second is the influence of neighboring spins. The Hamiltonian of the system (roughly, its energy) is

$$(1) \quad E(\{S_i\}) = -h \sum_i S_i - J \sum_{\langle ij \rangle} S_i S_j$$

where the second sum is over all nearest neighbors and $J > 0$ is the strength of pairwise interactions. The first sum measures the energy due to the interaction of spins with the external field (lowest when aligned, thus negative), and the second sum, the contribution of the nearest neighbor interactions (again lowest when aligned).

In three or more dimensions, this system is analytically intractable. To circumvent this problem, note that the contribution of a single spin to the Hamiltonian is

$$(2) \quad e(S_i) = -h S_i - J S_i \sum_j S_j$$

summing over all sites adjacent to S_i . The mean field approach replaces the individual S_j 's in the summation with their average value $\langle S_j \rangle$. The contribution of a single site can now be written as a linear function of that site,

$$(3) \quad e_{mf}(S_i) = -h S_i - J S_i \sum_j \langle S_j \rangle = -h_{mf} S_i, \text{ where } h_{mf} \equiv h + Jz \langle S_j \rangle$$

This simplified system is tractable, but neglects the interaction between spins. We have replaced the average of the interactions, $\langle S_i S_j \rangle$, with the interaction of the averages, $\langle S_i \rangle \langle S_j \rangle$. This replacement corresponds to an assumption that the spins are statistically independent. In fact, the problem is interesting just because the spins are *not* independent of one another. The consequences of the approximation vary with the dimension of the lattice. In one dimension, it yields qualitatively false results, but as the number of dimensions increases, in spite of the *a priori* unreasonableness of the mean-field assumption, it quantitatively approximates the exact result.

Following this terminology, we describe any model that relies on system-level averages over agent variables as a *mean-field model*. EBMs using extensive and averaged system variables are the most common example of mean-field models, but an agent-based model that uses global knowledge can also do mean-field reasoning. In both cases, the mean-field approach accepts an unrealistic assumption of independence among key variables in exchange for improved tractability.

As in physics, so in multi-agent systems, an MFM has limited accuracy. A given agent parameter may vary widely over the population, and encounters among agents may depend critically on that parameter. Replacing diverse values with a single aver-

age may qualitatively change the behavior of the system [13]. Still, MFMs can give concise insight into the behavior of a system that is obscured by discrete models, and researchers increasingly present both models for a system [4, 7].

3 Pheromone-Based Agents

Some social insects coordinate their activities by depositing and sensing chemicals called “pheromones” in a shared environment [15]. Some research in software agents [9] imitates these mechanisms, using a structured environment to implement fundamental information processing tasks. 1) Cells *aggregate* deposits by separate agents, a form of information fusion. 2) Pheromone *propagates* from cells of high concentration to those of low concentration, a form of communication. 3) *Evaporation* discards obsolete information, a primitive form of truth maintenance that runs in constant time (as opposed to NP-complete truth maintenance with non-trivial logics [5]).

Digital pheromones were originally used for spatial problems: agents move over a map of the world and the peak of the pheromone field converges to a shortest path [12]. Their quantitative nature makes them seem inappropriate for modeling more complex problems, such as social behavior and human decisions based on symbolic reasoning. In fact, pheromones can be applied to such problems by capturing the logical structure of the problem in the environment over which agents move and in which they deposit and sense pheromones. In traditional behavior modeling, each agent manipulates a logical model inside its head. Pheromone agents handle such problems by externalizing the logical structure and moving over it. An example of this approach is the rTÆMS system for reasoning about plans [2, 8]. In principle, pheromone agents can address any problem that can be represented as a graph.

These mechanisms resemble an MFM. In statistical physics, MFMs simplify computation. Instead of accounting for each particle’s interaction with many other particles, we replace the other particles by an average influence, and then compute the behavior of the particle of interest with that average influence. The pheromone field in a pheromone-based agent system is analogous to the average influence in a physical MFM. The field consists of deposits by individual agents. Up to a normalizing constant, this field gives the probability of encountering an agent of the type represented by the field at a given location. When one agent makes decisions based on the field, rather than on explicit representation of other agents, it is reasoning about a weighted average influence of the other agents—weighted because the field is generated by those agents as they move about, and reflects their recent locations.

This weighting is an important refinement over the naïve average in the mean-field approximation to the Ising model in Section 2. Consider one robot coordinating with four others in a 20x20 grid. A naïve mean-field estimate of the probability of encountering another robot in any given cell is $4/400 = 0.01$. This probability is isotropic, and so gives no guidance to the robot. At the other extreme, the robot could communicate directly with other robots, remember their most recent locations, and determine with high certainty whether or not a given cell contains another robot.

The pheromone approach is intermediate. Each agent contributes to a field in its vicinity. Other agents in the vicinity also contribute to the field. So the field is an

average over multiple agents, but a localized average. This localized average is not just over agents, but also over space and time.

- *Aggregation* collects contributions from multiple agents that are in the same area, averaging over multiple agents. The larger the cells into which the environment is divided, the more agents are included in the average.
- *Propagation* spreads out the deposits, averaging over space. The stronger the propagation, the larger the area of space over which the averaging occurs.
- *Evaporation* determines how long pheromones persist, averaging over time. The stronger the evaporation, the shorter the period over which averaging occurs.

Let's explore the interaction of these dynamics analytically, as they function with chemical pheromones.¹ A single stationary agent at the origin deposits pheromone at a constant rate D . The pheromone field $\varphi(r, t)$ is multiplied by a constant evaporation rate $E \in (0, 1)$ at each time step, thus removing pheromone at a rate $(1 - E)\varphi$, and propagation takes the form of a diffusion process with rate F . We assume a separable solution as a product of spatial and temporal terms of the form $\varphi(\vec{r}, t) = \varphi_s(\vec{r}) \varphi_t(t)$. Empirically, the resulting field reaches equilibrium with a constant strength φ_0 at the origin, providing the boundary condition $\varphi(0, \infty) = \varphi_0$, and this result together with conservation of matter implies $\varphi(\infty, t) = 0$.

The spatial component $\varphi_r(r)$ must satisfy

$$(4) \quad F \nabla^2 \varphi_r - (1 - E) \varphi_r = 0$$

That is, the gain in field strength at any location due to propagation of pheromone is balanced by the loss due to evaporation. (4) is the Helmholtz equation

$$(5) \quad \nabla^2 \varphi_r + k^2 \varphi_r = 0$$

with $k = iB$, $B = \sqrt{\frac{1-E}{F}}$. φ_r does not vary with angle, so we use polar coordinates $\vec{r} = (r, \theta)$, for which the solution is a Bessel function. Because k is pure imaginary, we have a modified Bessel function (in this case, of order zero). The boundary condition $\varphi_r(\infty) = 0$ means it is of the second type, $AK_0(Br)$, where A and B are scale factors that do not vary with location and K_n is the modified Bessel function of the second type of order n . In Cartesian coordinates, the corresponding solution would be Ae^{-Br} , and the shape of K_n is similar to e^{-r} , so we call A the amplitude and B the exponent.

This function is undefined at $r = 0$. To achieve the boundary condition $\varphi(0, \infty) = \varphi_0$, we assume that deposits occur in a disk of radius r_0 around the origin, within which density is constant. For comparison with a discrete implementation with cells of unit area, we choose $r_0 = 1/\sqrt{\pi}$, so that the disk corresponds to the cell where the deposit is made. Then we normalize the Bessel function to yield $\varphi(0)$ at the boundary, so

$$(6) \quad \varphi_r(\vec{r}) = \begin{cases} \varphi(0), & r < r_0 \\ \varphi(0) \frac{K_0(Br)}{K_0(Br_0)}, & r \geq r_0 \end{cases}$$

To get φ_0 , examine the temporal behavior at $r = 0$. Based on the observed convergence behavior of the system, $\varphi_t(t) = (\alpha - \beta e^{-\gamma t})$, where α is the asymptotic value we seek. So the overall solution will be of the form

¹ I am grateful to Robert Savit for suggestions in formalizing these dynamics. Unlike [1], and like natural systems, our propagation conserves total pheromone.

$$(7) \quad \varphi(\vec{r}, t) = \begin{cases} \alpha(1 - e^{-\gamma t}), & r < r_0 \\ \alpha(1 - e^{-\gamma t})^{\frac{K_0(B r)}{K_0(B r_0)}}, & r \geq r_0 \end{cases}$$

where the boundary condition $\varphi(r, 0) = 0$ requires $\beta = \alpha$.

As $t \rightarrow \infty$, the rate of deposit D into the deposit disk must equal the sum of two flows out of the disc: evaporation (with value $(1 - E)\alpha$), and diffusion (by Fick's law, $F\nabla\varphi$). Integrating the flux around the perimeter of the deposit disk yields

$$(8) \quad \begin{aligned} \int_0^{2\pi} F\nabla\varphi(r_0) r_0 d\theta &= 2\pi r_0 F\nabla\varphi(r_0) \\ &= 2\pi r_0 F\alpha \frac{K_1(B r_0)}{K_0(B r_0)} = 2\pi r_0 \sqrt{F(1-E)}\alpha \frac{K_1(B r_0)}{K_0(B r_0)} \end{aligned}$$

So

$$(9) \quad D = (1 - E)\alpha + 2\pi r_0 \sqrt{F(1-E)}\alpha \frac{K_1(B r_0)}{K_0(B r_0)}$$

Solving for the equilibrium concentration α ,

$$(10) \quad \alpha = \varphi_0 = \frac{D}{(1-E) + 2\pi r_0 \sqrt{F(1-E)} \frac{K_1(B r_0)}{K_0(B r_0)}}$$

Thus each agent's contribution to the field is peaked about its location, decaying exponentially with distance. The width of the peak (reflected in the exponent B) depends on the evaporation and propagation rates. Evaporation also localizes the peak temporally. If the agent is moving, it leaves a tear-drop shaped pheromone field, widest and highest near its current location and the tail tracing its recent history.

A pheromone-based agent reasons about other agents, not by tracking their individual locations, but by monitoring their aggregate pheromone field. If it wants to meet another agent, it climbs the field. If it wants to avoid other agents, it moves down the gradient. This process is simpler than monitoring individual agent locations. It is approximate, since each agent's location is represented only by a distribution about its actual location. But the exponential decay of $K_0(r)$ guarantees that that distribution does not become uniform until it evaporates completely. We might say that the agent is using, not a mean-field theory, but a "lumpy-field" theory.

Like the Ising example, the pheromone approach simplifies by discarding information about higher-order interactions. For example, the field representing the other agents does not distinguish them from one another, and so cannot coordinate a strategy that requires the agents to interact in a specified order. This information could be recovered, even in a pheromone model, by maintaining different fields for each agent, increasing processing complexity. Techniques based on parallel exploration of alternative sequences can also address this problem with pheromone-based agents, again requiring additional processing to recover the information lost in the aggregate field.

Now consider the polyagent technique of estimating alternative futures for a Red and Blue avatar and their ghosts. A single Red ghost makes successive moves that require it to decide whether, at each move, its avatar would meet Blue. It estimates the probability of meeting Blue at each move by sampling Blue's pheromone at its location. That pheromone field is generated by multiple Blue ghosts. Red's ghost may simulate successive encounters with Blue at successive time steps, representing a future in which the Red avatar and the Blue avatar repeatedly encounter one another. But the presence of a non-zero Blue pheromone field at successive locations visited by the single Red ghost does *not* guarantee the existence of *any* single future for Blue that visits those locations. Imagine a future in which the first encounter of Red and

Blue removes Blue from the system. In this case, it is clearly unrealistic for a Red ghost to simulate a later encounter with Blue. Yet it may encounter Blue pheromone later in its trajectory. The different portions of the field that the Red ghost encounters may have been generated by different Blue ghosts, whose distinct identities have been discarded by the field representation. As in the MFM of the Ising system, the pheromone representation simplifies computation by an independence assumption: it assumes that Blue's future is independent of the encounter with Red, even though the simulation is interesting just because its future is *not* independent of Red.

In spite of this simplification, polyagents can forecast the trajectory of complex systems better than other computational techniques with which they have been compared, and better than human experts [6]. The reason appears to be that most of the futures that one can enumerate are not in fact accessible. The domain itself constrains the ghosts' movements to a relatively limited number of trajectories.

4 An Experiment

We illustrate the lumpy-field nature of pheromone-based systems with a simulation experiment [13] that contrasts the behavior of a system of discrete agents with a set of differential equations that capture the mean-field behavior. A pheromone version shows behavior intermediate between the other two models.

4.1 Description of Experiment

A toroidal arena hosts two species of agents. Species *I* is immortal, uniformly distributed with average density n_I , and diffuses with diffusion coefficient D_I . Species *M* is mortal, with initial uniform density n_M . Mortal agents die at a constant rate μ , divide with rate λ when they encounter an immortal, and diffuse with coefficient D_M .

Continuity and symmetry predict that immortals will remain homogeneously distributed, $n_I(x) = n_I$. The time evolution of n_M is represented by

$$(11) \quad \frac{\partial n_M}{\partial t} = D_M \nabla^2 n_M + (\lambda n_I - \mu) n_M$$

For initially uniform spatial distributions of both species, the solution is

$$(12) \quad n_M(t) = n_M(0) e^{t(\lambda n_I - \mu)}$$

If $\lambda n_I < \mu$, mortals will become extinct.

An agent-based simulation (NetLogo, Table 1) shows different behavior.

Fig. 1 plots the population of mortals over time for $\mu - \lambda n_I = 0.3$, a regime in which Equation (12) predicts extinction. Instead, the mortal population explodes.

The difference between models is due to a mean-field assumption in the EBM. Overall, immortal agents are homogeneously distributed, but as sampled by mortal agents, they are not. Mortals are born only when a mortal encounters an immortal, and come into existence close to an immortal. The density of immortals *near mortals* is far greater than n_I . Fig. 2 shows a screen shot at

Table 1. General Parameters for NetLogo Simulations

Arena size	81 x 81 cells
Immortal population	50
Immortal diffusion	0.14
Mortal population	150
Mortal diffusion	0.11

the end of Fig 1. Three immortals form breeding clusters that generate mortals faster than they can die.

Not every run with $\lambda n_I < \mu$ explodes. The outcome depends on location in parameter space, and stochasticity. Several system parameters are implemented as stochastic choices by each agent, including birth rate, death rate, and diffusion. For example, a mortal who meets an immortal samples a uniform distribution on $[0, 1]$ and gives birth only if the result is less than or equal to λ . In different runs, these stochastic decisions can yield different outcomes. In addition, different parameters affect mortal survival, by determining the viability of a breeding cluster.

- With too few agents, a mortal may die before meeting an immortal and breeding.
- If the birth rate is too low, losses due to diffusion and death outpace births.
- If the mortal diffusion rate is too high, new mortals drift away from the immortal parent so fast that the breeding cluster disintegrates.

We observe the effects of these parameters by analyzing repeated runs. Empirically, if $n_M > 1000$, the system explodes. So we run until either $n_M = 0$ or $n_M > 1000$. We repeat each configuration 25 times with different random seeds, and record the percentage of trials in which n_M goes to zero.

Fig. 3 shows the percentage of runs (“pct”) in which mortals survive as a function of birth and death rates. $\lambda n_I - \mu$ ranges from 0 to -0.5, except when the death rate is zero, in which case the difference is in $[0.004, 0.008]$. Toward the front of the figure (low birth rate and high death rate), virtually all runs end in extinction. This region is nearly flat until one reaches a region around the line $\mu = 0.6\lambda - 0.2$. Then the probability rises rapidly to the corner with high birth rates and low death rates.

Survival or extinction does depend on a balance between birth and death rates (Fig. 3), but birth rate has more influence than in the MFM. $\partial\mu/\partial\lambda$ along the edge of the region where the agent-based model differs from the MFM is 0.6, while the corresponding value in the equation, n_A , is $50/81^2 \cong 0.0076$. The simulation behaves as

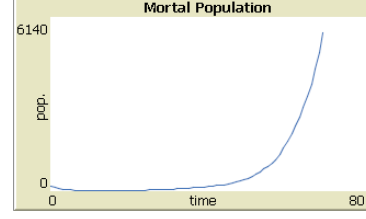


Fig. 1. Mortal Population, ABM

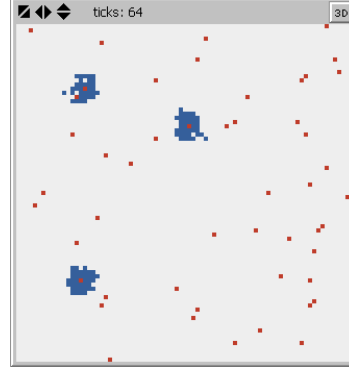


Fig. 2. Breeding clusters. The dark clumped agents are mortals; the scattered agents are immortals.

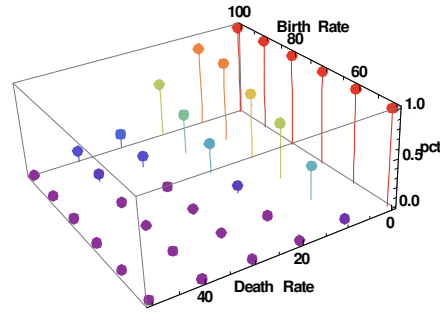
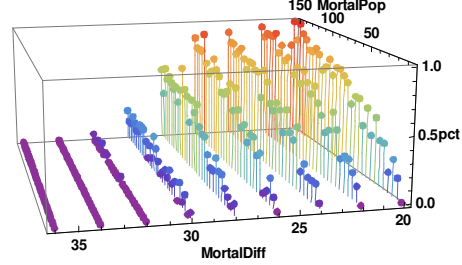


Fig. 3. Survival probability (“pct,” percent surviving) as function of birth and death rates

though the density of mortals were nearly 100 times greater than n_A . Factors like mortal population and rate of diffusion modulate this effect (Fig. 4). Low initial mortal population and high diffusion lead to extinction; the opposite corner leads to survival.



4.2 Adaptation for Pheromone-Based Agents

Fig. 4. Observed probability of survival as function of mortal diffusion rate and initial population

In both models, the probability $p(\text{birth})$ that a mortal gives birth is $p(\text{parent}) * \lambda$. λ is just the probability of a birth conditioned on the presence of an immortal parent. $p(\text{parent})$ is the probability that a parent is present. The models differ in how they estimate $p(\text{parent})$. For the MFM, $p(\text{parent}) = n_i$. In the discrete agent model,

$$p(\text{parent}) = \begin{cases} 1 & \text{if immortal is in cell} \\ 0 & \text{otherwise} \end{cases}$$

(If a cell contains multiple immortals, the computation is repeated for each.)

Pheromones offer an intermediate “lumpy-field” approach with simpler computation than the discrete model and more accuracy than the mean field approach.

Each immortal deposits one unit of pheromone per step at its location. The total deposit per step is n_i . Each mortal samples the field ϕ at its location, estimates $p(\text{parent})$, and computes $p(\text{birth})$. If $\phi > 1$, the mortal behaves as though it encountered $\lfloor \phi \rfloor$ immortals, plus one more with probability $\phi - \lfloor \phi \rfloor$. A single deposit by stationary immortals and no evaporation or propagation recovers the discrete model.

Constant deposit and exponential evaporation lead to an asymptotic fixed point. By setting the deposit rate to 1 per immortal and evaporation to 0.5, the total pheromone is constant, and equal to the immortal population. With stationary immortals and no propagation, this configuration also recovers the discrete model.

When immortals move, or when we allow propagation, the field extends beyond the cell containing the immortal. This spreading can allow an invalid birth: a mortal may think it is in the presence of an immortal when it is not. As in the Ising model, increased error is the cost of simpler computation.

We use NetLogo’s *diffuse* function for propagation. *diffuse* takes an argument ρ in $[0,1]$, subtracts $\rho * \phi$ from each cell, and distributes $\rho * \phi$ evenly among the cell’s eight neighbors, conserving the total field. The function is applied to all cells at once.

Propagation modeled on chemical diffusion in an insect system (4) is a function of the second derivative of ϕ , not its absolute strength. NetLogo’s *diffuse* operator (like most implementations of digital pheromones) depends only on the local strength of the field and relies on exchanges from neighboring cells to incorporate information analogous to the local derivative. So we need to qualify our claim that the NetLogo implementation represents propagation. The two processes share some features: the form of (6) is a good fit to the fixed point of the observed distribution. But quantitatively, the correspondence is weaker.

Under the dubious mapping $F = \rho$, *diffuse* yields an exponent related to the exponent $B = \sqrt{\frac{1-E}{F}}$, but an amplitude less clearly related to (10). Fig. 5 plots observed vs. theoretical amplitude. The theory is off by more than an order of magnitude. Fig. 6 plots the ratio of observed exponent *obs* to theoretical exponent *the*. This relationship is closer, a good fit to $\frac{obs}{the} = 2.27(e^{0.35 obs} - 1)$, but hardly constant, much less 1. Nevertheless, we press ahead.

As ρ increases, a pheromone model should behave less like a discrete model and more like an MFM. Because of the lumpiness of the field, the resulting error will never be as bad as in the MFM, as Fig. 7 shows. As ρ increases, the probability of survival approaches zero everywhere except when $\mu = 0$, as in the MFM.

We can quantify this difference. Each scenario (mean-field and pheromone with various diffusion rates) yields survival rates as a function of λ and μ that differ from the discrete agent system. We weight these differences by the differences observed in the MFM, and normalize by the sum of these weights. On this scale, the MFM scores 1, and the

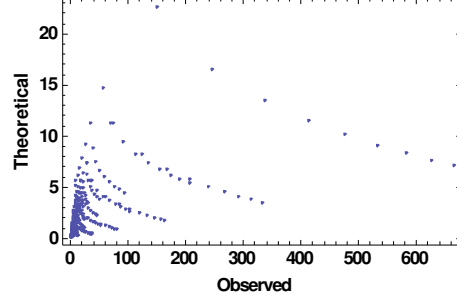


Fig. 5. Relation between observed and theoretical amplitude

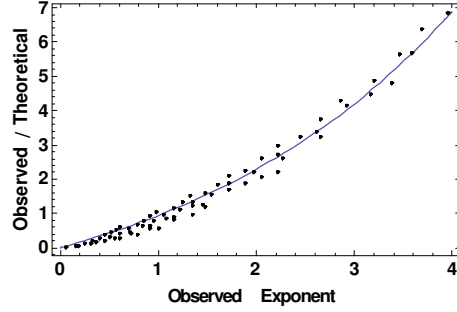


Fig. 6. Relation between observed and theoretical exponent

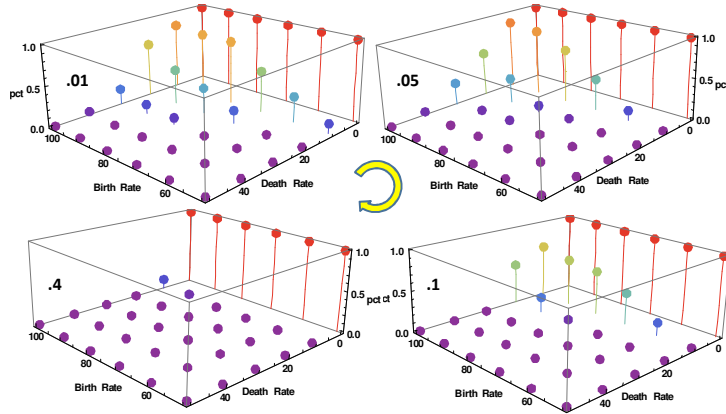


Fig. 7. Survival, propagation = (from top left, clockwise) 0.01, 0.05, 0.1, 0.4

discrete agent system scores 0. Fig. 8 shows this score as a function of propagation. As anticipated, the error grows with propagation rate, and asymptotes before it reaches the mean-field level.

In Fig. 8, error *increases* when diffusion is 0, compared with 0.01. We expect to recover the discrete model in this case. The difference is that propagation of 0 corresponds to the discrete model only if the agent depositing the pheromone is stationary. Our immortal agents move occasionally, and

may leave behind a deposit that can mislead mortals. Propagation, as well as evaporation, reduces this obsolete information. To confirm this effect, we run the system with stationary immortals, and obtain the error scores in Table 2. Now propagation 0 does indeed yield error less than 0.01, and almost achieves the 0 score of discrete agents. (The residual error is due to the difference between stochastic decisions in the pheromone agent and deterministic ones in the discrete agents.)

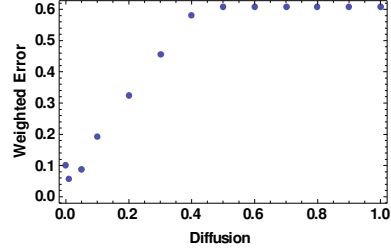


Fig. 8. Dependence of weighted error on diffusion

5 Discussion

Both analysis and simulation show that pheromone-based agents are intermediate between MFMs and conventional agents. Such mediating techniques provide a richer array of techniques to balance computational complexity and model fidelity. At the same time, the richer the toolbox, the more skills are required to use it. Though preliminary, our results suggest some guidelines concerning the use of pheromone systems.

Pheromones can reduce fidelity. When agents base decisions on fields rather than on direct observation, they make independence assertions that may violate the structure of the problem. Whether these violations are empirically meaningful is problem-dependent. In practice, pheromones often yield very useful results. We speculate that just as the mean-field analysis of the Ising model becomes more accurate as dimensionality increases, the high dimensionality of many domains to which pheromones are often applied makes the violations less significant.

Two conditions recommend pheromone methods: problem complexity and the need for statistical estimates.

A conventional ABM may require much knowledge engineering, computational effort, or both, making pheromones attractive. Still, the approach is a heuristic, and validation against domain-specific experiments (e.g., [6]) is essential.

Statistically, many domains require estimating alternative outcomes. The space of outcomes is often so large that even thousands of replications of a conventional agent model may not give a meaningful sample [10]. This requirement is a clear indicator for the polyagent technique of virtual or “ghost” agents that interact indirectly through fields summarizing the distribution

Table 2. Error for Stationary Agents.

Configuration	Weighted Error
Discrete agents	0 (by definition)
Propagation 0	0.079
Propagation 0.01	0.160

Table 3. Comparison of Types of Models.

Model Type	Representation of Non-Self Agents	Interactions	Computational Cost	Error
Classical MAS	Explicit	Direct	High	Low
Pheromones	Lumpy field	Indirect, non-uniform	Medium	Medium
Mean-field	Uniform field	Uniform	Low	High

over multiple futures of each agent.

One benefit of pheromones over an EBM is the ability to tune fidelity using propagation, which governs the degree of spatial averaging over agents. (We can also tune using the evaporation rate, adjusting the degree of temporal averaging.) Lower propagation improves fidelity (Fig. 8). Models more complex than ours may show a trade-off between fidelity and other characteristics (e.g., stability or convergence).

Finally, comparison of a representative computational analog of propagation (Net-Logo’s *diffuse* operator) with the physical diffusion that propagates insect pheromones shows qualitative similarity between the mechanisms, and suggestive though non-trivial quantitative relations with the exponent of the spatial decay and the asymptotic amplitude of the temporal stabilization that merit further study.

6 Conclusion

MFMs simplify the computational burden of tracking detailed multi-agent interactions by replacing individual interactions with statistical summaries of the population from the perspective of a single agent. This simplification imposes unwarranted independence assumptions. In spite of these assumptions, the results are useful often enough that these models continue to be widely used.

Conventional ABMs compute each interaction, yielding higher accuracy than an MFM, but the computational burden precludes thorough sampling.

Pheromone-based constructs such as the polyagent reduce the computational cost of modeling the space of interactions. There is no free lunch. This simplification (as in the corresponding physics theories) can usually be described as an independence assumption. Because the agent framework retains the discrete structure of the problem, the resulting error is often much less than in a complete mean-field treatment, and can be tuned by adjusting the degree of propagation of the pheromones.

Table 3 offers a summary comparison of the three approaches.

Recognizing the mediating position of pheromone models between conventional agents and equation-based MFMs allows modelers to make more appropriate use of this promising technology.

7 References

1. Brueckner, S.: Return from the Ant: Synthetic Ecosystems for Manufacturing Control. Thesis at Humboldt University Berlin, Department of Computer Science (2000)

2. Brueckner, S., Belding, T., Bisson, R., Downs, E., Parunak, H.V.D.: Swarming Polyagents Executing Hierarchical Task Networks. In Proceedings of Third IEEE International Conference on Self-Adaptive and Self-Organizing Systems (SASO 2009), pages 51-60, IEEE (2009)
3. Evans, M.: MP4/P4 Statistical Physics. University of Edinburgh, Edinburgh, UK (2009). <http://www2.ph.ed.ac.uk/~martin/sp/>
4. Glinton, R., Scerri, P., Sycara, K.: Exploiting Scale Invariant Dynamics for Efficient Information Propagation in Large Teams. In Hoek, W.v.d., Kaminka, G.A., editors, the Ninth International Conference on Autonomous Agents and Multi-Agent Systems (AAMAS 2010), pages 21-28, IFAAMAS, Toronto, Canada (2010)
5. Levesque, H.J., Brachman, R.J.: Expressiveness and Tractability in Knowledge Representation and Reasoning. *Computational Intelligence*, 3(2):78-93 (1987)
6. Parunak, H.V.D.: Real-Time Agent Characterization and Prediction. In Proceedings of International Joint Conference on Autonomous Agents and Multi-Agent Systems (AAMAS'07), Industrial Track, pages 1421-1428, ACM (2007)
7. Parunak, H.V.D.: A Mathematical Analysis of Collective Cognitive Convergence. In Proceedings of the Eighth International Conference on Autonomous Agents and Multi-Agent Systems (AAMAS09), pages 473-480, (2009)
8. Parunak, H.V.D., Bisson, R., Brueckner, S.A.: Agent Interaction, Multiple Perspectives, and Swarming Simulation. In Proceedings of the International Joint Conference on Autonomous Agents and Multi-Agent Systems (AAMAS 2010), pages 549-556, IFAAMAS (2010)
9. Parunak, H.V.D., Brueckner, S.: Ant-Like Missionaries and Cannibals: Synthetic Pheromones for Distributed Motion Control. In Proceedings of Fourth International Conference on Autonomous Agents (Agents 2000), pages 467-474, (2000)
10. Parunak, H.V.D., Brueckner, S.: Concurrent Modeling of Alternative Worlds with Polyagents. In Proceedings of the Seventh International Workshop on Multi-Agent-Based Simulation (MABS06, at AAMAS06), pages 128-141, Springer (2006)
11. Parunak, H.V.D., Savit, R., Riolo, R.L.: Agent-Based Modeling vs. Equation-Based Modeling: A Case Study and Users' Guide. In Proceedings of Multi-agent systems and Agent-based Simulation (MABS'98), pages 10-25, Springer (1998)
12. Sauter, J.A., Matthews, R., Parunak, H.V.D., Brueckner, S.A.: Performance of Digital Pheromones for Swarming Vehicle Control. In Proceedings of Fourth International Joint Conference on Autonomous Agents and Multi-Agent Systems, pages 903-910, ACM (2005)
13. Shnerb, N.M., Louzoun, Y., Bettelheim, E., Solomon, S.: The importance of being discrete: Life always wins on the surface. *Proc. Natl. Acad. Sci. USA*, 97(19 (September 12)):10322-10324 (2000)
14. Sterman, J.: *Business Dynamics*. New York, NY, McGraw-Hill (2000)
15. Wilson, E.O., Bossert, W.H.: Chemical communication among animals. *Recent Progress in Hormone Research*, 19:673-716 (1963)