

15th International Command & Control Research & Technology Symposium

“The Evolution of C2”

**Command & Control in Virtual Environments:
Tailoring Software Agents to Emulate Specific People**

Topics: 5, 6 and or 3

Danielle Martin Wynn, Evidence Based Research

Mary Ruddy, Parity Communications

Mark E. Nissen, US Naval Postgraduate School*

* Point of contact

Mary Ruddy

Parity Communications, Inc.

22 Bartlett Avenue, Arlington, MA 02476

Tel: (617) 290-8591, Fax: (617) 663-6165, E-mail: mary [at]parityinc.net

ABSTRACT

Development of and experimentation with ELICIT (Experimental Laboratory for Investigating Collaboration, Information-sharing and Trust) is an ongoing activity of the U.S. DoD Command and Control Research Program (CCRP) within the Office of the Assistant Secretary of Defense for Networks and Information Integration (OASD/NII). A recent CCRP-sponsored effort resulted in the development of a configurable sensemaking agent to enable agent-based ELICIT simulation experiments. A key step in the adoption of these configurable agents for use in research is demonstrating that these agents can be configured to behave as specific humans behave. This paper discusses how the behavior of humans in an actual ELICIT experiment is successfully modeled using sensemaking agents and provides suggestions for how the validated agents could be used in future ELICIT experiments.

Keywords: ELICIT, experimentation, edge organization, collaboration, agent

Acknowledgements: The research described in this article was funded in part by the Command and Control Research Program through the Center for Edge Power at the Naval Postgraduate School.

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE JUN 2010		2. REPORT TYPE		3. DATES COVERED 00-00-2010 to 00-00-2010	
4. TITLE AND SUBTITLE Command & Control in Virtual Environments: Tailoring Software Agents to Emulate Specific People				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Naval Postgraduate School,Center for Edge Power,589 Dyer Road,Monterey,CA,93943				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES Proceedings of the 15th International Command & Control Research & Technology Symposium, Santa Monica, CA (June 2010).					
14. ABSTRACT Development of and experimentation with ELICIT (Experimental Laboratory for Investigating Collaboration, Information-sharing and Trust) is an ongoing activity of the U.S. DoD Command and Control Research Program (CCRP) within the Office of the Assistant Secretary of Defense for Networks and Information Integration (OASD/NII). A recent CCRP-sponsored effort resulted in the development of a configurable sensemaking agent to enable agent-based ELICIT simulation experiments. A key step in the adoption of these configurable agents for use in research is demonstrating that these agents can be configured to behave as specific humans behave. This paper discusses how the behavior of humans in an actual ELICIT experiment is successfully modeled using sensemaking agents and provides suggestions for how the validated agents could be used in future ELICIT experiments.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 34	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

INTRODUCTION

Modern military organizations have adapted and evolved over many centuries and millennia, respectively. Hierarchical command and control (C2) organizations in particular have been refined longitudinally (e.g., through iterative combat, training and doctrinal development) to become very reliable and effective at the missions they were designed to accomplish. However, recent research suggests that the Hierarchy may not represent the best organizational approach to C2 in all circumstances (Nissen, 2005), particularly where the environment is unfamiliar or dynamic. Indeed, alternate, more flexible C2 organizational approaches such as the Edge have been proposed (Alberts & Hayes, 2003) to overcome Hierarchy limitations, but the same recent research suggests that the Edge may not represent the best organizational approach to C2 in all circumstances either, particularly where the environment is familiar and stable.

Of course, the Hierarchy and Edge both represent organizational archetypes (Orr & Nissen, 2006), each of which offers considerable latitude in terms of detailed organizational design and customization. For instance, recent research demonstrates further how the performance of both Hierarchy and Edge organizations is sensitive to factors such as network infrastructure, professional competency and other interventions that can be affected through leadership, management and investment (Gateau, Leweling, Looney, & Nissen, 2007). With incessant advances in information technology (IT) that appear to be continuing, one may be able to overcome the limitations inherent in Hierarchy, Edge or other organizations or even enable such organizations to adapt—through IT—to shifting conditions.

This notion is fundamental to Network Centric Operations (NCO), where people and organizations operate principally in network-enabled virtual environments as opposed to their physical counterparts. Unfortunately, empirical evidence to support the asserted superiority of NCO remains sparse, and the capability enhancing properties of virtual environments remain more in the domain of lore than empirical assessment. Indeed, drawing from substantial research in both Educational Psychology and Media Richness, our parallel paper (Bergin et al., 2010) presents the counter argument that performance in virtual environments may be *worse* than in physical counterparts, and such counter argument offers substantial merit and empirical support. Hence we find some controversy between the tenets of NCO and empirical evidence in related fields.

Building upon these separate streams of research, we continue a campaign of experimentation to assess the relative performance of different C2 organizational approaches across a diversity of environments and conditions. In this present study, we investigate explicitly the ability to develop semi-intelligent, sensemaking software agents to perform the kinds of information-sharing and processing tasks reserved historically for people. Specifically, capitalizing upon the excellent internal validity and control available, we begin a series of laboratory experiments to assess the relative performance of human and software agents. In this first phase, we tailor a set of software agents to emulate the performance of specific people in a counterterrorism intelligence organizational context.

In the balance of the paper, we provide background on the software agents and the virtual environment in which they perform. We then detail our activities to tailor such agents to emulate the performance of specific people who share and process information

in a comparable environment, and we report in turn the results of a head-to-head experiment with people and software agents performing the same set of tasks in the same task environment. The paper closes with a set of conclusions, recommendations for practice, and topics for future research along the lines of this campaign.

BACKGROUND

In this section we provide background on the software agents and the virtual environment in which they perform. This background begins with an overview of the ELICIT (Experimental Laboratory for Investigating Collaboration, Information-Sharing and Trust) multiplayer online counterterrorism intelligence game that serves as an instrumented platform for experimentation. The discussion continues with details regarding abELICIT, through which semi-intelligent, sensemaking software agents have been developed to play the ELICIT game, both autonomously and in conjunction with people. This background draws considerably from (Ruddy, Martin, & McEver, 2009).

ELICIT

The United States Department of Defense Command and Control Research Program (CCRP) of the Office of the Assistant Secretary of Defense for Networks and Information Integration (OASD/NII) is engaged in developing and testing principles of organization that enable transformation from traditional hierarchy-based command and control practices toward the transference of power and decision rights to the edge of the organization. The need for agility in Information Age militaries is becoming increasingly important. As discussed in *Understanding Command and Control* (Alberts & Hayes, 2006), in an era of complex, coalition, civil-military operations, understanding how to organize for agility not just within a specific organization but also across differing organizations and cultures is a key to success.

As noted above, there has been a shortage of formal experimentation data on the efficacy of different C2 organizational approaches. In order to remedy such shortage, the CCRP created and continues to sponsor and maintain the ELICIT experimentation environment. ELICIT is a Java-based software platform that can be used to run multi-user experiments focused on information, cognitive, and social domain phenomena. People participate in experiment sessions mediated by ELICIT by working together in teams that can be configured to reflect different organizational approaches (e.g., Hierarchy, Edge, others) and that can be subjected to a wide variety of experiment controls and manipulations.

This ELICIT experimentation platform has configurable scenarios that focus on the task of discovering the “who”, “what”, “where”, and “when” of a fictitious terrorist plot. Information in the form of “factoids” is provided periodically to each of the participants during an experiment session. The factoids and their distribution are structured so that no one participant receives all the information necessary to perform the task; thus, information sharing is required in order for any participant to be able to determine a solution to the ELICIT problem.

ELICIT provides an instrumented task environment that captures and time stamps participants’ information sharing activities. The environment generates detailed transaction logs summarizing such information; these, together with participant surveys

that can be administered either prior to a trial (for calibration), after a trial, or *in situ*, can be used to measure information sharing, collaboration behaviors and situational awareness, as well as a variety of value metrics including the ability of individuals and teams to correctly identify the future adversary attack and the time required to do so. Considerable research has been conducted to date using ELICIT (Leweling & Nissen, 2007; Powley & Nissen, 2009), and the interested reader is directed to the corresponding references for details and results.

abELICIT

The instantiation of semi-intelligent, sensemaking agents for use in ELICIT experimentation (abELICIT) expands the range of experiments immensely, and it enables a whole new campaign of experiments involving such agents, either in lieu of or in conjunction with human participants.

The goal of the design process for the sensemaking agents is to describe and instantiate a semi-intelligent, configurable agent. This agent not only needs to be able to take the place of a human participant in ELICIT experiments, but to actually form a mental model of the information in the factoids received and of the members of the group in which it operates. That is, as an agent participates in an experiment it needs to generate situational awareness that can be drawn upon to make decisions about behavior. The agent's behavior must also depend on the scenario in which it is operating, so that it behaves differently, as appropriate, in different scenarios (i.e., the agent must not be scripted; that is, it must be able to respond to the scenario as it unfolds according to its "personality"). The sensemaking agent needs to be able to formulate ELICIT messages based on awareness and understanding of the factoids to which it has been exposed.

A sub-objective of the sensemaking agent is that people interacting with properly configured sensemaking agents as part of ELICIT experiment trials should not be able to tell that some of the experiment participants are software agents: the agents should pass a "Turing test" ¹ in the ELICIT context. Since the interaction of ELICIT software agents with human participants is limited to those actions that are allowed by the ELICIT 2.1.1 platform (e.g., sharing factoids, posting factoids to information websites, receiving factoids from other participants), the software agent behavior needed is much more narrowly focused than that of any agent that would actually compete in a more broadly defined (and more traditional) Turing test.

In addition to having this awareness or sensemaking capability, the agents also have additional configurable variables to define their personalities and styles of social interaction with the other experiment participants. Using these variables, agents can be configured to operate in human timeframes (e.g., seconds and minutes), rather than just

¹ The phrase Turing test classically refers to a test of human intelligence proposed by Alan Turing in his 1950 paper "Computing Machinery and Intelligence". Wikipedia describes the Turing test as "A human judge engages in a natural language conversation with one human and one machine, each of which tries to appear human. All participants are placed in isolated locations. If the judge cannot reliably tell the machine from the human, the machine is said to have passed the test. In order to test the machine's intelligence rather than its ability to render words into audio, the conversation is limited to a text-only channel such as a computer keyboard and screen."

computer timeframes (e.g., nanoseconds and microseconds), show human levels of variability and human personality traits such as tendencies to hoard information, reciprocate favors and trust team members (among others).

In terms of architecture, Figure 1 delineates a high-level view of the sensemaking agent logic flow. The sense making agent explicitly models the major mental steps that humans take when performing ELICIT tasks. The model accounts for people taking these steps in differing orders and in varying time frames. Interested readers are directed to (Ruddy et al., 2009) for details.

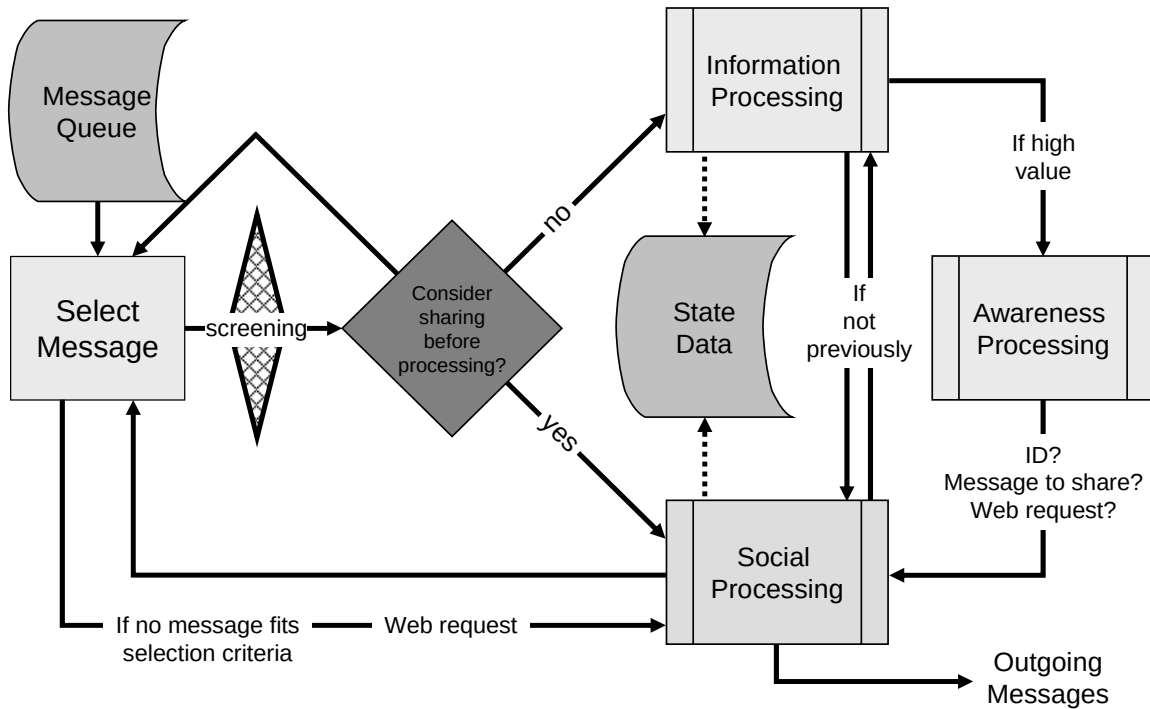


Figure 1. High-Level View of Sensemaking Agent Logic Flow

It is important to note that the sensemaking agent behaviors have been validated and calibrated in part through comparison with people's behaviors in numerous ELICIT experiments. Indeed, previous ELICIT based experiments provide a rich set of data on which to model agent behavior. These data include transaction logs that record and timestamp every action taken by each participant, scratch paper used by experiment participants, and information gleaned from participants in post experiment discussion sessions. This information and experience provide a strong basis for modeling human information-sharing and information-processing behaviors and illuminate insight into the mental models used by people when participating in the experiments. Through subsequent agent design refinement and comparison with ELICIT data, we develop semi-intelligent, sensemaking agents that emulate *generic* people relatively well.

However, the *specific* people participating in any given experiment will have their own sets of idiosyncratic behaviors, which may or may not match well with the generic behaviors developed and refined from the pool of previous ELICIT sessions. For direct comparison of software versus human agent performance, it is important to match the software agents' information-sharing and information-processing behaviors more closely

to those exhibited by the people participating in a particular experiment session. The abELICIT implementation includes nearly 50 configurable parameters that can be used to tailor the sensemaking agents' behavior. Hence each person who participates in an ELICIT experiment session can conceivably have a personal, corresponding software agent to emulate his or her behavior. This research represents the focus of the present study.

AGENT TAILORING APPROACH

We ran an ELICIT experiment at the Naval Postgraduate School (NPS) to examine people using ELICIT's customary, textual, network-based interface in comparison with physical, face-to-face interactions (Bergin et al., 2010). Participants completed a survey to solicit information about their behaviors during the experiment. The survey can be found in Appendix A. The survey responses, combined with the ELICIT transaction log file which records and timestamps all actions performed by participants using the software, are used to design individual agents whose behaviors match those of different, individual people. While the log file contains a great deal of information about actions performed by human subjects within the experiment, it is necessary to supplement the data with the survey responses; many behavior characteristics, such as whether a person reads newly received factoids before older factoids or how confident a person is in his identification attempt, cannot be gleaned from the transaction log. In comparing the human log file and corresponding surveys, people did not always accurately describe their behaviors in the experiment, and so where tension exists between the log file and survey data, we reference log file data.

We use these inputs to configure ELICIT agents to behave as specific humans behave. Our iterative process begins by populating the agent design parameters using the information collected. The following agent parameters, defined in Appendix B, are relevant to this work and discussed below:

messageQueueNewerBeforeOlder	shareWith
shareBeforeProcessing	shareWithWebSites
timeBeforeFirstIdentify	propensityToSeek
minSolutionAreas	minTimeBetweenPulls
hasSeenEnoughToIdentify	primary
idConfidencelevel	secondary
propensityToShare	awarenessProcessingThreshold
shareModalChoice	

As describe earlier a large number of agent parameters can be tuned within abELICIT, however we focus only on this subset and use the default setting for all other parameters. A sample agent configuration file can be found in Appendix C. Below we describe how we populate each of the agent settings used in this effort.

Participants in the human trial are instructed to identify only once within the ELICIT experiment. While players operate as part of a task area group, each can attempt to identify all four areas of the problem and so the *primary* agent parameter was set to *who*, *what*, *where*, *when* for all agents and thus agents have no *secondary* area of interest.

Analyzing the human trial log file, we are able to determine the appropriate agent settings for many parameters. The agent settings for *timeBeforeFirstIdentify* and *minSolutionAreas* reflect recorded values.

hasSeenEnoughToIdentify is set using both the survey data and log file. Survey answers are used to set this agent parameter unless the human participant had access to fewer factoids than listed in their survey answer, in which case the survey answer option corresponding to the number of factoids the individual actually had access to, and may have potentially seen is used.

Looking at the average time which passes from when each individual human receives a factoid through distribution until they Share or Post it, we are able to set each agent's *shareBeforeProcessing* value. Only one human, with pseudonym "Leslie," shares all four of the factoids received through distribution in a timely manner (averaging 37 seconds after receipt, while all other agents took more than 2 minutes if they even shared the factoid at all). The *shareBeforeProcessing* setting for the agent corresponding to Leslie is set to *true* while all other agents are set to *false*.

Using the log file and the ELICIT Log Analyzer (CCRP, 2009) the team analyzes the number of direct sharing and posting events the humans perform. If an individual does not perform any direct sharing events, the *shareModalChoice* setting for the corresponding agent is set to *post only*; likewise if an individual did not Post any factoids, the *shareModalChoice* setting for the corresponding agent is set to *peer to peer only*. If the number of direct Shares outweighs the number of Posts, the corresponding agent is set to *peer to peer dominant*. If the number of Posts outweighs the number of direct Shares or if it is reasonably larger than the number of Posts performed by the other participants, the corresponding agent is set to *post dominant*. Finally we conclude that the agents corresponding to all remaining participants have a *shareModalChoice* setting of *both*.

We also consider the list of websites to which an agent may Post. If we assign an agent a *shareModalChoice* setting other than *peer to peer only*, in this Hierarchy trial, its *shareWithWebSites* is set to those website(s) which the agent has access to (as aligned with their task areas) as dictated by the ELICIT organization file. Agents with a *shareModalChoice* set to *peer to peer only* do not have any websites on their *shareWithWebSites* list.

For those agents that do not have a *shareModalChoice* setting of *post only*, *shareWith* settings are assigned. The assignment of whom an agent must Share with is achieved by reviewing which other participants each human subject shared with. In the human experiment, subjects are able to Share factoids with multiple recipients in a single action; however the agents perform each direct Share as a separate action.

Analysis of the number of times each individual views a website (referred to as Pull) is used to determine the *propensityToSeek* setting for each corresponding agent. Using the ELICIT Log Analyzer we can see how many times a human pulls a given website (noting that in this trial pseudonym "Alex," the Cross Team Coordinator, has access to all four websites, while the other individuals only have access one of the four websites). All of the human participants have a relatively low number of Pulls with exception of pseudonym "Val," therefore the agent corresponding to Val is given a *propensityToSeek* setting of *moderate*, while all other agents are set to *low*.

The ELICIT Log Analyzer allows us to view Pulls over time to observe the frequency with which they occur in the human trial. In doing so, we set the agents' *minTimeBetweenPulls* to represent these frequencies.

Agent parameters settings for *messageQueueNewerBeforeOlder* and *idConfidencelevel* are taken directly from the survey answers.

For the two parameters which cannot be directly inferred from the log file and survey data, *awarenessProcessingThreshold* and *propensityToShare*, the team relies on its experience with abELICIT and its familiarity with the parameter rule sets to infer human-like settings. All agents are given an *awarenessProcessingThreshold* setting of 2, and a *propensityToShare* setting of low.

The agent trial uses the same organization file as the actual human run; a hierarchical structure, with limited access to other participants and websites, composed of 17 nodes (humans or agents), three of which are not active throughout the run. In addition the same names, roles, and task areas are used in the agent design. It is important to note that only one agent factoid set exists (set 1) which is different than the factoid set used in the human trial (set 2).

Each trial produces a transaction log. As explained above, the actions within the agent transaction log are a measure of agent behavior. Agent and human actions are compared across several variables to evaluate the relative fit of each agent design. Adjustments are made to tune the agent parameters accordingly for a new trial. This iterative process continues until the results match or the differences are minimal and further tuning does not improve the results.

RESULTS

In this section we summarize results of our agent tailoring followed by a summary of key implications. Details of the agent tailoring are presented in Appendix D.

Agent Tailoring

Using the reasoning described above in the Agent Tailoring Approach, we develop an initial trial design. Given that the agents are configurable, an iterative approach is taken: comparing the agent behaviors to the observed human behaviors, and adjusting the agent parameter settings as appropriate in the subsequent trial design. The initial agent trial design is depicted in Table 1 below. Each of the configurable parameters is shown across the top of the table, while the corresponding agent settings are listed in the rows below.

Table 1 Trial 1 Agent Design

Agent Number	Agent Name	Organization Type	Role	Group	shareBeforeProcessing	shareModalChoice	propensityToShare	propensityToSeek	shareWith	shareWithWebsites	messageQueueNewerBeforeOlder	awarenessProcessingThreshold	hasSeenEnoughToIdentify	timeBeforeFirstIdentify	minSolutionAreas	idConfidencelevel	minTimeBetweenPulls	primary
1	Alex	Hierarchical	Coordinator		FALSE	peer to peer only	low	low	2,3,4,5		TRUE	2	20	42	3	0.50	90000	who,what,when,where
2	Chris	Hierarchical	Team leader	Who	FALSE	post only	low	low		Who	TRUE	2	68	1000	4	0.50	3600000	who,what,when,where
3	Dale	Hierarchical	Team leader	What	FALSE	peer to peer dominant	low	low	1,2,4,5,9,10,11	What	TRUE	2	20	45	4	0.50	1500000	who,what,when,where
4	Francis	Hierarchical	Team leader	Where	FALSE	post dominant	low	low	1,12,13,14	Where	TRUE	2	20	40	3	0.50	900000	who,what,when,where
5	Harlan	Hierarchical	Team member	When	FALSE	both	low	low			TRUE	2	20	1000	4	1.00	3600000	who,what,when,where
6	Jesse	Hierarchical	Team member	Who	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
7	Kim	Hierarchical	Team member	Who	FALSE	post only	low	low		Who	TRUE	2	16	44	3	0.00	360000	who,what,when,where
8	Leslie	Hierarchical	Team member	Who	TRUE	both	low	low	2,6,7	Who	TRUE	2	16	40	3	0.25	900000	who,what,when,where
9	Morgan	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,10,11		TRUE	2	16	40	3	0.25	1500000	who,what,when,where
10	Pat	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,9,11		TRUE	2	16	1000	4	1.00	1200000	who,what,when,where
11	Quinn	Hierarchical	Team member	What	FALSE	peer to peer dominant	low	low	3,9,10	What	TRUE	2	16	45	3	0.50	3600000	who,what,when,where
12	Robin	Hierarchical	Team member	Where	FALSE	post only	low	low		Where	TRUE	2	16	46	4	0.50	2400000	who,what,when,where
13	Sam	Hierarchical	Team member	Where	FALSE	both	low	low	4,12,14	Where	TRUE	2	16	41	3	0.75	600000	who,what,when,where
14	Sidney	Hierarchical	Team member	Where	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
15	Taylor	Hierarchical	Team member	When	FALSE	post only	low	low		When	TRUE	2	6	49	4	0.50	600000	who,what,when,where
16	Val	Hierarchical	Team member	When	FALSE	peer to peer dominant	low	moderate	5,15,17	When	TRUE	2	6	49	3	0.25	60000	who,what,when,where
17	Whitley	Hierarchical	Team member	When	FALSE	post only	low	low		When	TRUE	2	6	46	3	0.50	2400000	who,what,when,where

Agent name refers to the participant pseudonym assigned to a particular agent in an experiment. Role refers to the position in the organization. In this case, the organization is a C2 hierarchy in which four Team Leaders, each with three Team Members of the same group or Task area, report to a cross team Coordinator. Those agents whose design parameters are shaded in grey in the table above represent the three absent/idle human players. This design yields the following results summarized in Table 2.

Table 2 Trial 1 Results Summary

Agent Trial 1 Results										
Agent Number	Agent	Agent Role	Agent Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		4	0	0	0	104	16	124
2	Chris	Team leader	Who	4	0	4	4	1	0	9
3	Dale	Team leader	What	4	0	4	4	2	28	38
4	Francis	Team leader	Where	4	0	4	4	4	16	28
5	Harlan	Team leader	When	4	0	0	0	1	0	5
6	Jesse	Team member	Who	4	0	0	0	1	0	5
7	Kim	Team member	Who	4	0	4	4	9	0	17
8	Leslie	Team member	Who	4	0	4	4	4	12	24
9	Morgan	Team member	What	4	0	0	0	2	12	18
10	Pat	Team member	What	4	0	0	0	3	12	19
11	Quinn	Team member	What	4	0	4	4	1	12	21
12	Robin	Team member	Where	4	0	4	4	2	0	10
13	Sam	Team member	Where	4	0	4	4	5	12	25
14	Sidney	Team member	Where	4	0	0	0	1	0	5
15	Taylor	Team member	When	4	0	4	4	5	0	13
16	Val	Team member	When	4	0	4	4	47	12	67
17	Whitley	Team member	When	4	0	4	4	2	0	10
Total	Total			68	0	44	44	194	132	438
Average	Average			4	0	2.59	2.59	11.41	7.76	25.765

The agent transaction log is compared to that of the human trial. The agent number, name, role and task are followed by a count of each type of event occurring during the trial. Note that all events with exception of *dist*, meaning the number of factoids received through distribution from the server, are actions performed by the agent. First Post is a count of how many times the agent/participant was the first one to *post* a factoid to a website. Given that these actions are already accounted for in the *post* count, First Post values are not part of the Total value reflected in the right hand column. Table 3 below displays the differences in the results (agent data subtracted from human data). From the discussion of our approach above, we seek to minimize such differences.

Table 3 Difference between Trial 1 and Human Trial

Difference Between Agent & Human Trial										
Number	Name	Role	Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		0	1	0	0	-50	13	-36
2	Chris	Team leader	Who	0	0	23	9	0	0	23
3	Dale	Team leader	What	0	1	5	0	-1	140	145
4	Francis	Team leader	Where	0	1	9	7	2	-4	8
5	Harlan	Team leader	When	0	0	0	0	-1	0	-1
6	Jesse	Team member	Who	0	0	0	0	-1	0	-1
7	Kim	Team member	Who	0	1	0	0	-1	0	0
8	Leslie	Team member	Who	0	1	1	0	0	-3	-1
9	Morgan	Team member	What	0	1	0	0	-1	0	0
10	Pat	Team member	What	0	0	0	0	0	0	0
11	Quinn	Team member	What	0	1	-1	-1	-1	24	23
12	Robin	Team member	Where	0	1	4	-1	-1	0	4
13	Sam	Team member	Where	0	1	-3	-4	4	-6	-4
14	Sidney	Team member	Where	0	0	0	0	-1	0	-1
15	Taylor	Team member	When	0	1	-2	-2	7	0	6
16	Val	Team member	When	0	1	-3	-4	0	-1	-3
17	Whitley	Team member	When	0	1	0	-2	-1	0	0
Total				0	12	33	2	-46	163	162
Average				0	0.71	1.94	0.12	-2.7	9.59	9.529

Notice that the Total row reflects considerable variation in differences. For instance, the Identify column shows a total difference of 12 between the software and human agent runs, whereas the Post column to its right shows a larger difference of 33, but the First Post column shows a tiny difference of only 2. The difference reflected in the Pull column (-46) is comparatively larger, and the Share difference (163) is comparatively very large. Since differences include both positive and negative values, they tend to cancel one another, and the Total column (162) summarizes the effect of such canceling. Likewise, the values shown in the Total column for each player reflect considerable variation also. For instance, the total differences for players Alex (-36), Chris (23) and Quinn (23) all reflect double-digit values, and differences for Dale (145) reflects triple digits, whereas all of the other differences are in single digits. The data summarized in this table help us to focus our tailoring.

Table 4 Trial 6 Agent Design

Agent Number	Agent Name	Organization Type	Role	Group	shareBeforeProcessing	shareModalChoice	propensityToShare	propensityToSeek	shareWith	shareWithWebSites	messageQueueNewerBeforeOlder	awarenessProcessingThreshold	hasSeenEnoughToIdentify	timeBeforeFirstIdentify	minSolutionAreas	idConfidencelevel	minTimeBetweenPulls	primary
1	Alex	Hierarchical	Coordinator		FALSE	peer to peer only	moderate	low	2,3,4,5		TRUE	2	20	15	1	0.25	150000	who,what,when,where
2	Chris	Hierarchical	Team leader	Who	FALSE	post only	moderate	low		Who	TRUE	2	68	1000	1	0.25	3600000	who,what,when,where
3	Dale	Hierarchical	Team leader	What	FALSE	peer to peer dominant	moderate	low	1,2,4,5,9,10,11		TRUE	2	20	15	1	0.25	1500000	who,what,when,where
4	Francis	Hierarchical	Team leader	Where	FALSE	post dominant	low	low	1,12,13,14	Where	TRUE	2	20	15	1	0.25	900000	who,what,when,where
5	Harlan	Hierarchical	Team leader	When	FALSE	both	low	low			TRUE	2	20	1000	4	1.00	3600000	who,what,when,where
6	Jesse	Hierarchical	Team member	Who	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
7	Kim	Hierarchical	Team member	Who	FALSE	post only	low	low		Who	TRUE	2	16	15	1	0.00	360000	who,what,when,where
8	Leslie	Hierarchical	Team member	Who	TRUE	both	moderate	low	2,6,7	Who	TRUE	2	16	15	1	0.25	900000	who,what,when,where
9	Morgan	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,10,11		TRUE	2	16	15	1	0.25	1500000	who,what,when,where
10	Pat	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,9,11		TRUE	2	16	1000	1	1.00	1200000	who,what,when,where
11	Quinn	Hierarchical	Team member	What	FALSE	peer to peer dominant	low	low	3,9,10	what	TRUE	2	16	15	1	0.25	3600000	who,what,when,where
12	Robin	Hierarchical	Team member	Where	FALSE	post only	moderate	low			TRUE	2	10	15	1	0.25	2400000	who,what,when,where
13	Sam	Hierarchical	Team member	Where	FALSE	peer to peer dominant	low	low	4,12,14	Where	TRUE	2	10	15	1	0.25	600000	who,what,when,where
14	Sidney	Hierarchical	Team member	Where	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
15	Taylor	Hierarchical	Team member	When	FALSE	post only	low	low		When	TRUE	2	6	15	1	0.25	360000	who,what,when,where
16	Val	Hierarchical	Team member	When	FALSE	peer to peer dominant	moderate	moderate	5,15,17	When	TRUE	2	6	15	1	0.25	60000	who,what,when,where
17	Whitley	Hierarchical	Team member	When	FALSE	post only	moderate	low		When	TRUE	2	6	15	1	0.25	2400000	who,what,when,where

With each iteration, we analyze such differences, modify the software agents, compare the subsequent results, and so forth until differences have been minimized. In total, six iterations are required to tailor the software agents to reduce differences sufficiently and hence emulate the performance of their human counterparts closely. Indeed, any further changes beyond those incorporated into the sixth iteration produce results with *greater differences*, so we stop at that point. For comparison with the initial agent design summarized in Table 1 above, the final corresponding design for trial six is shown in Table 4.

Likewise Table 5 summarizes results from trial six for direct comparison with those reported in Table 2 above, and differences are displayed in Table 6 for direct comparison with those reported in Table 3. Notice how small the differences in Table 6 are. Although the row totals for three players (i.e., Alex, Dale, Quinn) remain in double digits, most are in single digits, and the column totals are all less than or equal to one. This reflects substantial agent tailoring, and the corresponding design appears to mimic the actions performed by participants in the human trial quite well.

Table 5 Trial 6 Results Summary

Agent Trial 6 Results										
Agent Number	Agent Name	Agent Role	Agent Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		4	4	0	0	44	76	128
2	Chris	Team leader	Who	4	0	32	28	1	0	37
3	Dale	Team leader	What	4	2	0	0	2	148	156
4	Francis	Team leader	Where	4	1	4	4	4	16	29
5	Harlan	Team leader	When	4	0	0	0	1	0	5
6	Jesse	Team member	Who	4	0	0	0	1	0	5
7	Kim	Team member	Who	4	3	4	4	9	0	20
8	Leslie	Team member	Who	4	4	4	4	4	12	28
9	Morgan	Team member	What	4	1	0	0	2	12	19
10	Pat	Team member	What	4	0	0	0	3	12	19
11	Quinn	Team member	What	4	1	4	4	1	12	22
12	Robin	Team member	Where	4	0	5	2	2	0	11
13	Sam	Team member	Where	4	0	4	4	5	12	25
14	Sidney	Team member	Where	4	0	0	0	1	0	5
15	Taylor	Team member	When	4	0	4	4	9	0	17
16	Val	Team member	When	4	0	4	4	47	12	67
17	Whitley	Team member	When	4	2	8	4	2	0	16
Total				68	18	73	62	138	312	609
Average				4	1.06	4.29	3.65	8.12	18.35	35.82

Table 6 Difference between Trial 6 and Human Trial

Difference Between Agent & Human Trial										
Number	Name	Role	Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		0	-3	0	0	10	-47	-40
2	Chris	Team leader	Who	0	0	-5	-15	0	0	-5
3	Dale	Team leader	What	0	-1	9	4	-1	20	27
4	Francis	Team leader	Where	0	0	9	7	2	-4	7
5	Harlan	Team leader	When	0	0	0	0	-1	0	-1
6	Jesse	Team member	Who	0	0	0	0	-1	0	-1
7	Kim	Team member	Who	0	-2	0	0	-1	0	-3
8	Leslie	Team member	Who	0	-3	1	0	0	-3	-5
9	Morgan	Team member	What	0	0	0	0	-1	0	-1
10	Pat	Team member	What	0	0	0	0	0	0	0
11	Quinn	Team member	What	0	0	-1	-1	-1	24	22
12	Robin	Team member	Where	0	1	3	1	-1	0	3
13	Sam	Team member	Where	0	1	-3	-4	4	-6	-4
14	Sidney	Team member	Where	0	0	0	0	-1	0	-1
15	Taylor	Team member	When	0	1	-2	-2	3	0	2
16	Val	Team member	When	0	1	-3	-4	0	-1	-3
17	Whitley	Team member	When	0	-1	-4	-2	-1	0	-6
Total				0	-6	4	-16	10	-17	-9
Average				0	-0.35	0.24	-0.94	0.59	-1	-0.53

Key Implications

This work illustrates the ability to tailor the behaviors of individual agents within abELICIT to match those of different, individual people. While it takes a considerable number of parameter adjustments and several iterations to tailor the software agents so

their performance is sufficiently close to that of human participants, when comparing the results of the head-to-head experiments, each of the people participating in an ELICIT session has a corresponding abELICIT sensemaking agent which emulates his or her behavior in a parallel session.

Although somewhat technical in nature, this work extends the ELICIT experimentation infrastructure to enable specific human players to be emulated by software agents, and it paves the way for a whole new series of experiments, thereby enriching the campaign of experimentation. Nonetheless, there are several limitations to this work that are important to note. Specifically, the following differences between human and agent experiments affect our results directly:

- As mentioned earlier, in the human experiment, subjects are able to Share factoids with multiple recipients in a single action, while agents do not have this capability and must perform each direct Share as a separate action.
- Additional actions occur in the trials. *abSelect* is an action only performed by agents. In addition, the *add* action has different significance in human trials and agent runs as humans selectively *add* factoids to their MyFactoids list while agents must *add* factoids in order to process the information. Therefore *abselect* and *add* are not relevant when comparing the data to the human trial data.
- The human trial instructed participants to ID only once. In the agent trials we do not limit the number of ID attempts. This work uncovered a minor, but necessary code change which explains why we needed to make so many adjustments to the variables related to identification. Currently agents check to see whether they should ID when processing information, however, if an agent's *timeBeforeFirstIdentify* is set late in the trial and he does not receive any new information for processing after the *timeBeforeFirstIdentify* mark, the agent will never check to see whether he should ID. In the ELICIT 2.3 code revision, an agent is triggered to check whether he can ID when the experiment has lapsed beyond the *timeBeforeFirstIdentify*.
- Only one agent factoid set exists (factoid set 1), the human trial was conducted using factoid set 2, and so it is a bit difficult to tune the agents using a different fact set. That said we are currently translating factoid set 2 for use in future agent trials.

Clearly, addressing these differences through continued research along the lines of this investigation offers promise.

CONCLUSION

Development of and experimentation with ELICIT is an ongoing activity of the CCRP. A recent CCRP-sponsored effort resulted in the development of a configurable sensemaking agent to enable agent-based ELICIT simulation experiments. A key step in the adoption of these configurable agents for use in research is demonstrating that these agents can be configured to behave as specific humans behave. This paper discusses how the behavior of humans in an actual ELICIT experiment is successfully modeled using

sensemaking agents and provides suggestions for how the validated agents could be used in future ELICIT experiments.

This effort demonstrates that it is possible to configure specific ELICIT agents to emulate the behavior of specific individuals participating in an ELICIT experiment with a particular organizational structure. The completion of this validation exercise represents a key milestone in the evolution of the ELICIT platform. Validated ELICIT agents can be used to run more ELICIT experiments than is practical with human participants due to time and money constraints.

Validated ELICIT agents can also be used to run experiments on the effects of certain behaviors on overall group performance (e.g. which organization structures are more robust when there are many information hoarders present). These types of experiments are difficult to conduct as it is challenging to control for actual human behaviors. Agent and human experiment trials can also be used iteratively, as part of a campaign of experimentation. A number of alternative hypotheses can be explored relatively quickly using agents, and then the most promising hypotheses can be validated using human participants.

REFERENCES

- Alberts, D. S., & Hayes, R. E. (2003). *Power to the edge : Command and control in the information age*. Washington, DC: Command and Control Research Program.
- Alberts, D. S., & Hayes, R. E. (2006). *Understanding command and control*. Washington, D.C.: CCRP Publications.
- Bergin, R. D., Adams, A. A., Andraus, R., Hudgens, B. J., Lee, J. G., & Nissen, M. E. (2010). Command and control in virtual environments: Laboratory experimentation to compare virtual with physical. Santa Monica, CA.
- CCRP. (2009). *ELICIT log analyzer*. Retrieved 2009 <http://www.dodccrp.org/html4/elicit.php>
- Gateau, J. B., Leweling, T. A., Looney, J. P., & Nissen, M. E. (2007). Hypothesis testing of edge organizations: Modeling the C2 organization design space. *Proceedings International Command & Control Research & Technology Symposium*, Newport, RI.
- Leweling, T. A., & Nissen, M. E. (2007). Hypothesis testing of edge organizations: Laboratory experimentation using the ELICIT multiplayer intelligence game. *Proceedings International Command & Control Research & Technology Symposium*, Newport, RI.
- Nissen, M. E. (2005). Hypothesis testing of edge organizations: Specifying computational C2 models for experimentation. *Proceedings International Command & Control Research Symposium*, McLean, VA.
- Orr, R. J., & Nissen, M. E. (2006). Computational experimentation on C2 models. *Proceedings International Command and Control Research and Technology Symposium*, Cambridge, UK.
- Powley, E. H., & Nissen, M. E. (2009). Trust-mistrust as a design contingency: Laboratory experimentation in a counterterrorism context. Washington, DC.
- Ruddy, M., Martin, D., & McEver, J. (2009). Instantiation of a sensemaking agent for use with ELICIT experimentation. Washington, DC.

APPENDIX A: SURVEY

	ELICIT Survey
---	------------------------

1. When playing the experiment, what organization type were you a part of?

☐	Hierarchy Organization
☐	Edge Organization

2. If you were a part of a hierarchy organization, what was your role during the experiment?

☐	Team Member
☐	Team Leader
☐	Cross Team Coordinator

3. When I was participating in the experiment, I felt it was important to share my unique facts received from distribution with: (Scale: 1 = Completely Disagree, 2 = Disagree, 3 = Neutral, 4 = Agree, 5 = Completely Agree)

My team members	1	2	3	4	5
My team leader	1	2	3	4	5
The cross team coordinator	1	2	3	4	5
Everyone I could share with	1	2	3	4	5
Websites I could post to	1	2	3	4	5

4. During an ELICIT experiment, sharing information is important: (Scale: 1 = Completely Disagree, 2 = Disagree, 3 = Neutral, 4 = Agree, 5 = Completely Agree)

1 2 3 4 5

5. During your experience with ELICIT, how often did you share the following: (Scale: 1 = Never, 2 = Rarely, 3 = Sometimes, 4 = Frequently, 5 = Always)

New factoids you received by distribution	1	2	3	4	5
Facts received by another individual	1	2	3	4	5
Facts I saw on a website	1	2	3	4	5

6. When you share factoids during the experiment, how did you share the factoids? (Check all that apply.)

☐	I share the fact with my team members.
☐	I share the fact with my team leaders.
☐	I share the fact with the cross team coordinator.
☐	I share the fact with everyone I can.
☐	I post the fact to the relevant website.
☐	I did not share any factoids.

7. How soon in the experiment did you make your first identification attempt?

☐	0 – 10 minutes
☐	11 – 20 minutes
☐	21 – 30 minutes
☐	31 – 40 minutes
☐	41 – 50 minutes
☐	51 – 60 minutes
☐	I did not make any identification attempts.

8. Was your first identification attempt full or partial? If partial, how many areas did you identify?

☐	Full identification (I identified all areas)
☐	3 areas
☐	2 areas
☐	1 area
☐	I did not make any identification attempts.

9. How confident were you in your identification attempts? *(Scale 1 =Not at all confident, 3 = Somewhat confident, 5 = Very confident)*

1 2 3 4 5

10. How many factoids have you seen before your first identification attempt?

☐	0 – 5
☐	6 – 10
☐	11 – 15
☐	16 – 20
☐	More than 20 factoids
☐	I did not make any identification attempts.

11. When you check a website, how would you describe your behavior?
(Check all that apply.)

☐	I look at the newest information on the website.
☐	I look at the oldest information on the website.
☐	I look at all the information on the website.
☐	I look at only a few of the facts on the website.
☐	I look at the majority of the facts on the website.
☐	I did not look at any websites.

12. How often do you visit a website (by clicking a website tab)?

☐	Frequently (several times per minute)
☐	Often (every few minutes)
☐	Infrequently (every ten minutes)
☐	Never

13. Under what circumstances do you think it is appropriate to make an identification attempt?

--

APPENDIX B: PARAMETERS DEFINED

messageQueueNewerBeforeOlder|If true then newer messages are selected before older
[If false, then older messages are selected first. If true, then most recent message is selected first.]

shareBeforeProcessing|If true then share message before Processing
[If true, perform Share Social Processing before Information Processing the message. If False, then do Information Processing before Share Processing. Default is False.]

timeBeforeFirstIdentify|Time before the agent does its first identity (in minutes)
[Minimum number of minutes that need to pass before the agent does its first Identify. Default is 10 minutes]

minSolutionAreas|The minimum number of ID tables with some data
[The baseline ELICIT scenarios have 4 information areas (who, what, where and when.) If minSolutionAreas is set to 3, then agent must know something about 3 of these areas before it performs an Identify action.]

hasSeenEnoughToIdentify|HasSeenEnoughToIdentify
[Only identify if number of messages selected from queue is greater or equal to this number.]

idConfidencelevel|IdConfidencelevel
[Values range from 0-1. Agent only provides an answer in an Identify Action if the value in the relevant state table is greater or equal to this value.]

propensityToShare|PropensityToShare (possible values: low, moderate, high, very high)
[Controls agent's willingness to Share information. For example, if propensity to share is very high, agent Posts a message to all websites for which agent has permissions, sends a message to all entities where trust in recipient is not distrust, and sends a message to all entities with must_share flag set to on. If propensity to share is low, then agent sends a message only to all entities with must_share flag set to on. Default is low, and its effects depend on the value of the shareModalChoice variable.]

shareModalChoice|ShareModalChoice possible values (both, post dominant, post only, peer to peer dominant, peer to peer only)
[This variable governs the interaction and tradeoffs between Sharing a message with an individual and making it available by Posting it to a website. Its effects depend on the value of the propensityToShare variable. For example if shareModelChoice is set to both, and propensity to share is high, then Post message to area (of message website) and send message to all entities where trust in recipient is not distrust or no opinion, and send message to entities (participants and websites) with must-share flag set to on. Default is both.]

shareWith|List of players with whom agent may share (-1 means share with all from organization configuration file)

[List of players, specified by player role number, with which the agent must Share information for certain shareModalChoice propensityToShare combinations. A value of -1 means Share with all players with whom you have Share capability (as specified in the organization configuration file.) Default is -1.]

shareWithWebSites|List of websites with whom agent must share

[List of websites to which the agent must Post information for certain shareModalChoice propensityToShare combinations. Note that website must also be available to the agent as specified in the organization configuration file.]

propensityToSeek|PropensityToSeek (possible values: low, moderate, high, very high)

[Propensity to seek information via Pull action. Default is high. If propensityToSeek is very high, then agent pulls from all available sites, else it only pulls from the agent's primary area. PropensityToSeek also controls the minimum time between Pulls. For example if propensityToSeek is moderate, the minimum time between pulls is 3 minutes.]

minTimeBetweenPulls|If the time since the last pull is not $\geq$ minTimeBetweenPulls, do not Pull (in milliseconds), -1 means ignoring this parameter)

[Minimum time interval in milliseconds that the agent delays between subsequent Pulls (requests of data from a website.) Default depends on value of propensityToSeek. For example if propensityToSeek is moderate, the default is 3 minutes. If you would like the value of propensityToSeek to control the spacing between pulls, then do not include this variable in the configuration.]

primary|Primary areas of interest. (possible values: who, what, where, when)

[Area(s) of interest on which the agent initially focuses.]

secondary|Secondary areas of interest. (possible values: who, what, where, when)

[Other area(s) of interest on which the agent may subsequently focus.]

awarenessProcessingThreshold|If cumulative value of the perceived message value is more or equal to this variable, then start awareness processing.

[Cumulative value of additional newly processed information that can accumulate before a new cycle of awareness processing is initiated.]

APPENDIX C: SAMPLE AGENT CONFIGURATION FILE

```
SenseMaking_Agent_52_1
<begin agent configuration parameters>
SenseMaking_Agent_1.jar
net.parityinc.ccrp.web.agent.impl.SenseMaking_Agent_1
readyIntervalDelay|Time interval to click Ready button|10000
messageQueueCapacity|Capacity of queue (-1 means unlimited)|-1
messageQueueTimeRemainInQueue|Time a factoid can remain in queue (-1 means
unlimited)|-1
messageQueueNewerBeforeOlder|If true then newer messages are selected before
older|true
selectMessageFromQueueDelay|Select message from queue delay|1000
shareBeforeProcessing|If true then share message before Processing|false
postedTypes|PostedTypes|who,what,where,when
postFactor|PostFactor|1
postOutOfArea|PostOutOfArea|false
shareWithFactor|ShareWithFactor|1
sharedTypes|SharedTypes|who,what,where,when
shareRelevantAccordingToSiteAccess|ShareRelevantAccordingToSiteAccess|false
shareAccordingToSiteAccess|ShareAccordingToSiteAccess|false
isCompetitiveHoarder|IsCompetitiveHoarder|false
pullFactor|PullFactor|1
timeBeforeFirstIdentify|Time before the agent does its first identity (in minutes)|42
minSolutionAreas|The minimum number of ID tables with some data|3
hasSeenEnoughToIdentify|HasSeenEnoughToIdentify|20
isGuesser|IsGuesser|false
isFrequentGuesser|IsFrequentGuesser|false
idConfidencelevel|IdConfidencelevel|0.50
partialIdentify|Identify if there are no some answers|true
propensityToShare|PropensityToShare possible values (low, moderate, high, very
high)|low
shareModalChoice|ShareModalChoice possible values (both, post dominant, post only,
peer to peer dominant, peer to peer only)|peer to peer only
screeningSelectedMessageDelay|Screening selected message (message processing)
delay|1000
informationProcessingDelay|Information Processing delay|3000
socialProcessingDelay|Social Processing delay|4000
sharingPostingMessageDelay|Sharing/Posting each Message delay|8000
awarenessProcessingDelay|Awareness Processing delay|3000
determiningKnowledgeNeedsDelay|Determining Knowledge Needs delay|3000
idAttemptDelay|ID Attempt delay|20000
webRequestDelay|Web Request (Pull)|9000
shareWith|List of players with whom agent may share (-1 means share with all from
organization configuration file)|2,3,4,5
shareWithWebSites|List of websites with whom agent must share|
```

propensityToSeek|PropensityToSeek possible values (low, moderate, high, very high)|low
 minTimeBetweenPulls|If the time since the last pull is not $\geq$ minTimeBetweenPulls, do not Pull (in milliseconds, -1 means ignoring this parameter)|90000
 minTimeBetweenShares|If the time since the last Share is not $\geq$ minTimeBetweenShares, the agent should wait before it Shares (in milliseconds, -1 means ignoring this parameter)|5000
 trustInIndividuals|TrustInIndividuals possible values (high, medium, distrust, no opinion)|1=no opinion,2=no opinion,3=no opinion,4=no opinion,5=no opinion,6=no opinion,7=no opinion,8=no opinion,9=no opinion,10=no opinion,11=no opinion,12=no opinion,13=no opinion,14=no opinion,15=no opinion,16=no opinion,17=no opinion
 trustInWebSites|List of initial values of Trust for web sites. Possible values (high, medium, distrust, no opinion)|who=medium,where=medium,what=medium,when=medium
 reciprocity|Reciprocity possible values (high, low, medium, na, none)|1=none,2=none,3=none,4=none,5=none,6=none,7=none,8=none, 9=none, 10=none,11=none,12=none,13=none,14=none,15=none,16=none
 primary|Primary areas of interest. Possible values: who, what, where, when)|who,what,when,where
 secondary|Secondary areas of interest. Possible values: who, what, where, when)|
 propensityToShareExternal|If message is not in area of interest, then agent shares it according to sharing preferences with probability = propensityToShareExternal|1
 awarenessProcessingThreshold|If cumulative value of the perceived message value is more or equal to this variable, then start awareness processing.|2
 pullBetweenSitesDelay|Pull between sites delay|1000
 postBetweenSitesDelay|Post between sites delay|500
 provideRelevance|Provide relevance for posted and shared messages|false
 provideTrust|Provide trust for posted and shared messages|false

APPENDIX D: AGENT TAILORING ITERATIONS

In this appendix we include details of our agent tailoring iterations. Using the reasoning described above in the Agent Tailoring Approach, we develop an initial trial design. Given that the agents are configurable, an iterative approach is taken: comparing the agent behaviors to the observed human behaviors, and adjusting the agent parameter settings as appropriate in the subsequent trial design. The initial agent trial design is depicted in Table 7 below. Each of the configurable parameters is shown across the top of the table, while the corresponding agent settings are listed in the rows below.

Table 7 Trial 1 Agent Design

Agent Number	Agent Name	Organization Type	Role	Group	shareBeforeProcessing	shareModalChoice	propensityToShare	propensityToSeek	shareWith	shareWithWebSites	messageQueueNewerBeforeOlder	awarenessProcessingThreshold	hasSeenEnoughToIdentify	timeBeforeFirstIdentify	minSolutionAreas	idConfidencelevel	minTimeBetweenPulls	primary
1	Alex	Hierarchical	Coordinator		FALSE	peer to peer only	low	low	2,3,4,5		TRUE	2	20	42	3	0.50	90000	who,what,when,where
2	Chris	Hierarchical	Team leader	Who	FALSE	post only	low	low		Who	TRUE	2	68	1000	4	0.50	3600000	who,what,when,where
3	Dale	Hierarchical	Team leader	What	FALSE	peer to peer dominant	low	low	,2,4,5,9,10,1	What	TRUE	2	20	45	4	0.50	1500000	who,what,when,where
4	Francis	Hierarchical	Team leader	Where	FALSE	post dominant	low	low	1,12,13,14	Where	TRUE	2	20	40	3	0.50	900000	who,what,when,where
5	Harlan	Hierarchical	Team member	When	FALSE	both	low	low			TRUE	2	20	1000	4	1.00	3600000	who,what,when,where
6	Jesse	Hierarchical	Team member	Who	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
7	Kim	Hierarchical	Team member	Who	FALSE	post only	low	low		Who	TRUE	2	16	44	3	0.00	360000	who,what,when,where
8	Leslie	Hierarchical	Team member	Who	TRUE	both	low	low	2,6,7	Who	TRUE	2	16	40	3	0.25	900000	who,what,when,where
9	Morgan	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,10,11		TRUE	2	16	40	3	0.25	1500000	who,what,when,where
10	Pat	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,9,11		TRUE	2	16	1000	4	1.00	1200000	who,what,when,where
11	Quinn	Hierarchical	Team member	What	FALSE	peer to peer dominant	low	low	3,9,10	What	TRUE	2	16	45	3	0.50	3600000	who,what,when,where
12	Robin	Hierarchical	Team member	Where	FALSE	post only	low	low		Where	TRUE	2	16	46	4	0.50	2400000	who,what,when,where
13	Sam	Hierarchical	Team member	Where	FALSE	both	low	low	4,12,14	Where	TRUE	2	16	41	3	0.75	600000	who,what,when,where
14	Sidney	Hierarchical	Team member	Where	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
15	Taylor	Hierarchical	Team member	When	FALSE	post only	low	low		When	TRUE	2	6	49	4	0.50	600000	who,what,when,where
16	Val	Hierarchical	Team member	When	FALSE	peer to peer dominant	low	moderate	5,15,17	When	TRUE	2	6	49	3	0.25	60000	who,what,when,where
17	Whitley	Hierarchical	Team member	When	FALSE	post only	low	low		When	TRUE	2	6	46	3	0.50	2400000	who,what,when,where

Agent name refers to the participant pseudonym assigned to a particular agent in an experiment. Role refers to the position in the organization. In this case, the organization is a C2 hierarchy in which four Team Leaders, each with three Team Members of the same group or Task area, report to a cross team Coordinator. Those agents whose design parameters are shaded in grey in the table above represent the three absent/idle human players. This design yields the following results summarized in Table 8:

Table 8 Trial 1 Results Summary

Agent Trial 1 Results										
Agent Number	Agent	Agent Role	Agent Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		4	0	0	0	104	16	124
2	Chris	Team leader	Who	4	0	4	4	1	0	9
3	Dale	Team leader	What	4	0	4	4	2	28	38
4	Francis	Team leader	Where	4	0	4	4	4	16	28
5	Harlan	Team leader	When	4	0	0	0	1	0	5
6	Jesse	Team member	Who	4	0	0	0	1	0	5
7	Kim	Team member	Who	4	0	4	4	9	0	17
8	Leslie	Team member	Who	4	0	4	4	4	12	24
9	Morgan	Team member	What	4	0	0	0	2	12	18
10	Pat	Team member	What	4	0	0	0	3	12	19
11	Quinn	Team member	What	4	0	4	4	1	12	21
12	Robin	Team member	Where	4	0	4	4	2	0	10
13	Sam	Team member	Where	4	0	4	4	5	12	25
14	Sidney	Team member	Where	4	0	0	0	1	0	5
15	Taylor	Team member	When	4	0	4	4	5	0	13
16	Val	Team member	When	4	0	4	4	47	12	67
17	Whitley	Team member	When	4	0	4	4	2	0	10
Total	Total			68	0	44	44	194	132	438
Average	Average			4	0	2.59	2.59	11.41	7.76	25.765

The agent transaction log is compared to that of the human trial. The agent number, name, role and task are followed by a count of each type of event occurring during the trial. Note that all events with exception of *dist*, meaning the number of factoids received through distribution from the server, are actions performed by the agent. First Post is a count of how many times the agent/participant was the first one to *post* a factoid to a website. Given that these actions are already accounted for in the *post* count, First Post values are not part of the Total value reflected in the right hand column. Table 9 below displays the differences in the results (agent data subtracted from human data). From the discussion of our approach above, we seek to minimize such differences.

Table 9 Difference between Trial 1 and Human Trial

Difference Between Agent & Human Trial										
Number	Name	Role	Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		0	1	0	0	-50	13	-36
2	Chris	Team leader	Who	0	0	23	9	0	0	23
3	Dale	Team leader	What	0	1	5	0	-1	140	145
4	Francis	Team leader	Where	0	1	9	7	2	-4	8
5	Harlan	Team leader	When	0	0	0	0	-1	0	-1
6	Jesse	Team member	Who	0	0	0	0	-1	0	-1
7	Kim	Team member	Who	0	1	0	0	-1	0	0
8	Leslie	Team member	Who	0	1	1	0	0	-3	-1
9	Morgan	Team member	What	0	1	0	0	-1	0	0
10	Pat	Team member	What	0	0	0	0	0	0	0
11	Quinn	Team member	What	0	1	-1	-1	-1	24	23
12	Robin	Team member	Where	0	1	4	-1	-1	0	4
13	Sam	Team member	Where	0	1	-3	-4	4	-6	-4
14	Sidney	Team member	Where	0	0	0	0	-1	0	-1
15	Taylor	Team member	When	0	1	-2	-2	7	0	6
16	Val	Team member	When	0	1	-3	-4	0	-1	-3
17	Whitley	Team member	When	0	1	0	-2	-1	0	0
Total				0	12	33	2	-46	163	162
Average				0	0.71	1.94	0.12	-2.7	9.59	9.529

Reviewing the results of the original design, parameter setting adjustments are made to tune the agents within reason. In order to boost sharing the *propensityToShare* setting of each agent is raised from low to moderate. No identification attempts occur in the agent trial. Several parameters influencing the identification action are changed. First, *timeBeforeFirstIdentify* is decreased to forty minutes for all agents. Second, the *minSolutionAreas* needed to identify is dropped to 1. Finally, the max *idConfidencelevel* was changed to 0.49 for all agents previously set to 0.5 or greater as predetermined by the survey answer. Each of these changes is made to encourage agent identification as it is a primary objective for ELICIT participants. Please note that five subjects in the human experiment do not identify (the three absent/idle players, Chris, and Pat), therefore the ID related parameters for the agents representing these humans remain unchanged thus preventing the agent from identifying. The new agent design for the second trial is summarized in Table 10 as follows. The highlighted cells depict the changes made to the agent design.

Table 10 Trial 2 Agent Design

Agent Number	Agent Name	Organization Type	Role	Group	shareBeforeProcessing	shareModalChoice	propensityToShare	propensityToSeek	shareWith	shareWithWebSites	messageQueueNewerBeforeOlder	awarenessProcessingThreshold	hasSeenEnoughToIdentify	timeBeforeFirstIdentify	minSolutionAreas	idConfidencelevel	minTimeBetweenPulls	primary
1	Alex	Hierarchical	Coordinator		FALSE	peer to peer only	moderate	low	2,3,4,5		TRUE	2	20	40	1	0.49	90000	who,what,when,where
2	Chris	Hierarchical	Team leader	Who	FALSE	post only	moderate	low		Who	TRUE	2	68	1000	1	0.49	3600000	who,what,when,where
3	Dale	Hierarchical	Team leader	What	FALSE	peer to peer dominant	moderate	low	1,2,4,5,9,10,11	What	TRUE	2	20	40	1	0.49	1500000	who,what,when,where
4	Francis	Hierarchical	Team leader	Where	FALSE	post dominant	moderate	low	1,12,13,14	Where	TRUE	2	20	40	1	0.49	900000	who,what,when,where
5	Harlan	Hierarchical	Team member	When	FALSE	both	low	low			TRUE	2	20	1000	4	1.00	3600000	who,what,when,where
6	Jesse	Hierarchical	Team member	Who	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
7	Kim	Hierarchical	Team member	Who	FALSE	post only	moderate	low		Who	TRUE	2	16	40	1	0.00	3600000	who,what,when,where
8	Leslie	Hierarchical	Team member	Who	TRUE	both	moderate	low	2,6,7	Who	TRUE	2	16	40	1	0.25	900000	who,what,when,where
9	Morgan	Hierarchical	Team member	What	FALSE	peer to peer only	moderate	low	3,10,11		TRUE	2	16	40	1	0.25	1500000	who,what,when,where
10	Pat	Hierarchical	Team member	What	FALSE	peer to peer only	moderate	low	3,9,11		TRUE	2	16	1000	1	1.00	1200000	who,what,when,where
11	Quinn	Hierarchical	Team member	What	FALSE	peer to peer dominant	moderate	low	3,9,10	What	TRUE	2	16	40	1	0.49	3600000	who,what,when,where
12	Robin	Hierarchical	Team member	Where	FALSE	post only	moderate	low		Where	TRUE	2	16	40	1	0.49	2400000	who,what,when,where
13	Sam	Hierarchical	Team member	Where	FALSE	both	moderate	low	4,12,14	Where	TRUE	2	16	40	1	0.75	600000	who,what,when,where
14	Sidney	Hierarchical	Team member	Where	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
15	Taylor	Hierarchical	Team member	When	FALSE	post only	moderate	low		When	TRUE	2	6	40	1	0.49	600000	who,what,when,where
16	Val	Hierarchical	Team member	When	FALSE	peer to peer dominant	moderate	moderate	5,15,17	When	TRUE	2	6	40	1	0.25	60000	who,what,when,where
17	Whitley	Hierarchical	Team member	When	FALSE	post only	moderate	low		When	TRUE	2	6	40	1	0.49	2400000	who,what,when,where

This revised agent design yields the following results summarized in Table 11:

Table 11 Trial 2 Results Summary

Agent Trial 2 Results										
Agent Number	Agent	Agent Role	Agent Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		4	0	0	0	72	88	164
2	Chris	Team leader	Who	4	0	36	32	1	0	41
3	Dale	Team leader	What	4	0	28	16	2	172	206
4	Francis	Team leader	Where	4	0	28	24	4	102	138
5	Harlan	Team leader	When	4	0	0	0	1	0	5
6	Jesse	Team member	Who	4	0	0	0	1	0	5
7	Kim	Team member	Who	4	0	8	4	9	0	21
8	Leslie	Team member	Who	4	0	4	4	4	12	24
9	Morgan	Team member	What	4	0	0	0	2	60	66
10	Pat	Team member	What	4	0	0	0	3	60	67
11	Quinn	Team member	What	4	0	28	12	1	60	93
12	Robin	Team member	Where	4	0	32	4	2	0	38
13	Sam	Team member	Where	4	0	28	4	5	60	97
14	Sidney	Team member	Where	4	0	0	0	1	0	5
15	Taylor	Team member	When	4	0	8	4	5	0	17
16	Val	Team member	When	4	0	4	4	47	12	67
17	Whitley	Team member	When	4	2	8	4	2	0	16
Total				68	2	212	112	162.00	626	1070
Average				4	0.12	12.47	6.59	9.53	36.82	62.94

When the agent transaction log from the second design is compared to that of the human trial, we observe the following differences in results as summarized in Table 12.

Table 12 Difference between Trial 2 and Human Trial

Difference Between Agent & Human Trial										
Number	Name	Role	Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		0	1	0	0	-18	-59	-76
2	Chris	Team leader	Who	0	0	-9	-19	0	0	-9
3	Dale	Team leader	What	0	1	-19	-12	-1	-4	-23
4	Francis	Team leader	Where	0	1	-15	-13	2	-90	-102
5	Harlan	Team leader	When	0	0	0	0	-1	0	-1
6	Jesse	Team member	Who	0	0	0	0	-1	0	-1
7	Kim	Team member	Who	0	1	-4	0	-1	0	-4
8	Leslie	Team member	Who	0	1	1	0	0	-3	-1
9	Morgan	Team member	What	0	1	0	0	-1	-48	-48
10	Pat	Team member	What	0	0	0	0	0	-48	-48
11	Quinn	Team member	What	0	1	-25	-9	-1	-24	-49
12	Robin	Team member	Where	0	1	-24	-1	-1	0	-24
13	Sam	Team member	Where	0	1	-27	-4	4	-54	-76
14	Sidney	Team member	Where	0	0	0	0	-1	0	-1
15	Taylor	Team member	When	0	1	-6	-2	7	0	2
16	Val	Team member	When	0	1	-3	-4	0	-1	-3
17	Whitley	Team member	When	0	-1	-4	-2	-1	0	-6
Total				0	10	-135	-66	-14	-331	-470
Average				0	0.59	-7.9	-3.88	-0.8	-19.5	-27.6

Analyzing the results of the revised design, we see a lack of identification attempts and therefore need to adjust parameter settings accordingly. The *timeBeforeFirstIdentify* is decreased again to thirty minutes for all agents. The max *idConfidencelevel* is changed from 0.49 to 0.25 for all agents previously set to 0.25 or greater as predetermined by the survey answer. The new agent design for the third trial is summarized in Table 13 as follows:

Table 13 Trial 3 Agent Design

Agent Number	Agent Name	Organization Type	Role	Group	shareBeforeProcessing	shareModalChoice	propensityToShare	propensityToSeek	shareWith	shareWithWebSites	messageQueueNewerBeforeOlder	awarenessProcessingThreshold	hasSeenEnoughToIdentify	timeBeforeFirstIdentify	minSolutionAreas	idConfidencelevel	minTimeBetweenPulls	primary
1	Alex	Hierarchical	Coordinator		FALSE	peer to peer only	moderate	low	2,3,4,5		TRUE	2	20	30	1	0.25	90000	who,what,when,where
2	Chris	Hierarchical	Team leader	Who	FALSE	post only	moderate	low		Who	TRUE	2	68	1000	1	0.25	3600000	who,what,when,where
3	Dale	Hierarchical	Team leader	What	FALSE	peer to peer dominant	moderate	low	1,2,4,5,9,10,11	What	TRUE	2	20	30	1	0.25	1500000	who,what,when,where
4	Francis	Hierarchical	Team leader	Where	FALSE	post dominant	moderate	low	1,12,13,14	Where	TRUE	2	20	30	1	0.25	900000	who,what,when,where
5	Harlan	Hierarchical	Team member	When	FALSE	both	low	low			TRUE	2	20	1000	4	1.00	3600000	who,what,when,where
6	Jesse	Hierarchical	Team member	Who	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
7	Kim	Hierarchical	Team member	Who	FALSE	post only	moderate	low		Who	TRUE	2	16	30	1	0.00	360000	who,what,when,where
8	Leslie	Hierarchical	Team member	Who	TRUE	both	moderate	low	2,6,7	Who	TRUE	2	16	30	1	0.25	900000	who,what,when,where
9	Morgan	Hierarchical	Team member	What	FALSE	peer to peer only	moderate	low	3,10,11		TRUE	2	16	30	1	0.25	1500000	who,what,when,where
10	Pat	Hierarchical	Team member	What	FALSE	peer to peer only	moderate	low			TRUE	2	16	1000	1	1.00	1200000	who,what,when,where
11	Quinn	Hierarchical	Team member	What	FALSE	peer to peer dominant	moderate	low	3,9,10	What	TRUE	2	16	30	1	0.25	3600000	who,what,when,where
12	Robin	Hierarchical	Team member	Where	FALSE	post only	moderate	low		Where	TRUE	2	16	30	1	0.25	2400000	who,what,when,where
13	Sam	Hierarchical	Team member	Where	FALSE	both	moderate	low	4,12,14	Where	TRUE	2	16	30	1	0.25	600000	who,what,when,where
14	Sidney	Hierarchical	Team member	Where	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
15	Taylor	Hierarchical	Team member	When	FALSE	post only	moderate	low		When	TRUE	2	6	30	1	0.25	600000	who,what,when,where
16	Val	Hierarchical	Team member	When	FALSE	peer to peer dominant	moderate	moderate	5,15,17	When	TRUE	2	6	30	1	0.25	60000	who,what,when,where
17	Whitley	Hierarchical	Team member	When	FALSE	post only	moderate	low		When	TRUE	2	6	30	1	0.25	2400000	who,what,when,where

The third design iteration yields the following results summarized in Table 14:

Table 14 Trial 3 Results Summary

Agent Trial 3 Results										
Agent Number	Agent	Agent Role	Agent Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		4	0	0	0	76	88	168
2	Chris	Team leader	Who	4	0	36	32	1	0	41
3	Dale	Team leader	What	4	0	28	16	2	172	206
4	Francis	Team leader	Where	4	0	28	24	4	103	139
5	Harlan	Team leader	When	4	0	0	0	1	0	5
6	Jesse	Team member	Who	4	0	0	0	1	0	5
7	Kim	Team member	Who	4	2	8	4	9	0	23
8	Leslie	Team member	Who	4	3	4	4	4	12	27
9	Morgan	Team member	What	4	0	0	0	2	60	66
10	Pat	Team member	What	4	0	0	0	3	60	67
11	Quinn	Team member	What	4	0	28	12	1	60	93
12	Robin	Team member	Where	4	0	32	4	2	0	38
13	Sam	Team member	Where	4	0	28	4	5	60	97
14	Sidney	Team member	Where	4	0	0	0	1	0	5
15	Taylor	Team member	When	4	0	8	4	5	0	17
16	Val	Team member	When	4	0	4	4	47	12	67
17	Whitley	Team member	When	4	2	8	4	2	0	16
Total				68	7	212	112	166.00	627	1080
Average				4	0.41	12.47	6.59	9.76	36.88	63.53

The difference between the agent transaction log from the third design and human transaction log is displayed via Table 15 below.

Table 15 Difference between Trial 3 and Human Trial

Difference Between Agent & Human Trial										
Number	Name	Role	Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		0	1	0	0	-22	-59	-80
2	Chris	Team leader	Who	0	0	-9	-19	0	0	-9
3	Dale	Team leader	What	0	1	-19	-12	-1	-4	-23
4	Francis	Team leader	Where	0	1	-15	-13	2	-91	-103
5	Harlan	Team leader	When	0	0	0	0	-1	0	-1
6	Jesse	Team member	Who	0	0	0	0	-1	0	-1
7	Kim	Team member	Who	0	-1	-4	0	-1	0	-6
8	Leslie	Team member	Who	0	-2	1	0	0	-3	-4
9	Morgan	Team member	What	0	1	0	0	-1	-48	-48
10	Pat	Team member	What	0	0	0	0	0	-48	-48
11	Quinn	Team member	What	0	1	-25	-9	-1	-24	-49
12	Robin	Team member	Where	0	1	-24	-1	-1	0	-24
13	Sam	Team member	Where	0	1	-27	-4	4	-54	-76
14	Sidney	Team member	Where	0	0	0	0	-1	0	-1
15	Taylor	Team member	When	0	1	-6	-2	7	0	2
16	Val	Team member	When	0	1	-3	-4	0	-1	-3
17	Whitley	Team member	When	0	-1	-4	-2	-1	0	-6
Total				0	5	-135	-66	-18	-332	-480
Average				0	0.3	-7.9	-3.88	-1.1	-19.5	-28.2

Analyzing the results of the third trial, we look more closely at the differences in the quantity of actions each agent performs and make adjustments. The *timeBeforeFirstIdentify* is decreased again to twenty minutes for all agents. In order to better match the sharing actions performed in the human trial, the *propensityToShare* setting is set to low for several agents. Appropriate adjustments are also made to *minTimeBetweenPulls* for those agents who should Pull less frequently. The *shareModalChoice* setting of agent 13, pseudonym “Sam,” is changed from both to peer to peer dominant as the human counterpart performs more direct Shares than Posts. Finally, we observe too many website Posts in the agent trial as compared to the human trial we are mimicking. In order to decrease the number of Posts, we remove the area website from several agents’ *shareWithWebSites* list. With this change the agents are not forced to Post to websites. The new agent design for the fourth trial is summarized through Table 16 as follows:

Table 16 Trial 4 Agent Design

Agent Number	Agent Name	Organization Type	Role	Group	shareBeforeProcessing	shareModalChoice	propensityToShare	propensityToSeek	shareWith	shareWithWebSites	messageQueueNewerBeforeOlder	awarenessProcessingThreshold	hasSeenEnoughToIdentify	timeBeforeFirstIdentify	minSolutionAreas	idConfidencelevel	minTimeBetweenPulls	primary
1	Alex	Hierarchical	Coordinator		FALSE	peer to peer only	moderate	low	2,3,4,5		TRUE	2	20	20	1	0.25	120000	who,what,when,where
2	Chris	Hierarchical	Team leader	Who	FALSE	post only	moderate	low		Who	TRUE	2	68	1000	1	0.25	3600000	who,what,when,where
3	Dale	Hierarchical	Team leader	What	FALSE	peer to peer dominant	moderate	low	1,2,4,5,9,10,11		TRUE	2	20	20	1	0.25	1500000	who,what,when,where
4	Francis	Hierarchical	Team leader	Where	FALSE	post dominant	low	low	1,12,13,14	Where	TRUE	2	20	20	1	0.25	900000	who,what,when,where
5	Harlan	Hierarchical	Team leader	When	FALSE	both	low	low			TRUE	2	20	1000	4	1.00	3600000	who,what,when,where
6	Jesse	Hierarchical	Team member	Who	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
7	Kim	Hierarchical	Team member	Who	FALSE	post only	low	low		Who	TRUE	2	16	20	1	0.00	360000	who,what,when,where
8	Leslie	Hierarchical	Team member	Who	TRUE	both	moderate	low	2,6,7	Who	TRUE	2	16	20	1	0.25	900000	who,what,when,where
9	Morgan	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,10,11		TRUE	2	16	20	1	0.25	1500000	who,what,when,where
10	Pat	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,9,11		TRUE	2	16	1000	1	1.00	1200000	who,what,when,where
11	Quinn	Hierarchical	Team member	What	FALSE	peer to peer dominant	low	low	3,9,10		TRUE	2	16	20	1	0.25	3600000	who,what,when,where
12	Robin	Hierarchical	Team member	Where	FALSE	post only	moderate	low			TRUE	2	16	20	1	0.25	2400000	who,what,when,where
13	Sam	Hierarchical	Team member	Where	FALSE	peer to peer dominant	low	low	4,12,14	Where	TRUE	2	16	20	1	0.25	600000	who,what,when,where
14	Sidney	Hierarchical	Team member	Where	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
15	Taylor	Hierarchical	Team member	When	FALSE	post only	low	low		When	TRUE	2	6	20	1	0.25	360000	who,what,when,where
16	Val	Hierarchical	Team member	When	FALSE	peer to peer dominant	moderate	moderate	5,15,17	When	TRUE	2	6	20	1	0.25	60000	who,what,when,where
17	Whitley	Hierarchical	Team member	When	FALSE	post only	moderate	low		When	TRUE	2	6	20	1	0.25	2400000	who,what,when,where

The fourth design iteration yields the following results summarized in Table 17:

Table 17 Trial 4 Results Summary

Agent Trial 4 Results										
Agent Number	Agent	Agent Role	Agent Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		4	1	0	0	65	76	146
2	Chris	Team leader	Who	4	0	32	28	1	0	37
3	Dale	Team leader	What	4	0	0	0	2	148	154
4	Francis	Team leader	Where	4	0	4	4	4	16	28
5	Harlan	Team leader	When	4	0	0	0	1	0	5
6	Jesse	Team member	Who	4	0	0	0	1	0	5
7	Kim	Team member	Who	4	1	4	4	9	0	18
8	Leslie	Team member	Who	4	4	4	4	4	12	28
9	Morgan	Team member	What	4	0	0	0	2	12	18
10	Pat	Team member	What	4	0	0	0	3	12	19
11	Quinn	Team member	What	4	0	0	0	1	12	17
12	Robin	Team member	Where	4	0	5	2	2	0	11
13	Sam	Team member	Where	4	0	4	4	5	12	25
14	Sidney	Team member	Where	4	0	0	0	1	0	5
15	Taylor	Team member	When	4	0	4	4	9	0	17
16	Val	Team member	When	4	0	4	4	47	12	67
17	Whitley	Team member	When	4	2	8	4	2	0	16
Total				68	8	69	58	159.00	312	616
Average				4	0.47	4.06	3.41	9.35	18.35	36.24

The difference between the agent transaction log from the fourth design and human transaction log is displayed through Table 18 below.

Table 18 Difference between Trial 4 and Human Trial

Difference Between Agent & Human Trial										
Number	Name	Role	Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		0	0	0	0	-11	-47	-58
2	Chris	Team leader	Who	0	0	-5	-15	0	0	-5
3	Dale	Team leader	What	0	1	9	4	-1	20	29
4	Francis	Team leader	Where	0	1	9	7	2	-4	8
5	Harlan	Team leader	When	0	0	0	0	-1	0	-1
6	Jesse	Team member	Who	0	0	0	0	-1	0	-1
7	Kim	Team member	Who	0	0	0	0	-1	0	-1
8	Leslie	Team member	Who	0	-3	1	0	0	-3	-5
9	Morgan	Team member	What	0	1	0	0	-1	0	0
10	Pat	Team member	What	0	0	0	0	0	0	0
11	Quinn	Team member	What	0	1	3	3	-1	24	27
12	Robin	Team member	Where	0	1	3	1	-1	0	3
13	Sam	Team member	Where	0	1	-3	-4	4	-6	-4
14	Sidney	Team member	Where	0	0	0	0	-1	0	-1
15	Taylor	Team member	When	0	1	-2	-2	3	0	2
16	Val	Team member	When	0	1	-3	-4	0	-1	-3
17	Whitley	Team member	When	0	-1	-4	-2	-1	0	-6
Total				0	4	8	-12	-11	-17	-16
Average				0	0.24	0.47	-0.7	-0.6	-1	-0.94

These changes bring us very close to the behaviors of the human participants. After analyzing the results of trial four, we make a few adjustments. The

timeBeforeFirstIdentify is decreased to fifteen minutes for all agents. The value for *hasSeenEnoughToIdentify* is decreased for pseudonym “Robin” and “Sam” to 10 as they did not see 16 unique facts during the course of the trial 4. Pseudonym “Quinn’s” area website (*what*) is added back to his *shareWithWebSites* list to increase posting. “Dale’s” *shareModalChoice* setting is changed from peer to peer dominant to both. “Alex’s” *minTimeBetweenPulls* is increased to adjust for less frequent pulling. The new agent design for trial five is summarized in Table 19:

Table 19 Trial 5 Agent Design

Agent Number	Agent Name	Organization Type	Role	Group	shareBeforeProcessing	shareModalChoice	propensityToShare	propensityToSeek	shareWith	shareWithWebSites	messageQueueNewerBeforeOlder	awarenessProcessingThreshold	hasSeenEnoughToIdentify	timeBeforeFirstIdentify	minSolutionAreas	idConfidencelevel	minTimeBetweenPulls	primary
1	Alex	Hierarchical	Coordinator		FALSE	peer to peer only	moderate	low	2,3,4,5		TRUE	2	20	15	1	0.25	150000	who,what,when,where
2	Chris	Hierarchical	Team leader	Who	FALSE	post only	moderate	low		Who	TRUE	2	68	1000	1	0.25	3600000	who,what,when,where
3	Dale	Hierarchical	Team leader	What	FALSE	both	moderate	low	1,2,4,5,9,10,11		TRUE	2	20	15	1	0.25	1500000	who,what,when,where
4	Francis	Hierarchical	Team leader	Where	FALSE	post dominant	low	low	1,12,13,14	Where	TRUE	2	20	15	1	0.25	900000	who,what,when,where
5	Harlan	Hierarchical	Team leader	When	FALSE	both	low	low			TRUE	2	20	1000	4	1.00	3600000	who,what,when,where
6	Jesse	Hierarchical	Team member	Who	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
7	Kim	Hierarchical	Team member	Who	FALSE	post only	low	low		Who	TRUE	2	16	15	1	0.00	360000	who,what,when,where
8	Leslie	Hierarchical	Team member	Who	TRUE	both	moderate	low	2,6,7	Who	TRUE	2	16	15	1	0.25	900000	who,what,when,where
9	Morgan	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,10,11		TRUE	2	16	15	1	0.25	1500000	who,what,when,where
10	Pat	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,9,11		TRUE	2	16	1000	1	1.00	1200000	who,what,when,where
11	Quinn	Hierarchical	Team member	What	FALSE	peer to peer dominant	low	low	3,9,10	what	TRUE	2	16	15	1	0.25	3600000	who,what,when,where
12	Robin	Hierarchical	Team member	Where	FALSE	post only	moderate	low			TRUE	2	10	15	1	0.25	2400000	who,what,when,where
13	Sam	Hierarchical	Team member	Where	FALSE	peer to peer dominant	low	low	4,12,14	Where	TRUE	2	10	15	1	0.25	600000	who,what,when,where
14	Sidney	Hierarchical	Team member	Where	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
15	Taylor	Hierarchical	Team member	When	FALSE	post only	low	low		When	TRUE	2	6	15	1	0.25	360000	who,what,when,where
16	Val	Hierarchical	Team member	When	FALSE	peer to peer dominant	moderate	moderate	5,15,17	When	TRUE	2	6	15	1	0.25	60000	who,what,when,where
17	Whitley	Hierarchical	Team member	When	FALSE	post only	moderate	low		When	TRUE	2	6	15	1	0.25	2400000	who,what,when,where

The fifth design iteration yields the results reflected in Table 20:

Table 20 Trial 5 Results Summary

Agent Trial 5 Results										
Agent Number	Agent	Agent Role	Agent Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		4	4	0	0	44	76	128
2	Chris	Team leader	Who	4	0	32	28	1	0	37
3	Dale	Team leader	What	4	3	22	18	2	148	179
4	Francis	Team leader	Where	4	1	4	4	4	16	29
5	Harlan	Team leader	When	4	0	0	0	1	0	5
6	Jesse	Team member	Who	4	0	0	0	1	0	5
7	Kim	Team member	Who	4	3	4	4	9	0	20
8	Leslie	Team member	Who	4	5	4	4	4	12	29
9	Morgan	Team member	What	4	2	0	0	2	12	20
10	Pat	Team member	What	4	0	0	0	3	12	19
11	Quinn	Team member	What	4	2	4	4	1	12	23
12	Robin	Team member	Where	4	0	5	2	2	0	11
13	Sam	Team member	Where	4	0	4	4	5	12	25
14	Sidney	Team member	Where	4	0	0	0	1	0	5
15	Taylor	Team member	When	4	0	4	4	9	0	17
16	Val	Team member	When	4	0	4	4	47	12	67
17	Whitley	Team member	When	4	4	8	4	2	0	18
Total				68	24	95	80	138.00	312	637
Average				4	1.41	5.59	4.71	8.12	18.35	37.47

The difference between the agent transaction log from the fifth design and human transaction log is displayed in Table 21.

Table 21 Difference between Trial 5 and Human Trial

Difference Between Agent & Human Trial										
Number	Name	Role	Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		0	-3	0	0	10	-47	-40
2	Chris	Team leader	Who	0	0	-5	-15	0	0	-5
3	Dale	Team leader	What	0	-2	-13	-14	-1	20	4
4	Francis	Team leader	Where	0	0	9	7	2	-4	7
5	Harlan	Team leader	When	0	0	0	0	-1	0	-1
6	Jesse	Team member	Who	0	0	0	0	-1	0	-1
7	Kim	Team member	Who	0	-2	0	0	-1	0	-3
8	Leslie	Team member	Who	0	-4	1	0	0	-3	-6
9	Morgan	Team member	What	0	-1	0	0	-1	0	-2
10	Pat	Team member	What	0	0	0	0	0	0	0
11	Quinn	Team member	What	0	-1	-1	-1	-1	24	21
12	Robin	Team member	Where	0	1	3	1	-1	0	3
13	Sam	Team member	Where	0	1	-3	-4	4	-6	-4
14	Sidney	Team member	Where	0	0	0	0	-1	0	-1
15	Taylor	Team member	When	0	1	-2	-2	3	0	2
16	Val	Team member	When	0	1	-3	-4	0	-1	-3
17	Whitley	Team member	When	0	-3	-4	-2	-1	0	-8
Total				0	-12	-18	-34	10	-17	-37
Average				0	-0.7	-1.1	-2	0.6	-1	-2.18

The differences displayed in the table above are quite small, however trial 4 was a better match to the human trial. Only results for one agent, “Dale,” differ more from the

human results in trial five than in trial four. For this reason, Dale's *shareModalChoice* setting is changed back to peer to peer dominant. The new agent design for trial six is shown in Table 22:

Table 22 Trial 6 Agent Design

Agent Number	Agent Name	Organization Type	Role	Group	shareBeforeProcessing	shareModalChoice	propensityToShare	propensityToSeek	shareWith	shareWithWebsites	messageQueueNewerBeforeOlder	awarenessProcessingThreshold	hasSeenEnoughToIdentify	timeBeforeFirstIdentify	minSolutionAreas	idConfidencelevel	minTimeBetweenPulls	primary
1	Alex	Hierarchical	Coordinator		FALSE	peer to peer only	moderate	low	2,3,4,5		TRUE	2	20	15	1	0.25	150000	who,what,when,where
2	Chris	Hierarchical	Team leader	Who	FALSE	post only	moderate	low		Who	TRUE	2	68	1000	1	0.25	3600000	who,what,when,where
3	Dale	Hierarchical	Team leader	What	FALSE	peer to peer dominant	moderate	low	1,2,4,5,9,10,11		TRUE	2	20	15	1	0.25	1500000	who,what,when,where
4	Francis	Hierarchical	Team leader	Where	FALSE	post dominant	low	low	1,12,13,14	Where	TRUE	2	20	15	1	0.25	900000	who,what,when,where
5	Harlan	Hierarchical	Team leader	When	FALSE	both	low	low			TRUE	2	20	1000	4	1.00	3600000	who,what,when,where
6	Jesse	Hierarchical	Team member	Who	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
7	Kim	Hierarchical	Team member	Who	FALSE	post only	low	low		Who	TRUE	2	16	15	1	0.00	360000	who,what,when,where
8	Leslie	Hierarchical	Team member	Who	TRUE	both	moderate	low	2,6,7	Who	TRUE	2	16	15	1	0.25	900000	who,what,when,where
9	Morgan	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,10,11		TRUE	2	16	15	1	0.25	1500000	who,what,when,where
10	Pat	Hierarchical	Team member	What	FALSE	peer to peer only	low	low	3,9,11		TRUE	2	16	1000	1	1.00	1200000	who,what,when,where
11	Quinn	Hierarchical	Team member	What	FALSE	peer to peer dominant	low	low	3,9,10	what	TRUE	2	16	15	1	0.25	3600000	who,what,when,where
12	Robin	Hierarchical	Team member	Where	FALSE	post only	moderate	low			TRUE	2	10	15	1	0.25	2400000	who,what,when,where
13	Sam	Hierarchical	Team member	Where	FALSE	peer to peer dominant	low	low	4,12,14	Where	TRUE	2	10	15	1	0.25	600000	who,what,when,where
14	Sidney	Hierarchical	Team member	Where	FALSE	both	low	low			TRUE	2	6	1000	4	1.00	3600000	who,what,when,where
15	Taylor	Hierarchical	Team member	When	FALSE	post only	low	low		When	TRUE	2	6	15	1	0.25	360000	who,what,when,where
16	Val	Hierarchical	Team member	When	FALSE	peer to peer dominant	moderate	moderate	5,15,17	When	TRUE	2	6	15	1	0.25	60000	who,what,when,where
17	Whitley	Hierarchical	Team member	When	FALSE	post only	moderate	low		When	TRUE	2	6	15	1	0.25	2400000	who,what,when,where

The sixth design iteration yields the following results summarized in Table 23:

Table 23 Trial 6 Results Summary

Agent Trial 6 Results										
Agent Number	Agent Name	Agent Role	Agent Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		4	4	0	0	44	76	128
2	Chris	Team leader	Who	4	0	32	28	1	0	37
3	Dale	Team leader	What	4	2	0	0	2	148	156
4	Francis	Team leader	Where	4	1	4	4	4	16	29
5	Harlan	Team leader	When	4	0	0	0	1	0	5
6	Jesse	Team member	Who	4	0	0	0	1	0	5
7	Kim	Team member	Who	4	3	4	4	9	0	20
8	Leslie	Team member	Who	4	4	4	4	4	12	28
9	Morgan	Team member	What	4	1	0	0	2	12	19
10	Pat	Team member	What	4	0	0	0	3	12	19
11	Quinn	Team member	What	4	1	4	4	1	12	22
12	Robin	Team member	Where	4	0	5	2	2	0	11
13	Sam	Team member	Where	4	0	4	4	5	12	25
14	Sidney	Team member	Where	4	0	0	0	1	0	5
15	Taylor	Team member	When	4	0	4	4	9	0	17
16	Val	Team member	When	4	0	4	4	47	12	67
17	Whitley	Team member	When	4	2	8	4	2	0	16
Total				68	18	73	62	138	312	609
Average				4	1.06	4.29	3.65	8.12	18.35	35.82

The difference between the agent transaction log from the sixth design and human transaction log is displayed in Table 24 below.

Table 24 Difference between Trial 6 and Human Trial

Difference Between Agent & Human Trial										
Number	Name	Role	Task	dist	identify	post	First Post	pull	share	Total
1	Alex	Coordinator		0	-3	0	0	10	-47	-40
2	Chris	Team leader	Who	0	0	-5	-15	0	0	-5
3	Dale	Team leader	What	0	-1	9	4	-1	20	27
4	Francis	Team leader	Where	0	0	9	7	2	-4	7
5	Harlan	Team leader	When	0	0	0	0	-1	0	-1
6	Jesse	Team member	Who	0	0	0	0	-1	0	-1
7	Kim	Team member	Who	0	-2	0	0	-1	0	-3
8	Leslie	Team member	Who	0	-3	1	0	0	-3	-5
9	Morgan	Team member	What	0	0	0	0	-1	0	-1
10	Pat	Team member	What	0	0	0	0	0	0	0
11	Quinn	Team member	What	0	0	-1	-1	-1	24	22
12	Robin	Team member	Where	0	1	3	1	-1	0	3
13	Sam	Team member	Where	0	1	-3	-4	4	-6	-4
14	Sidney	Team member	Where	0	0	0	0	-1	0	-1
15	Taylor	Team member	When	0	1	-2	-2	3	0	2
16	Val	Team member	When	0	1	-3	-4	0	-1	-3
17	Whitley	Team member	When	0	-1	-4	-2	-1	0	-6
Total				0	-6	4	-16	10	-17	-9
Average				0	-0.35	0.24	-0.94	0.59	-1	-0.53

The differences displayed in the table above are sufficiently small for us to stop iterating. Trial six yields better results than any of the other iteration. Further, it appears that any additional changes will only counter adjustments made to reach these results: we appear to be as close as we can get with the parameters available to us. Therefore, this design, achieved with six iterations, appears to mimic the actions performed by participants in the human trial quite well.