

Influencing Agent Group Behavior by Adjusting Cultural Trait Values

Gaurav Tuli and Henry Hexmoor, *Senior Member, IEEE*

Abstract—Social reasoning and norms among individuals that share cultural traits are largely fashioned by those traits. We have explored predominant sociological and cultural traits. We offer a methodology for parametrically adjusting relevant traits. This exploratory study heralds a capability to deliberately tune cultural group traits in order to produce a desired group behavior. To validate our methodology, we implemented a prototypical-agent-based simulated test bed for demonstrating an exemplar from intelligence, surveillance, and reconnaissance scenario. A group of simulated agents traverses a hostile territory while a user adjusts their cultural group trait settings. Group and individual utilities are dynamically observed against parametric values for the selected traits. Uncertainty avoidance index and individualism are the cultural traits we examined in depth. Upon the user's training of the correspondence between cultural values and system utilities, users deliberately produce the desired system utilities by issuing changes to trait. Specific cultural traits are without meaning outside of their context. Efficacy and timely application of traits in a given context do yield desirable results. This paper heralds a path for the control of large systems via parametric cultural adjustments.

Index Terms—Command and control, intelligence, surveillance and reconnaissance (ISR), man on the loop, multiagent system, supervisory control, unmanned vehicles.

NOMENCLATURE

C2	Command and control.
CLV	Collectivism.
CRU	Capability resource unit.
DOD	Department of Defense.
ISR	Intelligence, surveillance, and reconnaissance.
IAS	ISR agent society.
IDV	Individualism.
LTO	Long-term orientation.
MOTL	Man on the loop.
NCW	Network centric warfare.
PDI	Power Distance Index.
PL	Preference list.
UAI	Uncertainty Avoidance Index.

Manuscript received January 3, 2009; revised July 14, 2009 and October 13, 2009; accepted November 3, 2009. First published January 8, 2010; current version published September 15, 2010. This work was supported by the Air Force Research Laboratory under Contract FA8750-06-C-0138. This paper was recommended by Associate Editor J. Liu.

G. Tuli was with the Department of Computer Science, Southern Illinois University, Carbondale, IL 62901 USA. He is now with the Computer Science Department, Virginia Polytechnic Institute and State University, Blacksburg, VA 24061 USA (e-mail: gtuli@cs.siu.edu).

H. Hexmoor is with the Department of Computer Science, Southern Illinois University, Carbondale, IL 62901 USA (e-mail: hexmoor@cs.siu.edu).

Color versions of one or more of the figures in this paper are available online at <http://ieeexplore.ieee.org>.

Digital Object Identifier 10.1109/TSMCB.2009.2036592

UAV	Unmanned aerial vehicle.
ULD	UAI level deciding.
USAF-SAB	United States Air Force Scientific Advisory Board.

I. INTRODUCTION

CORPORATE or organizational culture is a notion that is often considered to be nebulous and soft. Although it is commonly accepted to have a pivotal effect in performance, contemporary formulations lack rigor to enable construction of instrumental tools. Albeit in nascence, we start from the basic cultural dimensions and then take steps toward rigors of specifying cultural parameters for groups. Cultural parameters are fundamental engineering instruments for monitoring and directing culture in large groups. We will not be concerned with directly managing values and norms. By global changes in reasoning of how individuals relate to one another, collective behaviors are indirectly produced. We will examine the dynamics of organizational culture in the military context of C2.

The NCW paradigm initiated by the DOD prescribes requirements for the information sharing and human-automated system collaboration [4]. Operators are expected to leverage multiple information sources for decisions under significant time pressure and when there is uncertainty. However, an increase in the available information resources will place higher cognitive demands on the operators, which could become constraints that limit the success of network centric processes. Therefore, NCW advocates the socialization of (military) action through the socialization of each intent, capability, and awareness [5].

The DOD roadmap for unmanned aircraft system development has also predicted that future UAVs will work as independent robots, which can self-actualize to perform a given task instead of just working as teleoperated robots [6]. To lift the current unmanned vehicles (UVs) to this level, we selected the most widely explored cultural parameters given by social scientists, notably by Hofstede [10]. These parameters will socialize both military forces and warfare machinery, which provides leverage for performing high-level decision making to a human supervisory body as presented in the MOTL paradigm [2], [4], [8].

We are exploring the applicability of the cultural parameters in complex control systems, like ISR system, to make high-level human supervision possible. In order to model the dynamic nature of ISR-C2, the reasoning capabilities of the agents or ISR agents have been used in this paper as the key attribute, which distinguishes among agents. An agent is any software entity with predefined capabilities that can achieve a goal.

After following the recommendations given by several federal agencies, such as the DOD and the USAF-SAB, we have tailored our approach to promote the recommended changes in the present automation level of human supervisory control. The utilitarian approach has been utilized herein to create a few norms and rules for implementing the cultural parameters in a nonhuman cultural setting [7]. Moreover, special care has been given to design *Pareto optimal* ISR agents, which may sacrifice their personal benefits in favor of social benefits [18].

The paper is organized as follows. The next section outlines the details of MOTL and selected cultural parameters. Section III delineates our approach of socializing the agent community by forming and imposing laws using utility theory and role exchange, respectively, and finally presents our algorithm. Section IV presents the implementation of the algorithm in the form of simulated agent-based experiments, revealing several facets of our approach. In Section V, we present the results and conclusions of this paper. Lastly, Section VI includes the future research possibilities based on this work.

II. MOTL AND ITS CULTURAL PARAMETERS

In the shadow of demanding technological advancements in the field of defense like the NCW, many federal scientific agencies, including the USAF-SAB and DOD, have clearly urged the need for improvements in supervisory control [6], [23]. The level of human intervention for supervision can be broadly divided into two categories. If a human interference directly affects actions of a specific acting entity (i.e., an agent) in the control loop, it was traditionally termed as Man In The Loop (or MITL). On the other hand, if the human intervention indirectly affects actions at the unit or agent community level, it was coined as MOTL [8].

MOTL is a socio-psycho-cultural model to observe and mediate behaviors and interactions of complex control systems. It is characterized as a paradigm shift where a human supervisor employs psychological and social influences to control the system behavior. MOTL alleviates excessive dependence on the human supervisor. MOTL by its very nature could be considered as the Sheridan's strict form of supervisory control [20]. A system-level network policy management is an example for MOTL, where changing the policies in the policy management system by the system administrator will affect the configuration of the entire network without tediously altering each system parameter [22].

Hence, MOTL presents a foundation to a novel high-level human supervision that contrasts with typical micromanagement of tasks in complex control systems where large collections of communities are usually founded. Such complex control systems are customary in the NCW due to the notable changes in several war fighting trends. The shift in warfare paradigm from platform centric to network centric has significantly changed the entire human supervision role, both in mission planning and actual operations, as shown by Lambert and Scholz using a real-time example [14, pp. 6]. Instead of controlling the system manually, now, military operators are more involved in the higher levels of planning and decision making.

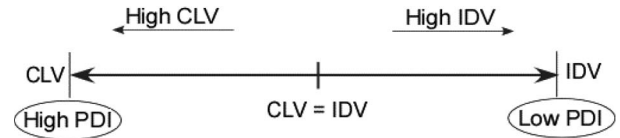


Fig. 1. IDV (or CLV) spectrum and IDV-PDI relationship.

The key attributes of MOTL that characterizes its novelty are based on well-founded theories in social sciences and fall into three main categories, namely, personality, social reasoning, and culture [8]. To avoid cultural differences, which may tend to decrease coordination in an agent society, MOTL uses vital cultural parameters given by Hofstede, namely, UAI, IDV, PDI, LTO, and masculinity [10].

Due to the nature of ISR, particularly during war with people and machine of differing cultures, uncertainty cannot be overlooked, and control over this parameter will eventually improve the whole system. Therefore, we included UAI in our research. It has also been shown that there is a strong interconnection between UAI, IDV, and PDI [1], [9]. Hence, the benefit of including IDV and PDI into control and command has also been explored in this research. However, LTO is worth exploring but is outside of our scope. Similarly, since we are mainly modeling artificial agents and they lack gender characteristics, *masculinity* was not considered in MOTL. Next, we briefly discuss the selected cultural parameters in more details.

A. UAI

Situational awareness is often reduced in a warlike scenario due to “*Fog of War*” that leads to uncertainty [17]. Uncertainty is also an integral characteristic of organizational decision making and any cultural makeup [3]. Hofstede explained uncertainty as the dimension that does not relate to risky situations but rather to unknown or unfamiliar situations [10]. Different cultures have different levels of tolerance for uncertainty and ambiguity that is measured and represented using UAI [10]. UAI indicates the extent to which a culture programs its members to feel either uncomfortable (i.e., high UAI) or comfortable (i.e., low UAI) in an unstructured situation. The level of imposed restrictions defines the respective level of comfort for the member of a particular culture. UAI is usually enforced by strict laws and rules, safety, and security measures. We have defined certain stringent rules in Section III for our agents to restrict their behavior while working with culturally similar or diverse members.

B. IDV

IDV and CLV describe the relationship between the individual and the collectivity that prevails in a given society, but both reside at the opposite ends, as shown in Fig. 1. The individualistic end of continuum construes a person as a separate entity, which is clearly distinguishable from social milieu [9]. On the other hand, in a collectivist culture, the distinction between the individual and the group becomes blurred, and the people regard themselves as the extension of a social system [13]. In the context of self-theory, the people who are a part of

more individualist cultures have selves or self-cognitions. The individualistic people refer themselves as independent, self-contained, autonomous, and distinct units, whereas people in collectivist cultures give more importance to interdependence or sociocentric identity [15].

C. PDI

PDI measures the extent to which the less powerful members of a society accept and expect that power to be distributed unequally [10]. This index has been incorporated into MOTL to accommodate inherited or developed inter- and intra-power differences among several working agent communities in a complex system. Hence, PDI will capture the degree of inequality among agents in an agent society and measure how much a society has respect for authority. High PDI ranking indicates that the inequalities of power and social values like autonomy, trust, reliance, and benevolence will be allowed to grow within a group [10]. In such a system, agents may view the supervisor as a benevolent dictator and obey her orders. Conversely, low PDI ranking indicates that the group de-emphasizes the differences between the agent's power and social values [10]. Hence, equality and opportunity for every agent are stressed in this case.

There exists an inherent relationship between IDV and PDI. Bochner and Hesketh found that only collectivist culture can work under pressure [1]. They studied individuals in a culturally diverse work setting and found that those who belong to high collectivist cultural background allow or feel comfortable while working under their boss' supervision. These individuals prefer a boss or a manager to be more autocratic and paternalistic, whereas individuals who belong to high individualistic culture prefer their boss or manager to be more hands-on and consultative [1]. Hence, only high CLV (or low IDV) cultures will allow high PDI. This spectrum of IDV–PDI relationship is also shown in Fig. 1.

III. SOCIALIZING AGENTS

Supervisory control paradigm is in need of a paradigm shift to invert the current many-to-one operator to UV ratio. Herein, we present our strides toward this new paradigm by socializing the agents. We replicated individual as agents, group of agents as culture, and group of cultures as society. Let us take an example of ISR operations. An ISR agent, denoted as agent in the literature, is a software component, which takes part in an ISR mission. For example, a software component responsible for the combat behavior, a component responsible for goals setting, or a software component responsible for communication with a ground station can be considered as an agent. Different cultural traits or properties of agents lead to different agent cultures, and a collection of agents with different or similar cultural traits leads to agent society, which we denoted as IAS.

In a multiagent environment, as in IAS, every agent would perform its role in collaborative actions to achieve a *well-defined goal*. For example, swiftly performing ISR in an adversary's territory with the lowest possible damage to the IAS is a *well-defined goal*. A C2 personnel or a human supervisor can transform or build such goals and can also supervise the

actions of the IAS agents using MOTL. The human supervisor or MOTL operator is addressed as *Moe* in this paper. Training *Moe* to identify the correspondence between the agents' cultural values and MOTL-based supervisory tools would improve her efficacy to administer IAS. However, discussion on such training is out of the scope of this paper.

The basic aim of this paper is to make agents of different cultures work collectively in an IAS. Hence, a sense of society is given to the agents that will allow the agents to take socially responsible decisions while working under the same umbrella of dirty, dull, and dangerous tasks. This will lead to the formation of a big IAS holding several small or big agent cultures. These IAS could either be *heterogeneous*, if one or more agent cultures belong to different cultural settings, or *homogeneous*, if the cultural settings of all the participating agent cultures are similar [10]. Moreover, similar to most social processes, many agents of an IAS have to bear personal sacrifices to increase social benefits [16]. Our approach mimics this social phenomenon using IAS that obeys the *Pareto optimal* principle [18].

A. Rule- and Norm-Based UAIs

In order to limit the inherent uncertainties in ISR, we imposed strict laws and rules over IAS. These rules provide us with capabilities for fine tuning or controlling uncertainty factors. Our paradigm aims the following: 1) reduce uncertainty for collaborative work and 2) make a control task to be less tiresome. To achieve these goals, we have designed our algorithm using concepts from *utility theory* and *joint benefits*.

To deal with the multidimensionality of alternatives in a complex decision-making system, we used *additive* and *lexicographic* properties while dealing with our agents' utility function [7]. The *additive* property of the utility function helps to represent an agent as a collection of various attributes, such as its functional capabilities, whereas the *lexicographic* property would allow an agent to identify more important attributes of the utility function among other attributes. The significance of using these features is explained with an example later in this section.

Similarly, to allow equal distribution of benefits to the participating agent cultures for collaborative achievements, we used the *joint-benefit* decision-making approaches given by Kalenka and Jennings [12]. The *joint-benefit* approach is the combination of *individual-* and *social-benefit* decision-making approaches. In the *individual-benefit-based* approach, an agent performing the lead role will always get the biggest piece of the benefit, whereas the *social-benefit* approach offers a proportional share of the benefits to all team members [12]. Hence, the *joint-benefit* approach prefers to maximize both the *individual* and *social benefits*. In our approach, we considered the *social benefit* as the summation of *social gain* and *social loss*, whereas the *individual benefits* as the summation of *individual gain* and *individual loss*, as given by

$$\begin{aligned} \text{Joint Benefit} = & \text{Individual gain} - \text{Individual loss} \\ & + \text{Social gain} - \text{Social loss}. \end{aligned} \quad (1)$$

1. Define each role as an array of predefined set of capabilities and randomly assign it to each agent.
2. Initialize default UAI Levels of each agent based on *ULD* capabilities.
3. Identify the role-exchange-seeker and role-exchange-helper agent class using user's desired UAI level and each agent's previous and base UAI levels.
4. Determine pairs of agents with maximum *Joint Benefit* (utility gain) as role exchange partners and perform role exchange between them.

Fig. 2. Algorithm (given at a high level).

B. Implementing Rules and Norms on IAS

In order to improve the *joint benefit*, we implemented the UAI in IAS using a role-exchange technique called the *utility-based role exchange* [19]. Hence, to achieve a *well-defined goal*, the agents will exchange their roles based on the *joint-benefit* value calculated according to (1). The reasoning for role swapping is discussed next.

The first reasoning that an agent requires is to answer the following question: *Whom do I ask for a role exchange?* The *rational* negotiation strategies defined by Rahman and Hexmoor are utilized to answer this question [19]. Using the *rational* approaches of negotiation, an agent will compute the *joint benefit* with respect to every agent that belongs to its priority list. After computing the utility value for the agents in the priority list, the next dilemma would be to decide: *Which utility value should I select for the role exchange?* We used the *individual optimal role-exchange* scheme to answer this question, which searches each agent's priority list and selects a maximum *joint-benefit-producing* agent as the role-exchange partner [24], [25]. Moreover, to make our approach *Pareto optimal*, we used such a utility function which will provide us with a socially committed priority list with every agent willing to exchange roles irrespective of its personal losses.

C. Role-Exchange Algorithm Using Joint Benefits

Let us now discuss the procedure that utilizes the concepts discussed so far and allows agents to encompass the diversity, dynamicity, and versatility of ISR. A four-step algorithm of this procedure is shown in Fig. 2. Details of the steps are described in text.

1) *Agents, Capabilities, and Roles*: An IAS is a collection of agents of similar or different cultural traits, where an agent can be represented as $\mathbf{a}_i$ and the collection of such agents by set **AGENT**. The proposed algorithm uses the agents' capabilities to distinguish among agents and among roles performed by these agents. The set **CAP** represents the collection of different capabilities $\mathbf{c}_j$.

We assume that each of these capabilities is generated by a CRU. Hence, a CRU can be considered as a collection of equipment and machineries, responsible to provide a capability to an agent. For example, an *attack* CRU of a combat vehicle would be a collection of advanced embedded systems and

several lethal weapons, which provides the *attack capability* to this combat vehicle. Every such physical unit that provides a capability to an agent is considered as a CRU in this paper. Hence, we can say that each capability $\mathbf{c}_j$ of an agent $\mathbf{a}_i$ is generated by a CRU, denoted as $CRU_{\mathbf{c}_j}^{\mathbf{a}_i}$. Moreover, these capabilities are generated in a wide spectrum of values, denoted by set $V_{\mathbf{c}_j}^{\mathbf{a}_i}$. At any time t , an agent $\mathbf{a}_i$ possesses a combination of capabilities, i.e., $\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_n$, and their corresponding values $\mathbf{v}_{\mathbf{c}_1, t}^{\mathbf{a}_i}, \mathbf{v}_{\mathbf{c}_2, t}^{\mathbf{a}_i}, \dots, \mathbf{v}_{\mathbf{c}_n, t}^{\mathbf{a}_i}$, where $1 \leq n \leq |\mathbf{CAP}|$ and $\mathbf{v}_{\mathbf{c}_j, t}^{\mathbf{a}_i} \in V_{\mathbf{c}_j}^{\mathbf{a}_i}$. $|\mathbf{CAP}|$ denotes the total number of capabilities in an IAS, and value $\mathbf{v}_{\mathbf{c}_j, t}^{\mathbf{a}_i}$ denotes the value of a capability $\mathbf{c}_j$ that an agent $\mathbf{a}_i$ is using at time t , where $0 \leq \mathbf{v}_{\mathbf{c}_j, t}^{\mathbf{a}_i} \leq \text{MaxCap}_{\mathbf{c}_j}^{\mathbf{a}_i}$. $\text{MaxCap}_{\mathbf{c}_j}^{\mathbf{a}_i}$ represents the maximum value of capability $\mathbf{c}_j$, which a $CRU_{\mathbf{c}_j}^{\mathbf{a}_i}$ can generate for an agent $\mathbf{a}_i$.

During a mission endeavor, each agent is required to perform some role(s). In an IAS, a role can be represented as $\mathbf{r}_k$ and the collection of such roles by set **ROLE**. To perform a role *fully*, an agent requires a set of capabilities with certain value. For example, at any time t , role $\mathbf{r}_k$ can be *fully* performed by agent $\mathbf{a}_i$, only if $\mathbf{a}_i$ possesses the capabilities $\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_n$ and uses them at the required value levels $\mathbf{v}_{\mathbf{c}_1, t}^{\mathbf{r}_k}, \mathbf{v}_{\mathbf{c}_2, t}^{\mathbf{r}_k}, \dots, \mathbf{v}_{\mathbf{c}_n, t}^{\mathbf{r}_k}$, respectively. Each of the $\mathbf{v}_{\mathbf{c}_j, t}^{\mathbf{r}_k}$ denotes the required value for a capability $\mathbf{c}_j$ to perform a role $\mathbf{r}_k$. This implies that, if $\mathbf{v}_{\mathbf{c}_j, t}^{\mathbf{r}_k} \leq \mathbf{v}_{\mathbf{c}_j, t}^{\mathbf{a}_i}$, then $\mathbf{a}_i$ could perform the role $\mathbf{r}_k$ *fully* using value $\mathbf{v}_{\mathbf{c}_j, t}^{\mathbf{a}_i}$ at $\mathbf{v}_{\mathbf{c}_j, t}^{\mathbf{r}_k}$; otherwise, it performs the role *partially* using a value below $\mathbf{v}_{\mathbf{c}_j, t}^{\mathbf{r}_k}$. Also, one point to be noted that, with the changing requirements of an ISR mission, both the agent $\mathbf{a}_i$'s role assignments and the required value of a capability $\mathbf{c}_j$ to perform a role $\mathbf{r}_k$ or $\mathbf{v}_{\mathbf{c}_j, t}^{\mathbf{r}_k}$ can change.

In an IAS, the number of roles assigned to an agent depends upon the mapping between the total number of agents or $|\mathbf{AGENT}|$ and the total number of roles to be performed or $|\mathbf{ROLE}|$. If $|\mathbf{AGENT}| = |\mathbf{ROLE}|$, then there will be a one-to-one mapping, and one agent will perform one role at any given time t . Similarly, if $|\mathbf{AGENT}| < |\mathbf{ROLE}|$, then more than one role would be assigned to an agent, and if $|\mathbf{AGENT}| > |\mathbf{ROLE}|$, then more than one agent will pursue a role. However, in order to avoid any complex situation that may reduce the clarity of this research, we restricted our experiments discussed in Section IV just to use $|\mathbf{AGENT}| = |\mathbf{ROLE}|$.

2) *ULD Capabilities*: There is a relationship between military operations such as attack and perceptions of unfolding

events. The skills or capabilities used in such military operations entail ambiguity and unfamiliarity into the system. We call this special class of capabilities as the ULD capabilities, denoted by set $\mathbf{UCAP}$, where $\mathbf{UCAP} \subseteq \mathbf{CAP}$. This implies that set $\mathbf{CAP}$ is a collection of non-ULD and ULD capabilities. Moreover, as these uncertainty-inducing capabilities are the integral parts of an ISR operation, therefore, $\mathbf{UCAP} \neq \emptyset$ in an IAS.

The level of uncertainty induced in an IAS depends upon the values at which the ULD capabilities are being used by its agents to perform their roles. Suppose that $UAI_t^{a_i}$ is the tolerance value for uncertainty and ambiguity (or UAI) that an agent a_i satisfies at time t . Therefore, $UAI_t^{a_i}$ depends upon the value of each of the ULD capability, denoted as c_l , used by a_i to perform a role r_k or $v_{c_l,t}^{a_i}$. Moreover, the bigger this value, the lower the $UAI_t^{a_i}$. Hence, $UAI_t^{a_i}$ can be calculated as the inverse function of the summation of all $v_{c_l,t}^{a_i}$, as shown in

$$UAI_t^{a_i} = f\left(\frac{1}{\sum_{l=1}^{|\mathbf{UCAP}|} v_{c_l,t}^{a_i}}\right). \quad (2)$$

3) *Role Exchange in IAS*: An IAS comprises of agents from different cultures with different UAIs. As per the UAI definition explained in Section II, an agent culture with the lowest UAI value dictates the UAI of an IAS or UAI_t^{IAS} . However, (2) implies that changing $v_{c_l,t}^{a_i}$ would change the $UAI_t^{a_i}$, and if $UAI_t^{a_i}$ is altered for every agent a_i , then eventually, we can change the UAI_t^{IAS} . *Moe* can enforce a UAI value, denoted by UAI_t^{Moe} , to collectively change an IAS behavior. If *Moe* specifies UAI_t^{Moe} for an IAS, then every agent a_i has to comply with this decision by keeping $UAI_t^{a_i}$ to this desired UAI level. However, if there is no UAI imposed by *Moe*, i.e., $UAI_t^{\text{Moe}} = 0$, or if agents are operating at $UAI_t^{\text{IAS}} = 0$, then each agent a_i is allowed to operate at its *base* UAI, denoted by $BaseUAI^{a_i}$. At this UAI level, if required, an agent a_i can use each c_j at $MaxCap_{c_j}^{a_i}$ and therefore can induce maximum uncertainty dictated by its culture into the system. Hence, the *base* UAI of an agent dictates the lowest UAI level up to which this agent can operate.

The UAI tuning can be achieved using a *utility-based role-exchange* technique. A role exchange at any time t between two IAS agents will swap their current roles at time $t + 1$. For example, if at time t an agent a_1 is performing role r_1 and a_2 is performing role r_2 , and if they initiate a role exchange, then at time $t + 1$, agent a_1 will perform role r_2 , whereas a_2 will perform role r_1 . This role exchange will also change the $UAI_t^{a_1}$ and $UAI_t^{a_2}$ according to (2).

Supposing that *Moe* specifies UAI_t^{Moe} to an IAS at time t , which is different than UAI_t^{IAS} , then role exchanges are required. Each role exchange requires two agents: *role-exchange-seeker* agent—agent that seeks a role to exchange, and *role-exchange-helper* agent—agent that helps the seeker agent by volunteering his role. If a role exchange between agents a_s and a_h changes, either increase or decrease, the $UAI_t^{a_s}$ such that $UAI_t^{a_s} = UAI_t^{\text{Moe}}$, then agents a_s and a_h are *role-exchange-seeker* and *role-exchange-helper* agents, respectively. When a role exchange is necessary to maintain UAI_t^{Moe} , then each agent a_s forms a PL, denoted by $PL_t^{a_s}$,

that represents a collection of agents a_h . Let us consider the following two scenarios where role exchange is required.

- a) If at time t $UAI_t^{\text{Moe}} > UAI_t^{\text{IAS}}$, then *Moe* is trying to increase the UAI of an IAS. In this case, an agent a_s that is operating below UAI_t^{Moe} , i.e., $UAI_t^{a_s} < UAI_t^{\text{Moe}}$, becomes a *role-exchange-seeker* agent to increase its UAI level, and an agent a_h with $UAI_t^{a_h} \geq UAI_t^{\text{Moe}}$ becomes a *role-exchange-helper* agent a_h . Hence, $PL_t^{a_s}$ is given by

$$PL_t^{a_s} = \{a_h | UAI_t^{a_h} \geq UAI_t^{\text{Moe}}\}. \quad (3a)$$

- b) If at time t $UAI_t^{\text{Moe}} < UAI_t^{\text{IAS}}$, then *Moe* is trying to decrease the UAI of an IAS. In this case, an agent a_s with *base* UAI that is less than UAI_t^{Moe} , i.e., $BaseUAI^{a_s} < UAI_t^{\text{Moe}}$, becomes a *role-exchange-seeker* agent to decrease its UAI level and forms $PL_t^{a_s}$, given by (3b), with the following *role-exchange-helper* agents.

- i) An agent a_h with $BaseUAI^{a_h} \geq UAI_t^{\text{IAS}}$.
- ii) An agent a_h with *base* UAI that is greater than or equal to UAI_t^{Moe} but less than UAI_t^{IAS} , i.e., $UAI_t^{\text{Moe}} \leq BaseUAI^{a_h} < UAI_t^{\text{IAS}}$, becomes a *role-exchange-helper* agent after resetting its UAI to *base* UAI, i.e., $UAI_t^{a_h} = BaseUAI^{a_h}$.

$$PL_t^{a_s} = \{a_h | BaseUAI^{a_h} \geq UAI_t^{\text{IAS}} \text{ or } UAI_t^{\text{Moe}} \leq BaseUAI^{a_h} < UAI_t^{\text{IAS}}\}. \quad (3b)$$

4) *Joint Benefit in IAS*: Each *role-exchange-seeker* agent a_s has the list of *role-exchange-helper* agents a_h in the form of $PL_t^{a_s}$ calculated either by (3a) or (3b). Now, a_s has to choose its *role-exchange partner* from this list. To reach this decision, agent a_s will calculate the *Joint Benefit* given by (1) for every agent a_h in the $PL_t^{a_s}$. After this calculation, an agent a_h with the maximum *Joint Benefit* will be selected as the *role-exchange partner* of an agent a_s . Calculating *Joint Benefit* for every agent in the $PL_t^{a_s}$ exemplifies the *rational* negotiation strategy, whereas selecting an agent a_h with maximum *Joint Benefit* as a *role-exchange partner* exemplifies the *individual optimal role-exchange* scheme. *Joint Benefit* is calculated by computing *individual gain*, *individual loss*, *social gain*, and *social loss*. Each of these parameters is explained next with respect to IAS.

Individual loss for a role-exchange agent pair $a_s - a_h$, denoted by $IndLoss_{a_s, a_h}$, is the measure of *capability sacrifice* required by the role exchange. The *capability sacrifice* is the wastage of available resources or capabilities of agents a_s and a_h incurred due to performing the exchanged roles. The resources of an agent get wasted if performing a role requires its capabilities at a lesser value than its current role. Hence, calculating *individual loss* helps us in identifying a role-exchange agent pair with the least resource wastage.

$IndLoss_{a_s, a_h}$ is calculated by first computing the *future* value of each capability for agents a_s and a_h , i.e., capability usage while performing the exchanged roles, and then by calculating the difference between the *current* values and the computed *future* values. The summation of this difference for each capability c_j will give $IndLoss_{a_s, a_h}$. The *future* value of a capability c_j for an agent a_i at time $t + 1$ is denoted by

$v_{c_j,t+1}^{a_i}$. It is equal to either the required value of c_j by a role r_k at time t , i.e., $v_{c_j,t}^{r_k}$, or the maximum value of capability c_j that the $CRU_{c_j}^{a_i}$ can generate for agent a_i , i.e., $MaxCap_{c_j}^{a_i}$, whichever is smaller. This is because, as we discussed earlier, an agent can provide a capability value up to the maximum limit which the agent's CRU can generate irrespective of the role's requirements. Hence, if at time t , agent a_s and a_h are performing the roles r_1 and r_2 , respectively, then at time $t + 1$, the future values of a capability c_j for the exchanged roles r_2 of agent a_s and the exchanged roles r_1 of agent a_h are given by (4) and (5), respectively

$$v_{c_j,t+1}^{a_s} = \text{Minimum} \left(v_{c_j,t}^{r_2}, MaxCap_{c_j}^{a_s} \right) \quad (4)$$

$$v_{c_j,t+1}^{a_h} = \text{Minimum} \left(v_{c_j,t}^{r_1}, MaxCap_{c_j}^{a_h} \right). \quad (5)$$

Next, the difference between the future and current values of each capability c_j is calculated. The summation of these differences for each capability c_j multiplied with the capability weight of c_j gives the $IndLoss_{a_s,a_h}$, as shown in (6). The capability weight of a capability c_j for an agent a_i , i.e., $Weight_{c_j}^{a_i}$, is the preference or importance of c_j in the role r_k , which agent a_i is performing. Considering every individual capability c_j of the agents and considering the preference of c_j in the form of capability weight preserve and present the additive and lexicographic attributes, respectively, of our joint-benefit utility. However, we will not consider the loss due to a decrease in the value of ULD capabilities in the individual loss calculation because this will be taken into account by social gain, which is discussed later in this section

$$IndLoss_{a_s,a_h} = \sum_{c_j \in CAP-UCAP} \left(Weight_{c_j}^{a_s} \left(v_{c_j,t}^{a_s} - v_{c_j,t+1}^{a_s} \right) + Weight_{c_j}^{a_h} \left(v_{c_j,t}^{a_h} - v_{c_j,t+1}^{a_h} \right) \right). \quad (6)$$

Individual gain for a role-exchange agent pair $a_s - a_h$, denoted by $IndGain_{a_s,a_h}$, is the measure of the personal benefit that an agent a_s receives due to a role exchange. In an ISR operation, a decrease in the damage of a $CRU_{c_j}^{a_i}$ can be considered as a personal benefit to agent a_i . In our experiments, the decrease in a CRU damage is calculated as the difference between the damage to the $CRU_{c_j}^{a_s}$ at time t while performing a role r_k at UAI_t^{IAS} , i.e., $CRUDamage_{c_j,t}^{a_s}(UAI_t^{IAS})$, and the average (or arithmetic mean) of damages to the $CRU_{c_j}^{a_s}$ at time t' while performing the same role at $UAI_t^{a_h}$, i.e., $AvgCRUDamage_{c_j,t'}^{a_s}(UAI_t^{a_h})$. Time t' represents all the previous instances of ISR operations, which are similar to the current operation. Hence, it gives the reference to the past experience in similar condition. The summation of these benefits for all the CRUs of the agent a_s will give us the $IndGain_{a_s,a_h}$, as shown in (7). Ideally, the $IndGain_{a_s,a_h}$ will be higher when a_h belongs to an agent culture with high UAI level because, at a high UAI level, an agent a_s will behave more stealthily and hence incur less damage due to less severe attacks and the counterattacks with the enemy units

$$IndGain_{a_s,a_h} = \sum_{c_j \in CAP} \left(CRUDamage_{c_j,t}^{a_s}(UAI_t^{IAS}) - AvgCRUDamage_{c_j,t'}^{a_s}(UAI_t^{a_h}) \right). \quad (7)$$

Social gain for a role-exchange agent pair $a_s - a_h$, denoted by $SocialGain_{a_s,a_h}$, is the measure of social benefit that an IAS receives due to the role exchange. Social gain is calculated with reference to the maximal social gain, which a human supervisor defines as one of the goals during a mission planning. For example, during an ISR mission planning, Moe can dictate that performing ISR on all the target locations at UAI_t^{Moe} will give maximal social gain or $MaximalSocialGain^{Moe}$. Moe also defines a social gain reduction or $SocialGainReduction^{Moe}$ that represents a per unit value by which the social gain would reduce if an agent a_s opts a role exchange with agent a_h , where $UAI_t^{a_h}$ is not exactly equal to UAI_t^{Moe} . Hence, the absolute difference between $UAI_t^{a_h}$ and UAI_t^{Moe} gives the factor by which $SocialGainReduction^{Moe}$ is deducted from $MaximalSocialGain^{Moe}$. Equation (8) shows how $SocialGain_{a_s,a_h}$ can be calculated using the maximal social gain and social gain reduction. Based on this equation, we can say that, the higher the closeness of UAI level of agent a_h to the UAI_t^{Moe} value, the larger the value of $SocialGain_{a_s,a_h}$

$$SocialGain_{a_s,a_h} = MaximalSocialGain^{Moe} - (|UAI_t^{Moe} - UAI_t^{a_h}|) SocialGainReduction^{Moe}. \quad (8)$$

Lastly, social loss for a role-exchange agent pair $a_s - a_h$, denoted by $SocialLoss_{a_s,a_h}$, is the measure of loss to the IAS incurred due to agent a_h leaving its current role unfinished due to the role exchange with agent a_s . Leaving a role unfinished may require rollback of some completed parts of a role, which leads to a hidden loss to an IAS.

D. IDV and PDI Inclusion in IAS

Incorporating IDV and PDI in our algorithm helped us magnify the effect of these cultural parameters on the overall benefit of the approach. Moreover, IDV and PDI appear together in a culture; therefore, experiment with one parameter will automatically reveal other parameter's effect on an IAS.

NCW emphasizes on making ISR components (or agents) collectivist [16]. Such collectivist agents will be very sensitive to the demand of their social context and more responsible to assumed the needs of others, as well as less insistent on pursuing personal goals that might jeopardize group benefits [16]. Also, collectivist agents similar to collectivist people are more likely to attribute their own and others behavior to situational rather than dispositional causes [11]. We can easily visualize the Pareto optimality nature in such collectivist culture agents.

The reason for including PDI in our research is to study the effect of hierarchical relationship or inequality, which is inevitable in a work setting where superior-subordinate relationships exist. According to Cummings *et al.*, the human supervisory control in UAV operation is hierarchical in nature [2]. Although we have not made these agents hierarchical with respect to each other in our approach, we did introduce hierarchy using agent-supervisor relation. Hence, by including this parameter, we could be able to see the effect of inequality, which is always certain while working at Sheridan's fifth and

sixth automation levels [21]. We discussed in Section II that the agents at high PDI (or low IDV) consider supervisor as a benevolent dictator and hence obey her directives.

IV. IMPLEMENTATION AND EXPERIMENTS

The collaboration of heterogeneous entities in ISR operations is customary. This leads to the scenarios where different components become responsible for different tasks, working toward a common goal. Our implementation of role-exchange algorithm replicates the nature of collaboration most commonly found in ISR. A simulation is designed and fully implemented using programming language C and OpenGL, which is used to generate data that illustrate several metrics. We have presented UVs as agents. In this section, we begin by explaining the nature of agent capabilities. Next, the details of simulator scenarios are offered.

A. Capabilities and Their ULD Nature

As discussed before, we used the capabilities of the ISR agents as the main attribute that distinguishes among agents in our approach. An agent can have any number and any type of capabilities. To make our experiments tractable and to mimic a battlefield scenario, we have used four most basic military operations, namely, *attack* (**A**), *patrol* (**P**), *reconnaissance* (**R**), and *communicate* (**C**), as capabilities available to an ISR agent.

Attack (**A**) in the warfare is mandatory for any manned or unmanned combat-capable ISR. A direct inevitable consequence of attack is ambiguity and unfamiliarity. Hence, with increase in **A**, uncertainty increases or UAI decreases in an IAS. Therefore, **A** is a member of ULD capability set **UCAP**, i.e., $\mathbf{A} \in \mathbf{UCAP}$. Equation (9) represents the relationship of **A** and UAI

$$\mathbf{A} \propto \frac{1}{\text{UAI}}. \quad (9)$$

Patrol (**P**) and *reconnaissance* (**R**) capabilities are also indispensable for any ISR operation and play vital roles. However, they generally do not lead to uncertainty. Therefore, **P** and **R** capabilities are non-ULD capabilities. Lastly, we have considered *communicate* (**C**) capability, which exemplifies the communicative nature of an ISR agent. Using **C**, battlefield awareness gets increased as it helps to share information about the adversary's activities and current status of available resources. This information will alleviate the state of confusion and ambiguity among peer group members and aid in the decision-making process of fellow ISR agents. This implies that, with increase in **C**, uncertainty decreases or UAI increases, as shown in

$$\mathbf{C} \propto \text{UAI}. \quad (10)$$

Hence, based on (10) and (11), the relationship between **C** and UAI is opposite to that of **A** and UAI or we can say that the inverse of **C**, i.e., **C'**, is directly proportional to **A**, as given by

$$\mathbf{C}' \propto \mathbf{A}. \quad (11)$$

Understanding the relation among the ULD capabilities is very important. Let us consider a situation where *Moe* wants to control the current UAI level of an IAS. Controlling several parameters requires that the direction of all the parameters is conjugate. If two ULD capabilities are growing in the reverse direction, for example, capabilities **A** and **C**, it is very difficult to increase or decrease both of them simultaneously using a single control tool. According to the proposed definition of supervisory control discussed in Section III, the more appropriate supervisory control would be the one that removes the micro-management of trivial tasks. However, to keep track of several control tools at a same time will lead to micromanagement. Hence, selecting **C'** instead of **C** will avoid such a situation and allows *Moe* to control these two ULD capabilities using a single control parameter. Moreover, the presence of **A** and **C'** at the same time also allows us to examine the effect of having more than one ULD capabilities in the ISR system.

B. Battle Space Simulation

For the simulation of our approach, we considered the three cultural parameters of MOTL, namely, UAI, IDV, and PDI, in a cumulative fashion and explored their impact and interrelationships in the ISR context. Our simulator also studies each parameter with respect to two different IASs, namely, *3cap* and *4cap*. A *3cap* IAS has three capabilities, whereas a *4cap* IAS has four.

A screenshot of our simulator is shown in Fig. 3. For simulation purposes, we assumed a *territory of interest* that spreads from the point **X** to the point **Y** in the simulated battlefield and is required to be explored using ISR. Several enemy units are also depicted in the simulated battlefield in the form of diamond-shape icons. The IAS agents are presented in the form of rectangular boxes. The intensity of the color of these boxes represents their current UAI levels; the darker shades symbolize agents with higher current UAI. The simulator also provides a control for *Moe* to explicitly tune the UAI level at run time using the command line. Hence, this control allows him to adjust $\text{UAI}_t^{\text{IAS}}$ by specifying $\text{UAI}_t^{\text{Moe}}$ at random points multiple times during an ISR operation. Using this control, an operator can act upon any changing or exigent situation by imposing a UAI to the IAS, which is traversing the *territory of interest*.

However, the screenshot shown in Fig. 3 presents a scenario when $\text{UAI}_t^{\text{Moe}} = \text{UAI}_t^{\text{IAS}} = 0$, i.e., no UAI is imposed over IAS, hence no role exchange would take place while traversing the *territory of interest*. Therefore, each agent $\mathbf{a}_i$ will operate on its $\text{BaseUAI}^{\mathbf{a}_i}$ and incur a high *total damage*. The *group loss* field in the screenshot is showing the reset value of zero, which should not be confused with no group loss. We can also observe a cleared region along the convoy's path in the *territory of interest* in this screenshot. This area represents the fully destroyed enemy forces by the attack of IAS agents. The encircled diamonds symbolize the damaged enemy forces that were attacked but not fully destroyed. The attacking incurs a total damage or total CRU damage of 10.5 units to the *3cap* IAS. This damage is the summation of the damages sustained by all the CRUs in the IAS due to the adversary's

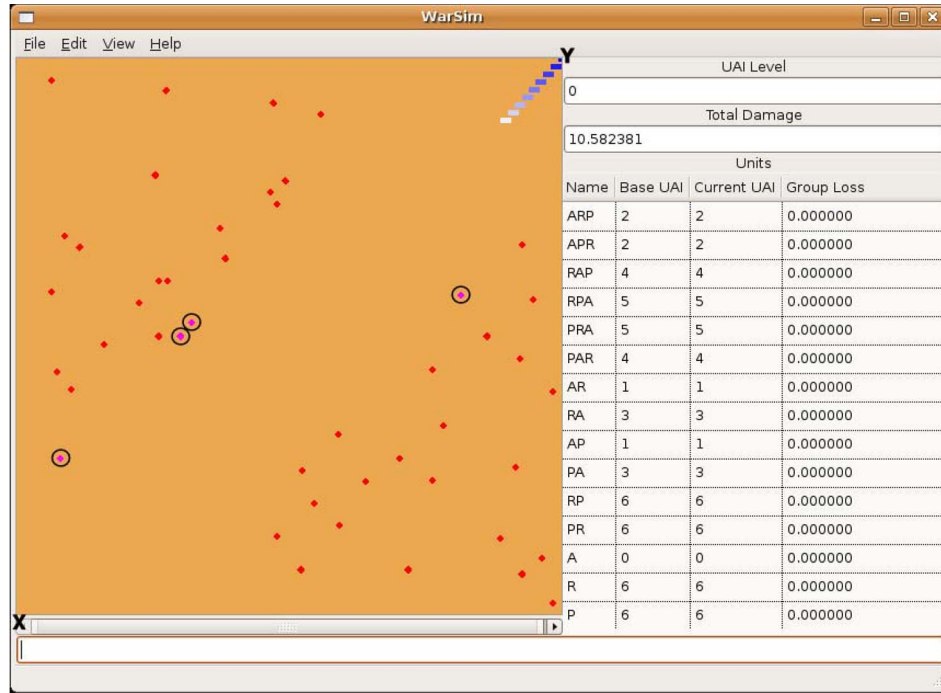


Fig. 3. Snapshot of simulator battlefield for 3cap IAS.

TABLE I
3capIAS

Agent (a_i)	Initial Role (r_k)	Maximum Attack Capability ($MaxCap_A^{a_i}$)	Maximum Patrol Capability ($MaxCap_P^{a_i}$)	Maximum Reconnaissance Capability ($MaxCap_R^{a_i}$)	Base UAI Level ($BaseUAI^{a_i}$)
A	A	100	-	-	0
AR	AR	60	-	40	1
AP	AP	60	40	-	
ARP	ARP	50	20	30	2
APR	APR	50	30	20	
RA	RA	40	-	60	3
PA	PA	40	60	-	
RAP	RAP	30	20	50	4
PAR	PAR	30	50	20	
RPA	RPA	20	30	50	5
PRA	PRA	20	50	30	
R	R	-	-	100	6
P	P	-	100	-	
RP	RP	-	40	60	
PR	PR	-	60	40	

counteractions. The total CRU damage for 4cap ISR is even higher, i.e., 13.7 units (not shown in the figure) because there are more agents per UAI-level-defined cultures.

In our experiments, the agents and roles are named or labeled using the combination of the capabilities **A**, **P**, **R**, and **C'**. For example, if a role r_k requires the capabilities $c_1, c_2, \dots, c_n$ at values $v_{c_1,t}^{r_k}, v_{c_2,t}^{r_k}, \dots, v_{c_n,t}^{r_k}$, respectively, and $v_{c_1,t}^{r_k} \geq v_{c_2,t}^{r_k} \geq \dots \geq v_{c_n,t}^{r_k}$, then r_k is labeled using capabilities such that the capabilities are arranged in the descending order of the values. Hence, the role label of r_k would be $c_1, c_2, \dots, c_n$. In addition, if $v_{c_j,t}^{r_k} = 0$, then c_j will not be a part of a role's label. Initially, each of these labeled roles is assigned to an agent a_i such that, for every capability c_j , $v_{c_j,t}^{a_i} = v_{c_j,t}^{r_k}$. Moreover, if we use the same naming convention for labeling the agents, then the functional label of the a_i

will exactly be same as that of the label of r_k . The roles, agents, and capabilities available in 3cap and 4cap IASs are discussed next, which will clearly present this naming convention technique.

1) 3cap IAS: In a 3cap IAS, capabilities **A**, **P**, and **R** were used. Hence, $CAP = \{A, P, R\}$ and $UCAP = \{A\}$. The total number of capability combinations, taken one, two, and three capabilities at a time from three capabilities, is ${}^3P_1 + {}^3P_2 + {}^3P_3 = 15$. Hence, 15 agents and 15 roles are initialized using the naming convention discussed earlier. Table I shows the agents along with their initially assigned roles and $MaxCap_{c_j}^{a_i}$ for each capability c_j . $BaseUAI^{a_i}$ is also shown in the table, which is calculated using (2). The 15 agents are divided into seven agent cultures using $BaseUAI^{a_i}$, each of which represents a different UAI level.

TABLE II
COMPUTATION OF $IndLoss_{AR, a_h}$ WHEN $UAI_t^{IAS} = 0$ AND $UAI_t^{Moe} = 3$

	A	R	-	R	A	-			A	R	-	P	A	-		
MaxCap	60	40	-	60	40	-			MaxCap	60	40	-	60	40	-	
CurrentCap	60	40	-	60	40	-			CurrentCap	60	40	-	60	40	-	
FutureCap	40	40	-	40	40	-			FutureCap	40	0	-	0	40	-	
IndLoss	-	0	0	12	0	0	12		IndLoss	-	16	0	36	0	0	52

	A	R	-	R	A	P			A	R	-	P	A	R		
MaxCap	60	40	-	50	30	20			MaxCap	60	40	-	50	30	20	
CurrentCap	60	40	-	50	30	20			CurrentCap	60	40	-	50	30	20	
FutureCap	30	40	-	40	30	0			FutureCap	30	20	-	0	30	20	
IndLoss	-	0	0	5	0	4	09		IndLoss	-	8	0	25	0	0	33

	A	R	-	R	P	A			A	R	-	P	R	A		
MaxCap	60	40	-	50	30	20			MaxCap	60	40	-	50	30	20	
CurrentCap	60	40	-	50	30	20			CurrentCap	60	40	-	50	30	20	
FutureCap	20	40	-	40	0	20			FutureCap	20	30	-	0	30	20	
IndLoss	-	0	0	5	9	0	14		IndLoss	-	4	0	25	0	0	29

	A	R	-	R	-	-			A	R	-	P	-	-		
MaxCap	60	40	-	100	-	-			MaxCap	60	40	-	100	-	-	
CurrentCap	60	40	-	100	-	-			CurrentCap	60	40	-	100	-	-	
FutureCap	0	40	-	40	-	-			FutureCap	0	0	-	0	-	-	
IndLoss	-	0	0	60	0	0	60		IndLoss	-	16	0	100	0	0	116

	A	R	-	R	P	-			A	R	-	P	R	-		
MaxCap	60	40	-	60	40	-			MaxCap	60	40	-	60	40	-	
CurrentCap	60	40	-	60	40	-			CurrentCap	60	40	-	60	40	-	
FutureCap	0	40	-	40	0	-			FutureCap	0	40	-	0	40	-	
IndLoss	-	0	0	12	16	0	18		IndLoss	-	0	0	36	0	0	36

Initially, when UAI_t^{Moe} is also zero, each agent a_i performs a role r_k using c_j at $v_{c_j,t}^{a_i} = MaxCap_{c_j}^{a_i}$, and the agent culture with the lowest UAI value, i.e., $UAI_t^A = 0$, dictates the $UAI_t^{IAS} = 0$. Furthermore, to mathematically elucidate the calculations in the experiments, for each c_j , an agent a_i is initialized with the capability value $v_{c_j,t}^{a_i}$ such that their summation $v_{c_1,t}^{a_i} + v_{c_2,t}^{a_i} + \dots + v_{c_n,t}^{a_i} = 100$.

2) *4cap IAS*: In a *4cap* IAS, in addition to **A**, **P**, and **R**, capability **C'** is also used. Hence, $CAP = \{A, P, R, C'\}$, and $UCAP = \{A, C'\}$. The total number of capability combinations is 64 in this case. Hence, 64 agents and 64 roles are initialized using the same naming convention. For each agent, $BaseUAI^{a_i}$ was calculated using (2). Based upon these values, the 64 agents are divided into nine agent cultures, each of which represents a different UAI level. The main reason to consider such a larger scenario is to study and reveal the impact of the presence of more than one ULD capabilities in an IAS. In the *4cap* IAS, 50% of the agent's capabilities are ULD in comparison with 33% of the *3cap* IAS. Moreover, the difference in ULD capabilities' preferences can also exist in an IAS, which might lead to other significant consequences. However, in our experiments with *4cap* IAS, we gave the same preference to both ULD capabilities.

To discuss the calculation involved in steps 3 and 4 of our algorithm (Fig. 2), let us consider the UAI tuning in both directions: low-to-high and high-to-low UAI levels.

3) *Tuning From Low UAI to High UAI Level*: In order to change the UAI_t^{IAS} of an IAS, *Moe* can impose a UAI_t^{Moe}

on an IAS. For instance, at time t when $UAI_t^{IAS} = 0$, if *Moe* specifies $UAI_t^{Moe} = 3$, then according to the role-exchange algorithm discussed, few agents have to change their UAI levels to this desired value. Since $UAI_t^{Moe} > UAI_t^{IAS}$, each agent a_s with $UAI_t^{as} < 3$ would become a *role-exchange-seeker* agent, whereas every agent a_h with $UAI_t^{ah} \geq 3$ would be a *role-exchange-helper* agent. Hence, in a *3cap* IAS, agents at zero, one, and two UAI levels, namely, agent **A**, **AR**, **AP**, **ARP**, and **APR**, are agents a_s . The PL_t^{as} will be computed for each of these agents, using (3a). For example, the PL of the agent **AR**, which is performing a role **AR**, at time t will contain the helper agents at UAI levels three, four, five, and six, i.e., $PL_t^{AR} = \{RA, PA, RAP, PAR, RPA, PRA, P, R, RP, PR\}$. After identifying a_s and a_h agents, the next step is to calculate the *Joint Benefit*, which requires *individual loss*, *individual gain*, *social gain*, and *social loss* to be calculated first.

For example, let us consider the *individual loss* calculation for agent **AR** or $IndLoss_{AR, a_h}$ of a *3cap* IAS, where a_h represents a helper agent in the PL_t^{AR} . The calculation of $IndLoss_{AR, a_h}$ is shown in Table II. The ten *individual losses* calculated using (6), one for each agent a_h , are shown in a separate subtable of Table II. In the subtables, *MaxCap* represents $MaxCap_{c_j}^{a_i}$, *CurrentCap* represents $v_{c_j,t}^{a_i}$, and *FutureCap* represents $v_{c_j,t+1}^{a_i}$.

For instance, let us look at the calculation performed to obtain the $IndLoss_{AR, RAP}$. First, the *FutureCap* (or $v_{c_j,t+1}^{a_i}$) is calculated for all the capabilities of agent **AR** and

RAP using $MaxCap$ (or $MaxCap_{c_j}^{a_i}$) and $CurrentCap$ (or $v_{c_j,t}^{a_i}$). As $UAI_t^{IAS} = 0$, therefore, $v_{c_j,t}^{a_i} = MaxCap_{c_j}^{a_i}$. Hence, for agent **AR**, the $MaxCap_A^{AR} = v_{A,t}^{AR} = 60$, $MaxCap_P^{AR} = v_{P,t}^{AR} = null$ (because capability **P** is not available with agent **AR**), and $MaxCap_R^{RAP} = v_{R,t}^{RAP} = 40$. Similarly, for agent **RAP**, $MaxCap_A^{RAP} = v_{A,t}^{RAP} = 30$, $MaxCap_P^{RAP} = v_{P,t}^{RAP} = 20$, and $MaxCap_R^{RAP} = v_{R,t}^{RAP} = 50$. The $FutureCap$ calculation of capability **R** for both the agents using (4) and (5) is performed as

$$\begin{aligned} v_{R,t+1}^{AR} &= \text{Minimum}(v_{R,t}^{RAP}, MaxCap_R^{AR}) \\ &= \text{Minimum}(50, 40) \\ &= 40 \\ v_{R,t+1}^{RAP} &= \text{Minimum}(v_{R,t}^{AR}, MaxCap_R^{RAP}) \\ &= \text{Minimum}(40, 50) \\ &= 40. \end{aligned}$$

After calculating the remaining $FutureCap$ in similar fashion, the $IndLoss_{AR,RAP}$ is calculated using $FutureCap$ and weight of all the capabilities as the second step. In our experiments, the weight of a capability in a role is calculated with respect to its required value in that role. The bigger the required value of a capability c_j in a role r_k , i.e., $v_{c_j,t}^{r_k}$, the higher the value of $Weight_{c_j}^{a_i}$. We have used 1% of the $v_{c_j,t}^{r_k}$ as $Weight_{c_j}^{a_i}$. For example, the weights of capabilities **R**, **A**, and **P** for agent **RAP** performing role **RAP** are 0.5, 0.3, and 0.2, respectively. However, according to (6), the ULD capabilities of the *role-exchange-seeker* agent are not considered in the *individual loss* calculation. The calculation for $IndLoss_{AR,RAP}$ using (6) is performed as

$$\begin{aligned} IndLoss_{AR,RAP} &= (Weight_R^{AR} (v_{R,t}^{AR} - v_{R,t+1}^{AR})) \\ &\quad + (Weight_R^{RAP} (v_{R,t}^{RAP} - v_{R,t+1}^{RAP})) \\ &\quad + (Weight_A^{RAP} (v_{A,t}^{RAP} - v_{A,t+1}^{RAP})) \\ &\quad + (Weight_P^{RAP} (v_{P,t}^{RAP} - v_{P,t+1}^{RAP})) \\ &= (0.4(40 - 40)) + (0.5(50 - 40)) \\ &\quad + (0.3(30 - 30)) + (0.2(20 - 0)) \\ &= 0 + 5 + 0 + 4 = 9. \end{aligned}$$

The value of $IndLoss_{AR,RAP}$, i.e. nine units, is minimum in all the *individual losses* shown in Table II. Therefore, if agent **AR** exchanges role with agent **RAP**, then the *individual loss* would be least. However, the remaining three components of the *joint benefit* will also influence the agent **AR**'s final partner selection.

The *individual gain* or the personal benefit of the same agent **AR** due to the role exchange with a *role-exchange-helper* agent a_h , i.e., $IndGain_{AR,a_h}$, is calculated using (7). For example, if $c_j = R$ and $a_h = RAP$, then $CRUDamage_{R,t}^{AR}(UAI_t^{IAS}) - AvgCRUDamage_{R,t}^{AR}(UAI_t^{IAS})$ gives the decrease in

the damage of CRU_R^{AR} due to the role exchange with agent **RAP**. The decrease in the damage of other CRUs of agent pair **AR** – **RAP** will also be calculated in similar fashion. The summation of all these values would give us $IndGain_{AR,RAP}$.

The $SocialGain_{AR,a_h}$ is calculated using (8). In our experiment, IAS performing ISR from point **X** to point **Y** at UAI_t^{Moe} is considered as the $MaximalSocialGain^{Moe}$. This implies that $SocialGain_{AR,RA}$ and $SocialGain_{AR,PA}$ will be maximum or equal to $MaximalSocialGain^{Moe}$. Also, the *social gain* reduces by the largest factor or to minimum if agent **AR** exchanges role with any one of the four agents at UAI level 6 or $UAI_t^{ah} = 6$.

Finally, for $SocialLoss_{AR,a_h}$, we used randomly generated values to mimic the percentage of a role that is already performed by agent a_h when it is asked for the role exchange with agent **AR**. However, if we keep IDV and PDI inactivated, then the *social loss* would be less influential toward the final selection of partners.

4) *Tuning from High UAI to Low UAI Level*: *Moe* can also decrease the UAI of an IAS during an ISR operation. Let us consider an intermediate instance, when $UAI_t^{IAS} = 5$ and *Moe* specifies $UAI_t^{Moe} = 2$. Since $UAI_t^{Moe} < UAI_t^{IAS}$, each agent a_s with $BaseUAI^{a_s} < UAI_t^{Moe}$ would become a *role-exchange-seeker* agent. Therefore, agents at UAI levels 0 and 1 or agent **A**, **AR**, and **AP** would become *role-exchange-seeker* agents a_s 's. To identify the *role-exchange-helper* agents, let us consider the agent **AP** that is performing the role **RPA** at time t . According to (3b), the PL_t^{AP} contains 12 agents: **ARP**, **APR**, **RA**, **PA**, **RAP**, **PAR**, **RPA**, **PRA**, **P**, **R**, **RP**, and **PR**. Agents with $BaseUAI^{ah} \geq UAI_t^{IAS}$, i.e., agents at UAI levels 5 and 6, would be included into PL_t^{AP} as they are. Therefore, $CurrentCap$ (or $v_{c_j,t}^{ah}$) for these agents is need not be equal to $MaxCap$ (or $MaxCap_{c_j}^{ah}$). The agents at UAI levels 2, 3, and 4 for which $UAI_t^{Moe} \leq BaseUAI^{ah} < UAI_t^{IAS}$ will be included to PL_t^{AP} after resetting their UAI to *base* UAI; therefore, $CurrentCap$ is equal to $MaxCap$ for these agents. The *joint benefit* will be calculated similar to the previous example.

V. RESULTS AND CONCLUSION

This section presents the results of the experiments discussed in Section IV as the validation of our approach. Generated data are plotted as graphs for two scenarios. The first scenario depicts the UAI influence on the CRU damage of *3cap* and *4cap* IASs, whereas in the second scenario, the UAI-IDV influence on *Average IndGain* of *3cap* IAS is presented.

A. UAI Influence on IAS

The adversary's counteractions on IAS agents while traversing the *territory of interest* have damaged the CRUs of IAS. The summation of the damages sustained by all CRUs of each agent in an IAS is called the *total CRU damage* or $TotalCRUDamage^{IAS}$. We used our simulator to illustrate the effect of different UAI levels of *3cap* and *4cap* IASs on $TotalCRUDamage^{IAS}$. The range of UAI levels for *3cap*

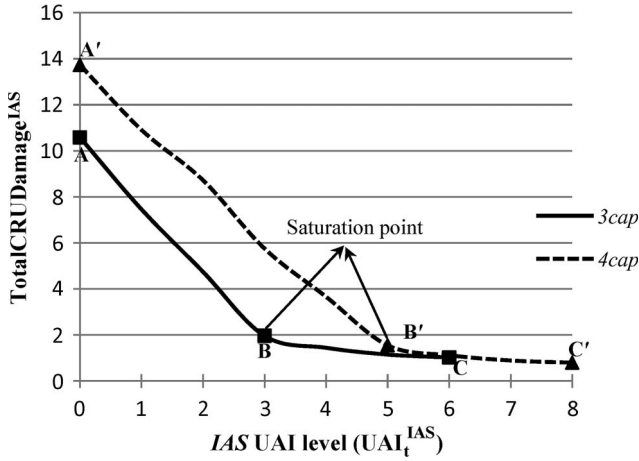


Fig. 4. $TotalCRUDamage^{IAS}$ at different UAI levels for 3cap and 4cap IASs.

IAS is zero to six, whereas it is zero to eight for 4cap IAS. In Fig. 4, the curves ABC and $A'B'C'$ represent the total CRU damages incurred to 3cap and 4cap IASs, respectively, at different UAI levels.

$TotalCRUDamage^{IAS}$ values for both 3cap and 4cap IASs are maximum at $UAI_t^{IAS} = 0$, as shown by points A and A' in the graph. Since no UAI is enforced on IASs at this level, therefore, every agent is performing its role using $MaxCap_{c_j}^{a_i}$. We can also observe that the $TotalCRUDamage^{4cap} > TotalCRUDamage^{3cap}$ for all the UAI values. This difference is due to the difference in the numbers of agents and the difference in the numbers of ULD capabilities in two IASs. Therefore, we can conclude that, the larger the group of agents, the less stealthy it can keep its moves and, similarly, the larger number of ULD capabilities are difficult to control.

In spite of these differences between the two IASs, the gradual decrease in their CRU damage can be observed with an increase in UAI_t^{IAS} . This shows the benefit of using our approach. The decrease is very sharp for 3cap IAS from point A to point B , i.e., from CRU damage of 10.3 units at $UAI_t^{IAS} = 0$ to 2 units at $UAI_t^{IAS} = 3$, but the changes in UAI value from point B to C do not lead to significant benefits. Similarly, the decrease in CRU damage became almost saturated after point B' for 4cap IAS. Hence, $UAI_t^{IAS} = 3$ and $UAI_t^{IAS} = 5$ are the saturation points of the UAI values for 3cap and 4cap IASs, respectively, beyond which the CRU damage did not decreased significantly.

B. IDV Influence on IAS

In order to illustrate the effect of IDV on the IAS, we have augmented IDV with UAI and have examined its influence on the individual gain of agents in an IAS. We considered three different UAI_t^{Moe} , namely, 1, 3, and 6, for 3cap IAS, which represent the UAI spectra at lowest, median, and highest values, respectively. We computed the effect of IDV values that range from -5 to $+5$ with each of these three UAI levels on the individual gain. The IAS at $IDV = -5$ represents a collectivist IAS that possesses agents with a highest collectivist belief, whereas $IDV = +5$ represents an individualistic IAS where agents possess a highest individualistic belief.

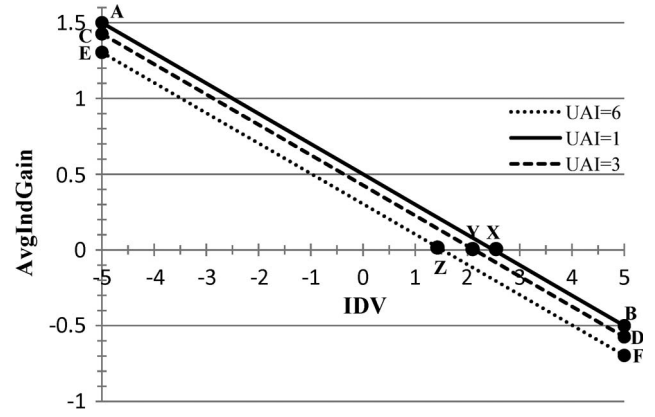


Fig. 5. $AvgIndGain$ versus IDV for 3cap IAS.

First, the maximum value of $IndGain_{a_s, a_h}$ is identified for every role-exchange-seeking agent a_s at a particular UAI_t^{Moe} and an IDV, and then, the arithmetic mean of these individual gain values or $AvgIndGain$ is computed. Fig. 5 shows the $AvgIndGain$ of a 3cap IAS at different values of IDV for three different UAI_t^{IAS} values. The line segments AB , CD , and EF represent $AvgIndGain$ at UAI_t^{IAS} : 1, 3 and 6, respectively. The $AvgIndGain$ at $IDV = -5$ for all the three UAI_t^{Moe} values, shown as points A , C , and E , is maximum. However, with the increase in IDV value, it decreased and even became negative after points X , Y , and Z . Since the role-exchange-helper agents in a collectivist IAS implicitly possess high PDI that makes them strictly obey the UAI_t^{Moe} , therefore, a higher individual gain is attained as compared to individualistic IAS, where the low PDI allows agents to work in a more unrestrained manner. This dictates that the agents in a collectivist IAS generate more positive individual gain than the agents in an individualistic IAS.

We can also observe that the $AvgIndGain$ is less for higher UAI_t^{Moe} because a high UAI_t^{Moe} leads to more numbers of role-exchange-seeker agents in an IAS and hence reduces the $AvgIndGain$ due to greater competition among these agents for the most beneficial role-exchange-helper agents. For instance, at $UAI_t^{Moe} = 1$ in 3cap IAS, there was only one a_s agent. Hence, without any contender, the role exchange occurs with an agent that leads to maximum $IndGain_{a_s, a_h}$. On the contrary, when $UAI_t^{Moe} = 3$ and $UAI_t^{Moe} = 6$, the numbers of a_s agents were 5 and 11, respectively, which lead to a race condition and keep agents away from the most benefited matches for their role exchange.

VI. FUTURE WORK

Reported research can be extended in many directions. For instance, a power-based hierarchy of agents can be introduced on the top of UAI and IDV in an agent society. This would merge the advantages of MOTL with other proposed hierarchical control models [2]. Furthermore, the hierarchical division of agents would also allow us to establish supervisor-subordinate relationship among agents. In such an agent society, a supervisor agent would control or manage its subordinate agents on the commands and directions of human supervisor.

Implementing such a supervisor–subordinate agent relationship using MOTL would lead us to another futuristic concept that we can call “agent on the loop” (AOTL). AOTL will be useful to engender a simplified supervisory control for envisioned fully automated multiplatform cooperative UV swarms.

Similarly, the IDV can be more deeply explored in a platform diverse agent community. This would reveal the consequences when a member of “newcomer” agents or out-group agents interacts with the member of “historical” or in-group agents in an agent society [1]. Multiplatform software coordination is an example of such interaction, where culturally alike group of software, which is a pivotal part that forms in-group—interacts with many culturally dissimilar out-group software for a collective goal attainment in a cooperative task.

Lastly, the priority given to the two ULD capabilities A and C' was equal in this paper; however, ULD capabilities with different priorities can also be studied. An extensive exploration of complex decision-making systems, such as defense or aerospace systems, would unleash other ULD and non-ULD capabilities with their relative importance and hence help materializing real-world MOTL-based supervisory control.

REFERENCES

- [1] S. Bochner and B. Hesketh, “Power distance, individualism/collectivism, and job-related attitudes in a culturally diverse work group,” *J. Cross-Cult. Psychol.*, vol. 25, no. 2, pp. 233–257, 1994.
- [2] M. L. Cummings, S. Bruni, S. Mercier, and P. J. Mitchell, “Automation architecture for single operator, multiple UAV Command and Control,” *Int. Command Control J.*, vol. 1, no. 2, pp. 1–24, 2007.
- [3] R. M. Cyert and J. G. March, *A Behavioral Theory of the Firm*. Englewood Cliffs, NJ: Prentice-Hall, 1963.
- [4] Network Centric Warfare: Department of Defense Report to Congress. Washington, DC: OSD, U.S. Dept. Defense, 2001.
- [5] The Implementation of Network-Centric Warfare. Washington, DC: Office Force Transformation, OSD, Dept. Defense, 2003.
- [6] Unmanned Aircraft Systems (UAS) Roadmap, 2005–2030. Washington, DC: OSD, Dept. Defense, 2005.
- [7] P. C. Fishburn, “Utility theory,” *Manage. Sci.*, vol. 14, no. 5, pp. 335–378, Jan. 1968.
- [8] H. Hexmoor, B. McLaughlan, and G. Tuli, “Towards a natural human role in supervising complex control systems,” *J. Exp. Theor. Artif. Intell.*, vol. 21, no. 1, pp. 59–77, Mar. 2009.
- [9] G. Hofstede, *Culture's Consequences: International Differences in Work-Related Values*. Beverly Hills, CA: Sage, 1980.
- [10] G. Hofstede, *Culture's Consequences: Comparing Values, Behaviors, Institutions and Organizations Across Nations*. Thousand Oaks, CA: Sage, 2001.
- [11] J. Jaspers and M. Hewstone, “Cross-cultural interaction, social attributions and inter-group relations,” in *Cultures in Contact: Studies in Cross-Cultural Interaction*. Oxford, U.K.: Pergamon, 1982.
- [12] S. Kalenka and N. R. Jennings, “Socially responsible decision making by autonomous agents,” in *Proc. Int. Conf. Complex Syst.*, 1997, pp. 135–149.
- [13] U. Kim, H. C. Triandis, C. Kaqitcibasi, S. Choi, and G. Yoon, Eds., *Individualism and Collectivism: Theory, Method, and Applications*, vol. 18. Thousand Oaks, CA: Sage, 1994, ser. Cross-Cultural Research and Methodology Series.
- [14] D. Lambert and J. Scholz, “A Dialectic for Network Centric Warfare,” in *Proc. 10th Int. Command Control Res. Technol. Symp.*, 2005. [Online]. Available: http://www.dodccrp.org/events/10th_ICCRTS/CD/papers/016.pdf
- [15] H. R. Markus and S. Kitayama, “Culture and the self: Implications for cognition, emotion, and motivation,” *Psychol. Rev.*, vol. 98, no. 2, pp. 224–253, 1991.
- [16] Y. K. Ng, *Welfare Economics: Towards a More Complete Analysis*. Hampshire, NY: Palgrave MacMillan, 2004, ch. 2.
- [17] W. A. Owens, *Lifting the Fog of War*. Baltimore, MD: Johns Hopkins Univ. Press, 2001.
- [18] C. J. Petrie, T. A. Webster, and M. R. Cutkosky, “Using Pareto optimality to coordinate distributed agents,” *Artif. Intell. Eng., Des., Manuf.*, vol. 9, no. 4, pp. 269–281, 1995.
- [19] A. Rahman and H. Hexmoor, “Negotiation to improve role adoption in organizations,” in *Proc. Int. Conf. Artif. Intell.*, 2004, pp. 476–480.
- [20] T. B. Sheridan, *Telerobotics, Automation, and Human Supervisory Control*. Cambridge, MA: MIT Press, 1992.
- [21] T. B. Sheridan and W. L. Verplank, “Human and computer control of undersea teleoperators,” MIT Man-Mach. Lab., Cambridge, MA, Tech. Rep., 1978.
- [22] J. C. Strassner, *Policy-Based Network Management: Solution for the Next Generation*. San Francisco, CA: Morgan Kaufmann, 2004.
- [23] U.S. Air Force Scientific Advisory Board, *Air Force Operations in Urban Environments*, vol. 1, Washington, DC, DC: U.S. Air Force 2005.
- [24] X. Zhang and H. Hexmoor, “Utility based role exchange,” in *Proc. Central Eastern Eur. Multi-Agent Syst.*, 2001, pp. 332–340.
- [25] X. Zhang and H. Hexmoor, “Algorithms for utility-based role exchange,” in *Proc. 7th Int. Conf. Intell. Auton. Syst.*, 2002, pp. 396–403.



Gaurav Tuli was born in Mehsana, India, on December 12, 1983. He received the B.E. degree in computer engineering from the University of Rajasthan, Jaipur, India, in 2005 and the M.S. degree in computer science from Southern Illinois University, Carbondale, IL, in 2008. He is currently working toward the Ph.D. degree with the Computer Science Department, Virginia Polytechnic Institute and State University, Blacksburg.

He was a Graduate Research Assistant with the Computer Science Department and the Biochemistry and Molecular Biology Department, Southern Illinois University, from January 2007 to December 2007 and from June 2007 to October 2008, respectively. His research interests include multiagent system, artificial intelligence, ontology mediation, and bioinformatics.



Henry Hexmoor (M'04–SM'05) received the M.S. degree from Georgia Tech, Atlanta, and the Ph.D. degree in computer science from the State University of New York, Buffalo, in 1996.

He taught at the University of North Dakota before a stint at the University of Arkansas. He is currently an Assistant Professor with the Computer Science Department, Southern Illinois University, Carbondale, IL. He has published widely in artificial intelligence and multiagent systems. His research interests include multiagent systems, artificial intelligence, cognitive science, and mobile robotics.