

17th ICCRTS “Operationalizing C2 Agility”

Command and Control of Teams of Autonomous Units

Topics

Topic 8: Networks and Networking
Topic 4: Collaboration, Shared Awareness, and Decision Making
Topic 7: Architectures, Technologies, and Tools

Authors

Douglas S. Lange
Phillip Verbancsics
Robert Gutzwiller*
John Reeder**
Space and Naval Warfare Systems Center – Pacific
San Diego, CA 92152

*Robert Gutzwiller is a student intern from Colorado State University.
**John Reeder is a student intern from the University of Central Florida.

Point of Contact

Douglas S. Lange
Space and Naval Warfare Systems Center – Pacific
(619)553-6534
doug.lange@navy.mil

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 2012		2. REPORT TYPE		3. DATES COVERED 00-00-2012 to 00-00-2012	
4. TITLE AND SUBTITLE Command and Control of Teams of Autonomous Units				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) Space and Naval Warfare Systems Center, 53560 Hull Street, San Diego, CA, 92152-5001				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 Presented at the 17th International Command & Control Research & Technology Symposium (ICCRTS) held 19-21 June, 2012 in Fairfax, VA.					
14. ABSTRACT Command and Control (C2) has always been the practice of directing teams of autonomous units. These units have included individual soldiers, aircraft directed by a pilot, and ships maneuvered and fought by the combined intelligence of an entire crew. It is into this already populated universe that we are now working to add autonomous unmanned systems (AUS). In this paper, we explore the C2 of teams of autonomous units that include human, human populated, and uninhabited systems. Beyond the simple one-for-one substitution of a manned vehicle by an unmanned vehicle, we consider the reality that AUS share information differently among themselves and their inhabited counterparts. Similarly, humans and human populated units will provide information to their commanders in a different fashion than machines will, because each are uniquely capable of different observations and understandings. This paper also describes the sparse supervisory control that must be exercised over highly autonomous units, and considers what it means for a commander to supervise AUS that employ machine learning and cooperative autonomy.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 27	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

Command and Control of Teams of Autonomous Units

Douglas S. Lange
Phillip Verbancsics
Robert Gutzwiller
John Reeder

Space and Naval Warfare Systems Center – Pacific

Abstract. Command and Control (C2) has always been the practice of directing teams of autonomous units. These units have included individual soldiers, aircraft directed by a pilot, and ships maneuvered and fought by the combined intelligence of an entire crew. It is into this already populated universe that we are now working to add autonomous unmanned systems (AUS). In this paper, we explore the C2 of teams of autonomous units that include human, human populated, and uninhabited systems. Beyond the simple one-for-one substitution of a manned vehicle by an unmanned vehicle, we consider the reality that AUS share information differently among themselves and their inhabited counterparts. Similarly, humans and human populated units will provide information to their commanders in a different fashion than machines will, because each are uniquely capable of different observations and understandings. This paper also describes the sparse supervisory control that must be exercised over highly autonomous units, and considers what it means for a commander to supervise AUS that employ machine learning and cooperative autonomy.

The Boundary of Command and Control and Autonomy

Autonomous unmanned systems (AUS) development is based on the idea of having a system operate without continuous human intervention. This autonomy is not the same as avoiding all human direction. Human commanders decide when to deploy systems, when to end their deployment, and for those systems capable of making changes to their plans while deployed, commanders may change their tasks in some way. Teams of heterogeneous autonomous units will require command and control, just as they always have. That the units are now unmanned does not necessarily change that.

Controlling such a complex team requires several critical capabilities. First, the goals and constraints for the AUS team must be communicated to the various decision making nodes. These nodes may include all of the AUS, as they all may possess sufficient autonomous capability to decide how to act under many situations given the goals. A central controller, or more generally several distributed controllers, must have confidence (particularly if human operated) that the goals and constraints have been received and correctly interpreted by the autonomous units. Second, the control units must have sufficient situational awareness of the environment and the behaviors of the team members in order to decide if changes to orders are required. The control must have the ability to determine if any error conditions are

present and must be able to distinguish between aberrant behavior and what may be a plausible but unpredicted solution. Finally, team strategies must be selected to accomplish goals, and these strategies may need to be altered as the environment changes.

This paper explores the control of heterogeneous autonomous systems. In particular we look at the requirement for human controllers to influence the operation of these systems as well as addressing the need for autonomic control in situations where time constraints do not allow human decision making. This combination of requirements poses interesting demands on how AUS are constructed and how they incorporate adaptation.

Control of Complex Systems

The control exercised by a military commander over forces is described as “...guiding the operation” [1]. The presumption is that there is a mission statement, a set of assets with which to perform the mission, and an environment to operate in that may include an opposing force. According to Willard in [1], there are several ways in which a commander guides an operation.

- **Maintain alignment:** The commander must ensure that all decisions remain aligned with the operation’s mission and the commander’s intent.
- **Provide situational awareness:** The commander must assess the status of plan execution constantly, utilizing a common operational picture (COP).
- **Advance the plan:** The commander must monitor the

status of plan execution against the plan’s timeline.

- **Comply with procedure:** The commander oversees compliance with warfighting procedures to avoid mistakes (e.g., friendly fire engagements or collateral damage) and achieve efficiencies.
- **Counter the enemy:** The commander must be responsive to emerging intelligence, surveillance, and reconnaissance information that differ significantly from expectations.
- **Adjust apportionment:** Changes to asset availability or changes to requirements and priorities may require reapportionment of assets.

Military organizations are essentially complex cyber-physical systems. The end –nodes may be aircraft with pilots, or aircraft without pilots. These units whether manned or unmanned can be viewed as autonomous systems that cooperate to achieve a mission under the command of a human commander.

The tasks of the military commander are also clearly analogous to what would be required in many non-military systems. The only difference may be that no enemy exists, but the environment is nevertheless capable of surprising, therefore emerging information that can alter assumptions is still possible. Units may become inoperable just as in military operations, and adjustments to plans must often be made.

As the complexity of the system increases, the commander must work at higher, more abstract levels. The units of the system must also exhibit higher levels of autonomy so that decision

making is moved further down and is more immediate [2]. Based on the level of autonomy exhibited by the units, we can model the size of an operation that a single person can manage, provided the situation can be adequately described to the commander.

Sparse Supervisory Control

As we invert the ratios from many people working to manage a single unmanned air vehicle to a single person managing a team of heterogeneous AUS, we must move from direct control to supervisory control. In order to highlight the fact that with many vehicles to manage, the human must spend very little time on each, we call the mode we envision as sparse supervisory control.

Human Management of Automation.

Automation can easily be called ubiquitous. We interface with it daily, even if we do not immediately recognize it. In years past for example, elevators required human operators, but now we simply press a button to reach our floor. Even highly complex systems integrate automation; commercially flown airplanes have autopilots that are capable of landing the plane, and some models of cars have automated systems which bypass the driver if safety is in doubt (i.e., automatic braking systems, vehicle headway monitoring). The latest innovations in self-driving cars take autonomy even further. These systems integrate automation, but still rely heavily on a human component for routine performance and supervision. While automation is becoming more common, and more reliable, it rarely replaces or removes a human with experience from the overall task [3].

Automation has also been shown to result in phenomena such as complacency which results in operators failing to detect failures of automated systems [4,5,6] and automation bias, which results in operators blindly following automation recommendations or failing to act unless the automation requests the human action in decision making systems [7,8,9,10,11].

Leli and Filskov [12] suggest that it is specifically the integration that plays a large role in determining the effectiveness of a system outcome. In their work, automated diagnostic systems consistently outperformed clinicians when in isolation; however decision accuracy decreased as a direct result of integrated clinician application of the aid to diagnose psychological conditions. This result suggests that perhaps the most critical aspect of automation is not the engineering behind the automation itself, but the interaction between any automation and the operator who is expected to work together with it.

Parasuraman and Wickens [3] also identified issues related to the ability of human operators to understand the actions of the automation. The operators trust and ability to evaluate the performance of autonomous systems comes, in part, from an ability to recognize behaviors as correct or incorrect. AUS that have been programmed to perform in a particular fashion may or may not exhibit behaviors that are recognizably correct while optimal for the given situation. Knowing that such situations can exist may also push an operator to show complacency when observing odd behaviors because they can be explained as possibly correct if not humanlike.

It bears mention that the pitfalls of automation do not suggest that automation will always be harmful or should be abandoned. The dual nature of automation pitfalls and windfalls has begun to be explored within unmanned systems, with some success in improving situational awareness (SA), workload, and performance in the case of adaptive automation [13].

Higher levels of automation also have been shown to drive in part the number of vehicles a single operator is able to control, with higher levels of automation increasing that number [14, 15, 16]. How operators implement team strategies and switch between vehicles also appears to be influenced by the level of automation (Miller & Parasuraman, 2007; Parasuraman et al., 2005; Squire & Parasuraman, 2010)

Achieving Sparse Supervisory Control

Where we previously had one operator (or more) manually controlling only one AUS we plan to have one operator now oversee the workings of several semi or fully automated vehicles, changing the role to one of supervision and overseeing automated functions.

This role shift is where the control of automation merges back with command and control in terms of goals and where what has been learned about the management of automation must be applied to succeed in exercising control of AUS under command.

The commander must be able to alter the operation to counter decisions by the enemy. These changing needs mean that the tactics employable by a team of AUS must be flexible and robust. The environment will rarely look exactly as it

does to the programmers scripting behaviors. In this paper we describe initial experiments into the machine learning of team tactics. This approach is robust to changes in the composition of the team in ways that scripted actions are not.

Commanders must be able to communicate intent to AUS teams. There must be confidence that the teams understand this intent and must be able to recognize when decisions made by the team are out of alignment with that intent. In this paper we will discuss one approach we propose for aiding in that recognition.

The commander must maintain situational awareness of the entire environment being operated in by both AUS and autonomous human-occupied systems (AHOS). This SA must be done despite the mismatch of information coming from AUS and AHOS. Human occupied units will provide information differently than will machines directly. This information disparity will be a topic for another paper.

Learning Team Tactics

We present two ideas in developing team tactics by machine learning rather than programmer scripting. The first approach learns through the evolutionary development of artificial neural nets that take into account the geometry of the force and the opposition. By computing policy as a function of geometry, policies are developed that can adjust to changes in force composition and changes to threat tactics. The second idea employs computer observation of human behavior to develop robot actions that are recognizable to humans.

Developing Robust Tactics

A common approach to the design of autonomous systems is to design the entire system to be scripted. That is, a human decides on the action the autonomy takes for any given state the system is in. Such systems face many challenges. The first is the significant investment in human resources in the design because every part of the autonomy must be thought out and scripted. Furthermore, the autonomous capability is dependent upon the incorporated knowledge. Therefore, the investment of human resources necessarily includes subject matter experts in the task the autonomy is addressing. Another difficulty is that scripted autonomy is brittle. Unexpected situations can cause the autonomy to fail. This lack of robustness is due to the expansive state space that exists in the real world. Human designers will not be able to test or even anticipate every situation the autonomous system will be exposed to, resulting in a number of states that will not be addressed by the autonomy or be addressed with limited effectiveness. For example, consider a scripted autonomous system designed to deter piracy that assumes there exists only a single pirate threat at any given time. Once such autonomy encounters a situation where there exists more than one pirate, the script will degrade in effectiveness because the situation was not anticipated. Furthermore, once the autonomous system is introduced, pirates can adapt their tactics to counter

the system and render ineffective the specific design of the autonomy. Finally, scripted systems often lack scalability. In particular, the designs will be tied to a particular number of autonomous agents, or a particular autonomous system, meaning each time the number or types of unmanned vehicles change, the autonomy for the system must be redesigned.

In a proof of concept experiment, such a scripted system was compared to a cutting edge multiagent learning method developed at the University of Central Florida (UCF): Multiagent HyperNEAT. Multiagent HyperNEAT approaches the problem of multiagent learning by focusing on the geometric relationships among agent policies [17]. The policy geometry is the relationship among policies located at particular positions and the team behavior. Because multiagent HyperNEAT is built upon HyperNEAT, it can exploit the same patterns that HyperNEAT exploits, such as regularities. Furthermore, because HyperNEAT can encode repetition with variation, it can encode agent policies that share skills and vary in significant ways. Conventional multiagent learning cannot capture such regularities to enable sharing of skills and variation of policy. For a full description of how multiagent HyperNEAT encodes a team of policies see [17,18]. The main idea is to place a whole set of policies within a team geometry and compute their individual policies as a function of their location within the team.

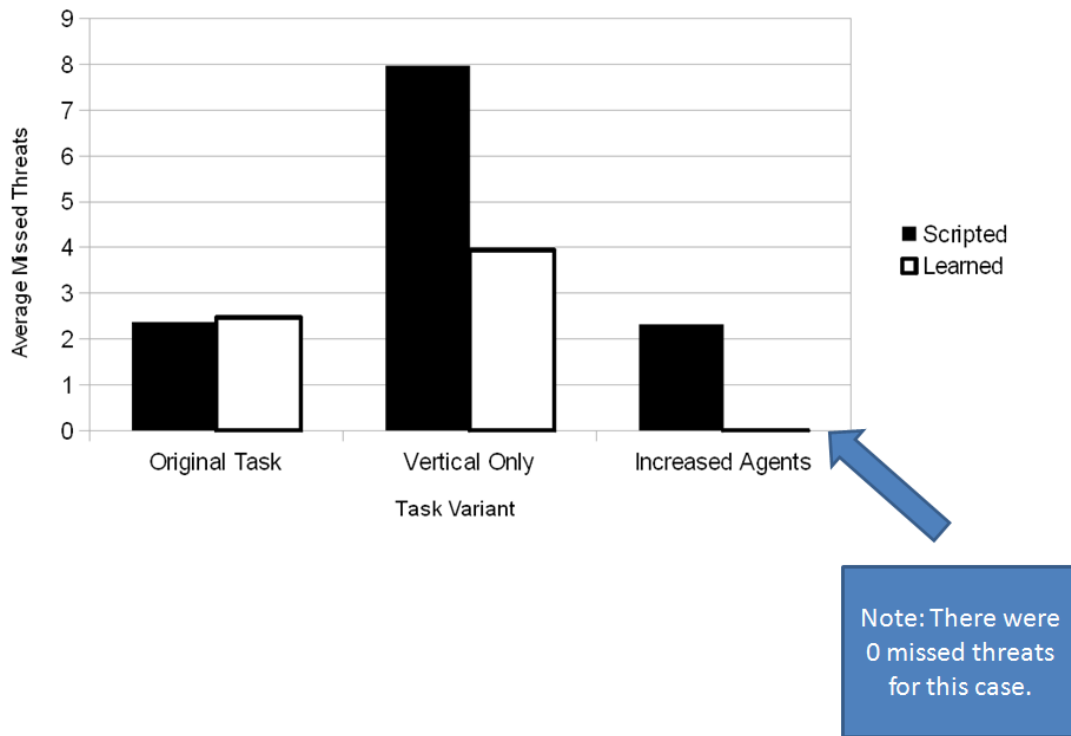


Figure 1 Performance of Scripted Search versus Learned Policy

Moving from simulated domains to real robots will mean that situations may occur where the number of agents varies, due to malfunctions or replacements and therefore ideally team size should be dynamically adjustable. The multiagent HyperNEAT approach allows such scaling because it represents team policies indirectly as a function of team geometry. Thus new agents can be added by simply generating the policy for their assigned team position.

The results of the proof of concept experiment are illustrated in the Figure 1. Overall, the results show that policies created by multiagent learning approaches are more robust to change. The scripted parallel search and learned multiagent HyperNEAT policies are compared on three variations of a threat detection task. In each of the variations, the policies are tested over 100

evaluations and the results averaged. The learned policy tested is trained solely on the first variation. The first variation is the training task for multiagent HyperNEAT, in which there are seven simulated unmanned vehicles patrolling and the threats can randomly appear along any of the four edges of the operational area. The AUS team is tasked with surveillance of the area and detecting all threats. The scripted version employs a classic *lawnmower* pattern to cover the area. In this task, the learned policy has statistically the same performance as the scripted policy, resulting in the patrols missing 2.37 threats, and the learned patrol policy missing 2.47.

In the second variation, the tactics employed by the threats are altered such that they now appear from two of the four sides at random, thus increasing the

density of the attacks along that vector and testing the robustness of the approaches. The learned policy missed 3.94 threats, significantly ($p < 0.001$) outperforming the scripted policy, which decreased significantly in performance to 7.98 threats missed.

In the third variation, threats can appear along all four edges, but the number of simulated unmanned vehicles in the team is increased from seven to eleven. The learned policy exploits the increased number of vehicles, decreasing the missed threats to 0. However, the scripted policy is unable to take advantage of the new vehicles and only insignificantly decreased missed threats to 2.32. These results demonstrate that learning can produce more robust and scalable policies.

Recognizing Correct Tactical Behavior

Learning team tactics through HyperNEAT will create more robust and scalable policies and behaviors. However, we must also be concerned with whether or not the human controller will recognize the behaviors as being safe and correct. As the HyperNEAT approach produces Artificial Neural Nets (ANN), we can only look at the team tactics as black boxes, and even within the proof of concept experiment, it took a fair amount of observation of the units to interpret (essentially guess) why they were behaving as they did. A human controller in such a system however, must be able to decide if the tactics being employed are aligned to the mission and whether or not they are properly countering the enemy or handling arising complications in the environment.

One of the primary draw backs to learning behaviors is that in the search for optimal actions the agents can behave in ways that seem foreign and unintelligible to the human operators. It is most likely the case, and something that should be tested, that agents that behave in a more humanlike fashion are more easily trusted by human observers. The development of humanlike agents is possible through hand coding and expert systems, but it is a tedious and complicated process. It is, however possible to learn humanlike behaviors through observation. FALCONET [19], also from UCF, is such a system designed to create high performance humanlike agents through human observation.

Humans learn through several different processes. Learning through observation entails watching the process as performed by some other individual or agent. Learning through experience involves repetitive practice of the process with feedback on performance. Learning can and does occur under each process individually, but it is the combination of observation and experience that generally leads to the highest levels of performance. For instance when learning a new sport humans typically observe others already proficient in the activity before beginning to practice themselves. Observation bootstraps the learning process of experience enabling faster learning speed and higher peak performance.

There is a long history in machine learning of borrowing from biological systems. Examples include knowledge representations like neural networks, optimization algorithms such as genetic algorithms and ant colony optimization,

and learning paradigms like reinforcement learning. In the particular method discussed below the observational-experiential learning cycle is replicated in machine learning to achieve the same goals for simulated learning that are achieved in biological learning.

FALCONET is a method of agent training that follows the biologically inspired cycle of observation and experiential learning. It was designed to enable the creation of high performing, humanlike agents for real time reactionary control systems [19]. Typically, the building of humanlike agents involves the complicated process of interviewing knowledge experts and then codifying that knowledge into a format that is machine readable. This process is complicated and time consuming and has led to the slow adoption of this technique in real world systems, despite the success that can be achieved. This problem is known as the “knowledge engineering bottleneck” [20]. FALCONET was designed to automate the agent creation process from human observation thereby sidestepping the bottleneck.

Previous work has been done with observational data alone to train agents, stopping once an acceptable level of performance is reached on training and validation sets [21,22,23,24,25]. While this approach might produce humanlike agents, it ignores the possibility that the observational agents will perform poorly in situations not covered in the observational data. The agents when presented with novel situations could perform in unpredictable and unintelligent ways. The experiential phase can fill in these gaps by providing

feedback on the agent’s performance in novel situations.

The training in FALCONET follows a two phase training approach. The initial phase is a supervised observational phase, followed by an unsupervised experiential phase. During the observational phase the objective of the learning is to be similar to the actions of a human trainer. Human trainers run through the selected tasks starting from many different scenarios to generate the observational training set. The agents are then trained on this data set while being graded on how closely they mimic the decisions of the human. In the experiential phase the agents are trained further against a particular measure of performance on the task. In FALCONET all training is done by a hybrid genetic algorithm (GA) particle swarm optimization (PSO) algorithm called PIDGION-alternate. This training algorithm is an ANN optimization technique that generates efficient ANN controls from simple environmental feedback. FALCONET has been tested showing that it can produce agents that perform as well or better than experiential training alone while incorporating humanlike behaviors. The results from FALCONET also state that unique human operator traits can be incorporated and be evident in the final highest performance controls, that is to say that agents sourced from different trainers have slightly different behavioral quirks.

As part of the validation of the FALCONET method, experiments were conducted with only the experiential learning phase. High performance controls were created in this manner, but they showed several “improper” quirks, that while more optimal in the

performance metric, seem foreign to human operators. These quirks, like driving backward or slamming the controls left and right very quickly, can be programmed out by a human designer, but it requires the a priori knowledge of all “improper” behaviors that would be undesirable. The FALCONET method bypasses this need by bootstrapping the process with human training.

Autonomic Control

So far, we have discussed the basic needs that will allow a human commander/controller to exercise command and control over a network of autonomous units that include highly autonomous unmanned systems (AUS). We have recognized that the controller must be able to develop adequate situational awareness of the environment, any enemies, and the behaviors and status of assets available. This SA requires the commander to recognize the behaviors being displayed as aligned with the mission, commander’s intent, and applicable procedures. It also requires that these behaviors be robust to the variation found in the environment.

We have also recognized that the human controller is subject to many difficulties inherent to managing automation. Humans are prone to complacency and automation bias. They also can only work at human speeds and can only handle a finite level of complex information. Abstraction and supervisory control are therefore essential to success if many rapid decisions will need to be made in controlling the network.

We are beginning to model AUS teams utilizing the Rainbow Framework [26] from Carnegie Mellon University. Rainbow groups commands into tactics and strategies and directs the system with those instead of individual actions. This approach allows an automated controller to move the system out of local maximums that it may encounter in utility functions. Additionally, the grouping of actions into tactics and strategies allows for the system to leverage learning techniques and previous human experience in dealing with situations.

Rainbow will provide an autonomic command and control in the sense that it assists with the same set of six tasks, only faster.

Maintain alignment: The mission goals and the commander’s intent will be modeled as a set of utility functions within Rainbow. Rainbow evaluates current readings from probes and gauges as well as tactics for changing the resource allocations against these utility functions to select an action.

Provide situational awareness: Rainbow’s framework of probes and gauges provides situational awareness into how well the current plan is meeting mission goals.

Advance the plan: The autonomic systems is continuously evaluating the readings from the probes and gauges against the plan and makes changes to adjust in the event that desired goals are not being met.

Comply with procedure: Procedural guidelines can be coded in the tactics employed by Rainbow in the *stitch*

language. It is our intention to also link tactical procedures to learned behaviors.

Counter the enemy: As the environment or enemy actions impinge on success of the goal, Rainbow adjusts the operation based on evaluating tactics against the likelihood of success. The evaluation processes will need to be robust enough to estimate how the changes will effect operations.

Adjust apportionment: The basic types of tactics employed in Rainbow to date have been apportionment decisions. In [22] experiments were done on video teleconferencing services where additional servers were brought online to solve problems that occurred during operations.

In the application of Rainbow, AUS teams are represented in a similar fashion as a network of servers would be. However, we include probes into the physical world providing information both on the AUS and on the environment. Many of these probes relate directly to the sensors that are onboard typical AUS. Strategy decisions involving costs include physical costs of fuel as well as risks found only in systems that interact with the physical world. Likewise, rewards are based on the ability of the AUS to effect a positive change on the environment, often in the form of achieving a probability of detection of other physical entities in a portion of the environment.

We have developed an initial proof of concept in autonomic control of unmanned systems by applying the Rainbow Framework to a simulated domain. In this domain, a number of AUS must maximize the probability of detection, $P(d)$, in an environment by

maximizing sensor coverage across the area. In this case, $P(d)$ is a simple metric defined as the fraction of horizontal and vertical paths across the space that do not have sensor coverage across them, i.e. straight paths that can be traversed without detection. Thus each AUS has a location (x,y-coordinate) and sensor range along with other parameters. Because AUS are conceptually similar to computational services (e.g. servers), they can be similarly modeled in the Acme architecture model language that defines a system architecture in the Rainbow Framework.

Properties, such as the geographic location and fuel state, represent values that are probed from the (simulated) world. Such values inform the Rainbow model manger, allowing it to accurately reflect and gauge the real system within the Rainbow defined model.

A key feature of Rainbow is the definition of constraints that the system must follow. For example, a web business may desire the minimization of response time for its customers and define a constraint that the response time experienced by any customer if below some threshold. In turn, Rainbow would probe these response time values from the real system and then evaluate the model for constraint violations. If a constraint is violated, Rainbow adapts the model through predefined strategies and then executes these strategies on the real system through effectors. In this proof of concept, the constraint is that the value $P(d)$ must be above a given threshold of 0.8. To satisfy this constraint, Rainbow implements a simple strategy. If an overlap in sensor coverage exists, move the active AUS to minimize the overlapping coverage. If

there is no overlap, but the constraint is still violated, then add a new AUS to the domain. This strategy implements two tactics that we described in *stitch: move* and *enlistAUS*.

In the *enlistAUS* tactic, the condition first checks whether the constraint is violated and then if there are any AUS not currently active. If both these conditions are satisfied, a random free AUS is chosen and activated. The effect being that the addition of the new AUS increases the sensor coverage, thus improving $P(d)$ over the current levels.

In the domain, the AUS exist in a two-dimensional plane with coordinate values in the range $[0,1]$. Initially, no AUS are active, thus the constraint is violated at the start. This compels Rainbow to adapt the system with the above strategies and tactics to achieve the pre-determined desired $P(d)$ level. When each AUS is activated, they are placed at location $(0,0)$ and then move from there. Each AUS moves at the same fixed speed of 0.01, have a sensor range of 0.1, and begin with 100% fuel.

Results demonstrate that Rainbow can be implemented to effectively control such systems. Figure 2 shows that the system begins at a low $P(d)$, indicative of the initial state of the system. However, by time step 150, Rainbow has successfully adapted the system to achieve the desired $P(d)$ value.

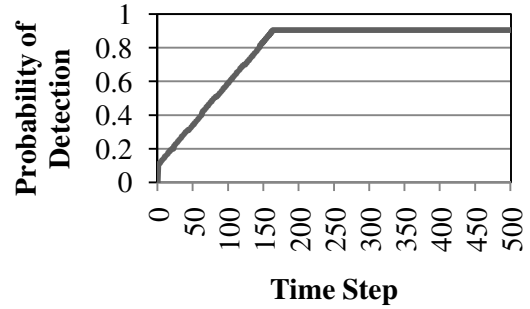


Figure 2 Probability of Detection in Proof-of-concept experiment

Not only is the result interesting, but the behavior of the system is as well. Figure 3 shows the final configuration of the system, when it achieves the $P(d)$ threshold required. Each circle represents an AUS sensor coverage.

Through the composition of simple strategies and tactics, organization emerges that effectively minimizes the probability of anyone passing through the region undetected. In this simulation, utility of the system is equal to the probability of detection. However, Rainbow includes the capability to calculate system utility as a function of multiple variables. For example, maintaining sensor coverage may be only one important aspect; another may be reducing fuel consumption or minimizing the AUS required. Rainbow weights each of these contributions of utility to determine the overall utility of the system. Through these relative weightings, different aspects can be emphasized. A suite of strategies to address these differing concerns may be required, forming a pareto-front of performance depending on particular user needs.

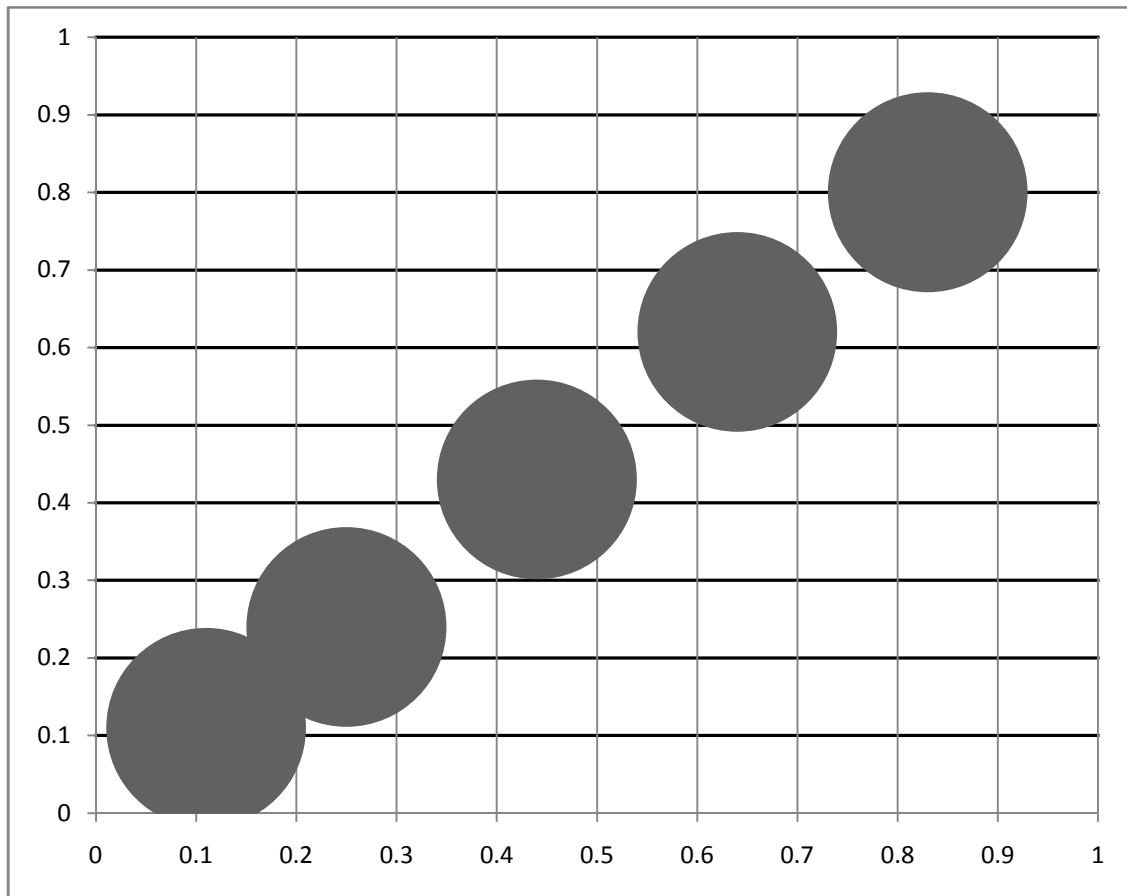


Figure 3 Final AUS Positions with Sensor Radius

The ability of Rainbow to automatically and quickly implement strategies frees up the controller to focus on macro level concerns, such as overall probability of detection, fuel levels, and costs, rather than micro-managing individual AUS. Thus the controller can make decisions about required detection levels versus preferred fuel levels and leave it up to Rainbow and its strategies to implement the decisions. Many avenues remain for exploration in the Rainbow Framework including, but not limited to, performance with “human-in-the-loop” changing the system constraints and goals, integrating machine learning into tactics and strategies, extending Rainbow to be able to dynamically acquire system architecture, and evaluating robustness to failures in the system, such as an AUS malfunctioning,

being destroyed, running out of fuel, or being reassigned.

Conclusions

We have both found and produced, proof-of-concept level experiments that demonstrate possible solutions to some of the challenges we perceive for the successful command and control of teams that include AUS. Our goal is to continue to pursue these possible solutions.

The command and control of teams requires that commanders be able to work at a suitable level of abstraction. Commanders must be able to recognize when changes to a plan are required and must have the ability to affect such a change. The dynamic nature of the

military environment indicates the need for robust adaptable capabilities for decision making in the individual AUS, but also the ability for their actions to be recognizable to human controllers. Autonomic capabilities are a likely approach to allow commanders to handle very large teams that may require rapid decision making, but the autonomic strategies must also be made more adaptable and in doing so also maintain the property of being recognizable by a commander.

References

1. Willard, R.F.: Rediscovering the Art of Command & Control. Proceedings of the US Naval Institute. (2002)
2. Rodas, M.O., Szatkowski, C.X., Veronda, M.C.: Modeling Operator Cognitive Capacity in Complex C2 Environments. 16th International Command and Control Research and Technology Symposium. (2011)
3. Parasuraman, R., Wickens, C.D.: Humans: Still vital after all these years of automation. *Human Factors*. Vol. 50. No. 3. 511-520. (2008)
4. Parasuraman, R., Molly, R., Singh, I.L.: Performance consequences of automation induced "complacency". *The International Journal of Aviation Psychology*. Vol. 3. No. 1. 1-23. (1993)
5. Wiener, E.L. Cockpit automation. In Wiener, E.L., Nagel, D.C. (Eds.): *Human factors in aviation*. Pp. 433-461. Academic. San Diego. (1988)
6. Parsasuraman, R., Manzey, D.H.: Complacency and Bias in Human Use of Automation: An Attentional Integration. *Human Factors*. Vol. 52. 381-410. (2010)
7. 32nd Army Air and Missile Defense Command: *Patriot Missile Defense Operations during Operation Iraqi Freedom*, Washington, DC. (2003)
8. Chen, T.L., Pritchett, A.R.: Development and evaluation of a cockpit decision-aid for emergency trajectory generation. *Journal of Aircraft*. Vol. 38. 935-943. (2001)
9. Johnson, K., Ren, L., Kuchar, J., Oman, C.: Interaction of automation and time pressure in a route replanning task. *International Conference on Human-Computer Interaction in Aeronautics*, 132-137. (2002)
10. Layton, C., Smith, P. J., McCoy, E.: Design of a cooperative problem-solving system for en-route flight planning: An empirical evaluation. *Human Factors*, Vol. 36, 94-119. (1994)
11. Mosier, K. L., Skitka, L. J., Dunbar, M., & McDonnell, L.: Aircrews and automation bias: The advantages of teamwork? *The International Journal of Aviation Psychology*, Vol. 11. No. 1, 1-14. (2001)
12. Leli, D., Filskov, S.: Clinical detection of intellectual deterioration associated with brain damage. *Journal of Clinical Psychology*. Vol. 40. No. 6. 1435-1441. (1984)
13. Parasuraman, R., Cosenzo, K. A., & De Visser, E. (2009). Adaptive automation for human supervision of multiple uninhabited vehicles: Effects on

- change detection, situation awareness, and mental workload. *Military Psychology*, 21(2), 270-297.
14. Cummings, M. L., & Guerlain, S. (2007). Developing operator capacity estimates for supervisory control of autonomous vehicles. *Human Factors*, 49(1), 1-15.
 15. Cummings, M. L., & Mitchell, P. J. (2005). Predicting controller capacity in supervisory control of multiple UAVs.
 16. Ruff, H. A., Narayanan, S., & Draper, M. H. (2002). Human interaction with levels of automation and decision-aid fidelity in the supervisory control of multiple simulated unmanned air vehicles. *Presence*, 11(4), 335-351
 17. D'Ambrosio, D.B., Stanley, K.O.: Generative encoding for multiagent learning. *Proceedings of the Genetic and Evolutionary Computation Conference*. (2008)
 18. D'Ambrosio, D.B., Lehman, J., Risi, S., Stanley, K.O.: Evolving policy geometry for scalable multiagent learning. *Proceedings of the Ninth International Conference on Autonomous Agents and Multiagent Systems*. (2010)
 19. Stein, G. FALCONET: Force-feedback approach for learning from coaching and observation using natural and experiential training. Ph.D. Thesis, University of Central Florida. (2009)
 20. Feigenbaum, E.A.: Knowledge Engineering: The Applied Side of Artificial Intelligence. *Annals of the New York Academy of Sciences*. 426 (1 Computer Culture: The Scientific, Intellectual, and Social Impact of the Computer). 91-107. (1984)
 21. Dejong, G., Mooney, R.: Explanation-based learning: An alternative view. *Machine Learning*. Vol. 1. No. 2. 145-176. (1986)
 22. Lee, S., Shimoji, S.: Machine acquisition of skills by neural networks. *International Joint Conference on Neural Networks*, IEEE. Vol II, pages 781-788. (1991)
 23. Sammut, C., Hurst, S., Kedzier, D., Michie, D.: Learning to fly. *Proceedings of the ninth International Workshop on Machine Learning*, pp. 385-393. (1992).
 24. Henninger, A.E., Gonzalez, A.J., Georgiopoulos, M., DeMara, R.F.: The limitations of static performance metrics for dynamic tasks learned through observation. *Ann Arbor*, 1001:43031. (2001)
 25. Fernlund, H.K.G., Gonzalez, A.J., Georgiopoulos, M., DeMara, R.F.: Learning tactical human behavior through observation of human performance. *Systems, Man, and Cybernetics, Part B: Cybernetics*, IEEE, Vol. 36. No. 1. 128-140, (2006)
 26. Garlan, D., Cheng, S., Huang, A., Schmerl, B., Steenkiste, P.: Rainbow: Architecture-Based Self Adaptation with Reusable Infrastructure. *Computer*. Vol. 37, No. 10, 46-54 (2004)



Command and Control of Teams of Autonomous Units

Douglas S. Lange, Phillip Verbancsics, Robert S. Gutzwiller, John Reeder

Presented to:
ICCRTS 17
Fairfax, VA
19 June 2012

Doug Lange
Research and Applied Science
C2 Experimentation and Technology

Autonomy and Command and Control (C2): Definitions for Our Purposes

Autonomy: Mobile robots, including autonomous vehicles, can be characterized by three tasks they perform – sense the environment around them; make a decision based on a predefined task and the environment it senses; and finally act in order to perform the predefined task by adapting to its environment [IEEE-Robotics 101].

- A mobile robot is comprised of sensors, autonomy algorithms and actuators
- Autonomy is therefore the decision making based on the task and the robot's model of the environment.
- Cooperative Autonomy is the ability of a group of AUS to collaboratively make task assignments and interpret and execute the intent of the system operator [Brizzolara 2011]



Autonomy and C2: Definitions for Our Purposes

“Command is the doctrinal assignment of authority”. One must possess a measure of command in order to exert control, which is defined as “...guiding the operation”. Control of forces can be described through the following contributions that a commander may make to an operation [ADM Willard 2002]

- ▼ Maintain alignment: The commander must ensure that all decisions remain aligned with the operation's mission and the commander's intent.
- ▼ Provide situational awareness: The commander must assess the status of plan execution constantly, utilizing a common operational picture (COP).
- ▼ Advance the plan: The commander must monitor the status of plan execution against the plan's timeline.
- ▼ Comply with procedure: The commander oversees compliance with warfighting procedures to avoid mistakes (e.g., blue-on-blue engagements or collateral damage) and achieve efficiencies.
- ▼ Counter the enemy: The commander must be responsive to emerging intelligence, surveillance, and reconnaissance information that differ significantly from expectations.
- ▼ Adjust apportionment: Changes to asset availability or changes to requirements and priorities may require reapportionment of assets.

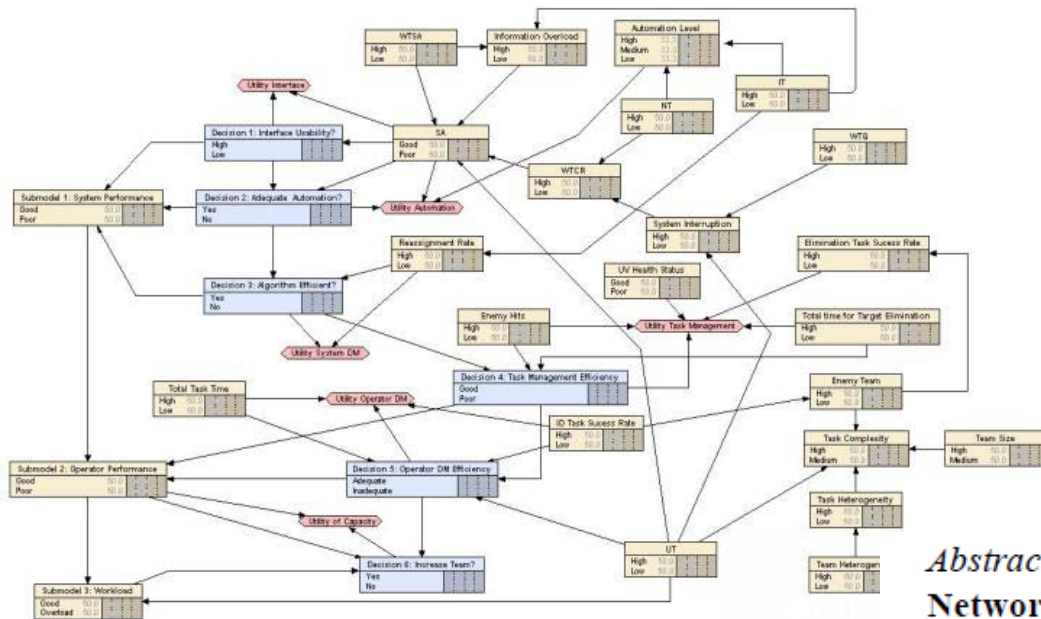
Requirements for C2 of AUS

- ▼ Communicate task and commander's intent to AUS. Commander must have confidence that AUS can accomplish mission and that all tasks and constraints are understood.
 - Vocabulary of tactics
 - Adapt to new tactics and constraints efficiently and effectively
 - Seamlessly adjust to team composition and geometry
- ▼ Maintain SA at appropriate level of abstraction
 - Ability to control multiple AUS in dynamic team arrangements
 - Recognize the difference between correct and aberrant behaviors relative to tasks and constraints provided. When must new tasking be given.
 - Recognize opportunities and requirements for changes to resource apportionment.

Sparse Supervisory Control

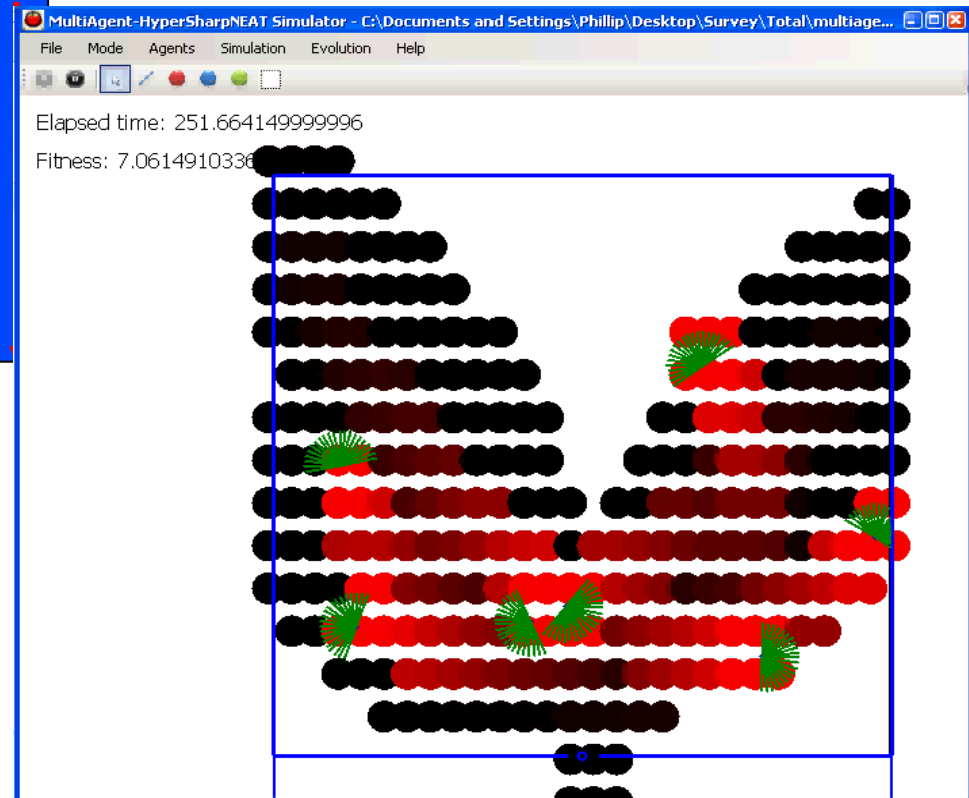
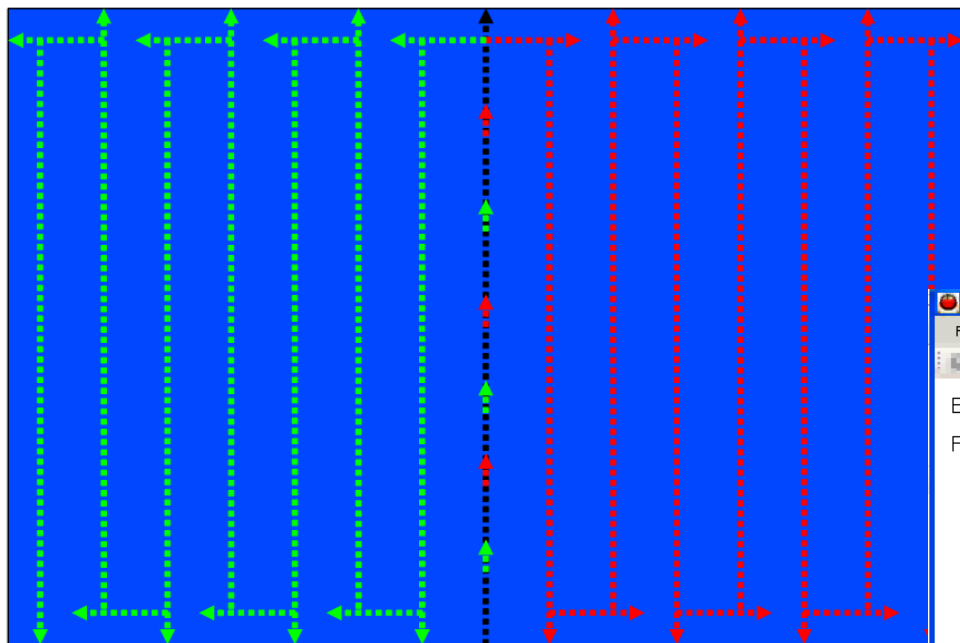
- 1 human operator controlling many (20-30+) AUS
→ 'inverting' the ratio
- Operator role now becoming *supervisory*:
 - issue orders & supervise: versus manually direct, navigate, survey, investigate, etc.
 - occasionally intervene: approve/disapprove of actions, change action/goals
 - In general; monitoring the actions of the AUS, stepping in when necessary
- Situation awareness important for both operator and AUS
- Human recognition of proper or aberrant behavior by the AUS

Situational Awareness and Operator Load

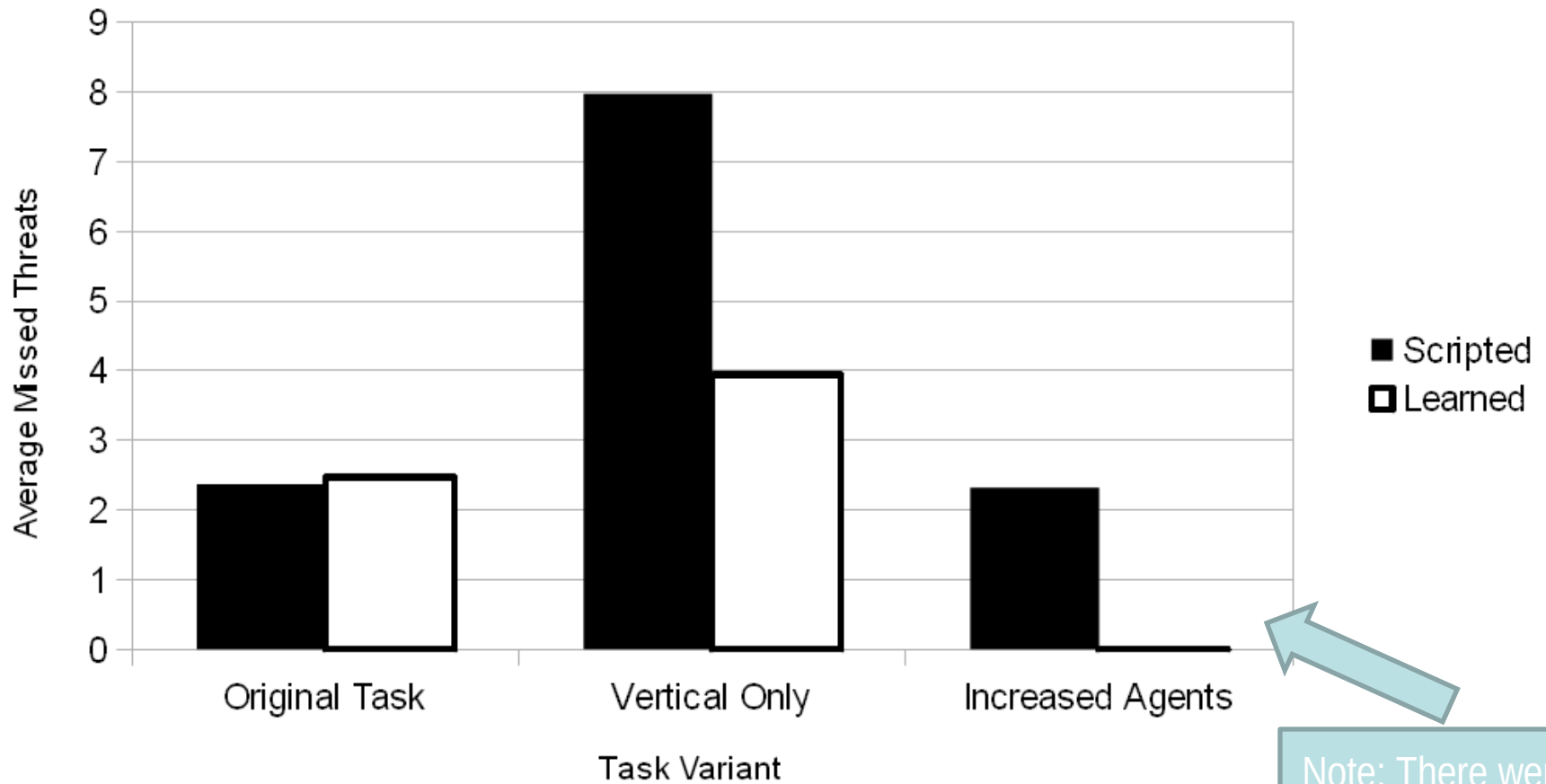


Abstract—The Department of Defense's future vision for Network Centric Operations (NCO) will increase combat power by networking relevant entities across the battlefield. This will result in highly complex mission scenarios in which the operator's workload will be easily overloaded if the system is not designed to support the mission requirements. New technologies for these complex command and control environments are currently being developed. Evaluating the adequacy of a particular technology for specific mission requirements is critical for military decision makers. This paper will introduce a new approach to model operator and system performance.

Scripted Tactics vs Multi-Agent HyperNEAT

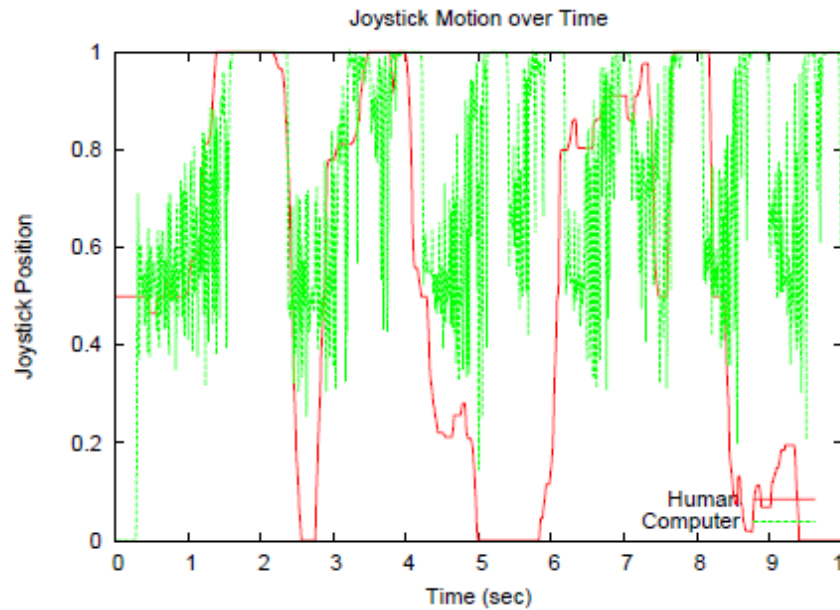


Experiments with HyperNEAT

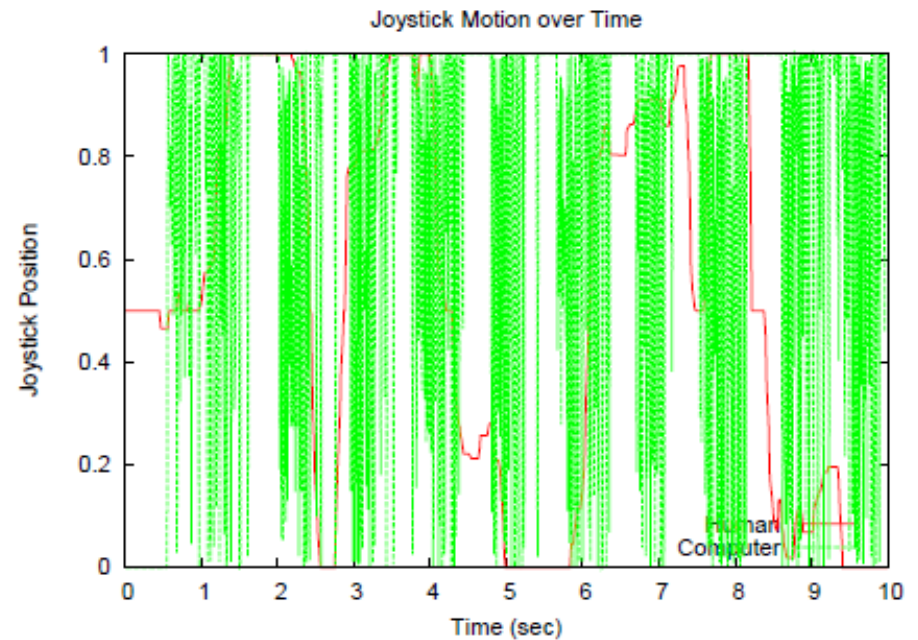


Note: There were 0 missed threats for this case.

Human Recognition of Behavior

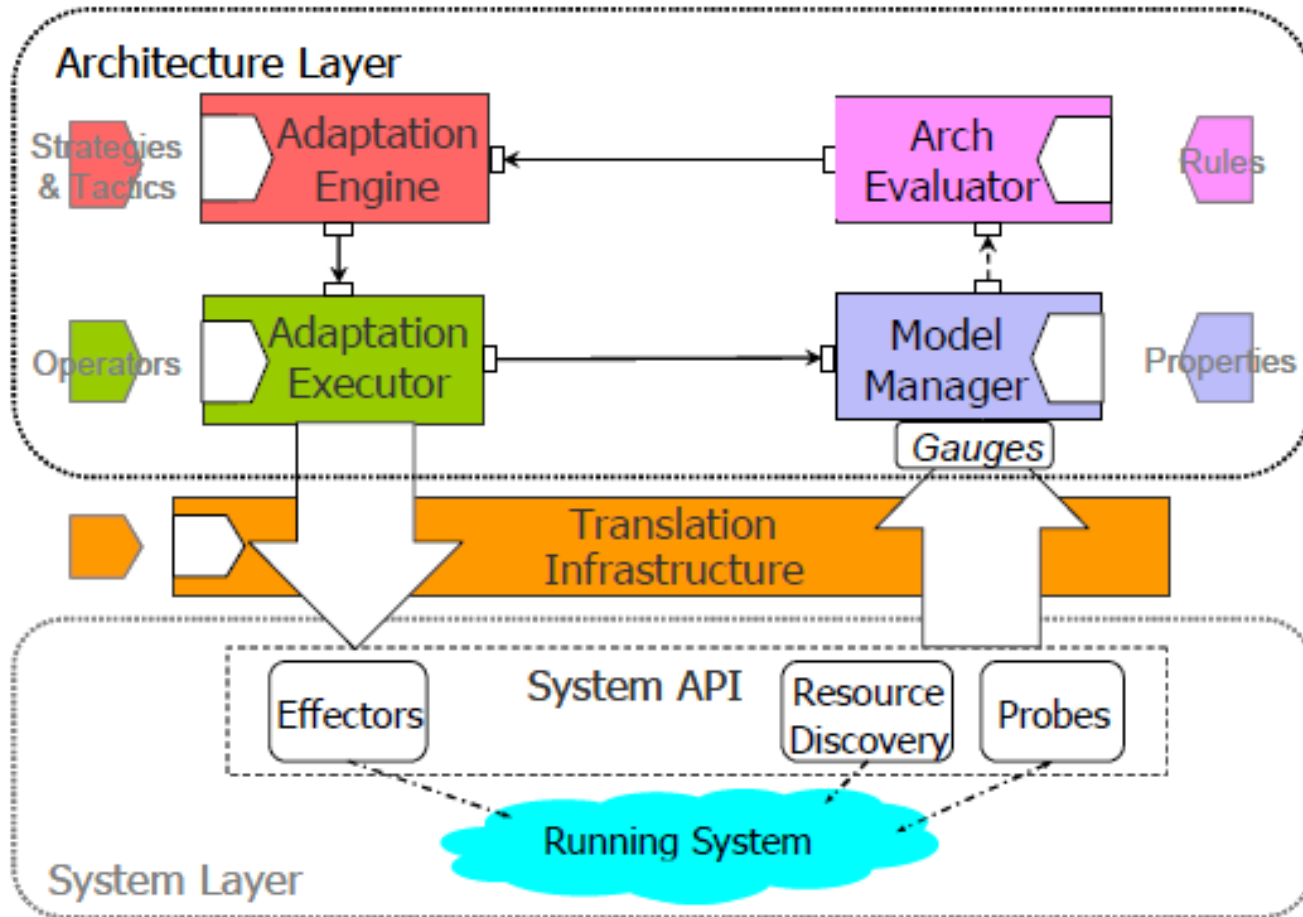


Observational + Experiential



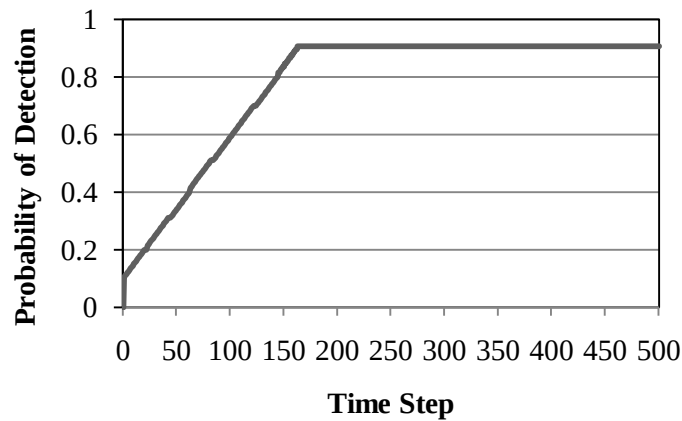
Experiential Only

AUS C2 Enhanced by Rainbow



Cheng, Garlan, and Schmerl, "Making Self-Adaptation an Engineering Reality", In Self-Star Properties in Complex Information Systems, Springer-Verlag, 2005

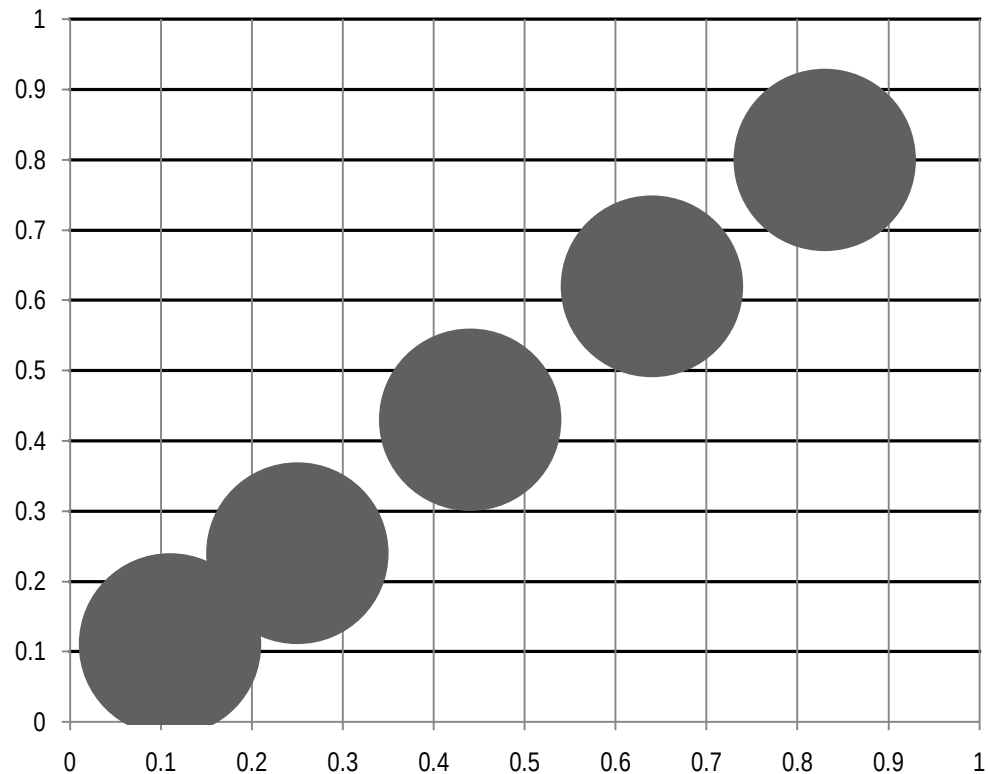
Emergent Behavior Using Rainbow



```

strategy BruteDetection
[styleApplies as cViolation] {
  t0: (overlapExists) -> move();
  t0a: (default) -> done;
}
t1: (cViolation as ! overlapExists) -> enlistAUS() {
  t1a: (overlapExists) -> move();
}

tactic enlistAUS () {
  condition {
    // Probability of detection is below a threshold
    ModelAlt.probabilityDetection(M.components) < M.MIN_PDTECT;
    // there should be enough available AUS
    ModelAlt.availableServices(T.AUST) >= 1;
  }
  action {
    set aus - Set.randomSubset(ModelAlt.findServices(T.AUST), 1);
    for (T.AUST freeAUS : aus) {
      S.activateAUS(freeAUS);
    }
  }
  effect {
    // Probability of detection rising should result
    ModelAlt.probabilityDetection(M.components) >= M.PDTECT;
  }
}
  
```



Implications for Cyber-Physical Systems

- ▼ Large Complex CPS result in “Human On-the-Loop” rather than “in-the-loop”.
 - Sparse Supervisory Control
 - Command and Control
- ▼ We are likely to require the use of machine learning.
 - Adaptation == under-specification
 - Trust comes from experimentation and observation. Is that enough in a safety critical application.
 - How will techniques of proving properties of composed systems of black-boxes work when the boxes adapt and affect each other.
- ▼ Observation of machine optimized policies.
 - AUS teams are composed and will exhibit cooperative autonomy
 - Commanders need to understand when to step in