

Comparative Analysis of Computer Generated Forces' Artificial Intelligence

Nacer Abdellaoui, Adrian Taylor, and Glen Parkinson

DRDC Ottawa and xplorenet

3701 Carling Avenue, Ottawa, Ontario Canada

Nacer.Abdellaoui@drdc-rddc.gc.ca, Adrian.Taylor@drdc-rddc.gc.ca, and gparkinson@xplorenet.com

ABSTRACT

Computer Generated Forces (CGFs) are a key component in constructive simulations and are being increasingly used to control multiple entities in Synthetic Environments (SEs). Being a cost-effective way to providing extra players in SEs, they are becoming a possible alternative in various activities, such as Concept, Development and Experimentation (CD&E), analysis, training, tactic development, and mission rehearsal. The predictable nature of many current CGFs behaviour is one of their biggest problems, making it easy for the trainee to distinguish between human-controlled and computer-controlled entities in the simulation environment. This can result in negative or ineffective training as the trainee quickly learns to predict the behaviour of the CGF entity and easily defeats it in a way that would not happen with a human opponent. This results in a requirement for humans to control synthetic entities, thus limiting simulation exercises by the availability of operators. If instead the Artificial Intelligence (AI) of these entities could be improved, the number of operators required will, thus, be reduced. The first step in such an effort is evaluating the AI capabilities commonly available in CGFs. Such an analysis was performed at the Defence Research & Development Canada (DRDC), revealing the common strengths and weaknesses of available CGFs, and suggesting which might be most useful as a platform for further AI research. This document presents the methods and results of this analysis.

1.0 INTRODUCTION

Modelling and Simulation (M&S) are extensively used in a wide range of military applications, from development, testing and acquisition of new systems and technologies, to operation analysis and provision of training and mission rehearsal for combat situations. In military areas, such as Exercise, Defence Planning and Support to Operations, and processes' effectiveness; the importance of M&S is steadily increasing.

M&S allows for exercises to be run with few real people involved with the remaining hundreds or even thousands of other battlefield entities being computer simulations. The possibility to deploy simulated entities possessing specific and distinct characteristics and behaviours whose parameters can be user-definable per any single entity, is of major importance as it contributes in solving some very significant problems inherent to the lack of personification of these simulated entities.

1.1 Background

In the last few years, significant advances have been made by the Computer Generated Forces (CGF) and Semi-Automated Forces (SAF) communities to make synthetic military environments more realistic. However, human reaction, adaptability and decision making in these environments are still far from being fully understood, and their modelling is still fairly simplistic. To overcome these limitations in current CGFs, synthetic entities are either controlled directly by a human or have their behaviour managed by a

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 OCT 2009	2. REPORT TYPE N/A	3. DATES COVERED -
4. TITLE AND SUBTITLE Comparative Analysis of Computer Generated Forces Artificial Intelligence		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) DRDC Ottawa and xplorenet 3701 Carling Avenue, Ottawa, Ontario Canada		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 See also ADA562563. RTO-MP-MSG-069 Current uses of M&S Covering Support to Operations, Human Behaviour Representation, Irregular Warfare, Defence against Terrorism and Coalition Tactical Force Integration (Utilisation actuelle M&S couvrant le soutien aux opérations, la représentation du comportement humain, la guerre asymétrique, la défense contre le terrorisme et l'intégration d'une force tactique de coalition). Proceedings of the NATO RTO Modelling and Simulation Group Symposium held in Brussels, Belgium on 15 and 16 October 2009., The original document contains color images.		
14. ABSTRACT Computer Generated Forces (CGFs) are a key component in constructive simulations and are being increasingly used to control multiple entities in Synthetic Environments (SEs). Being a cost-effective way to providing extra players in SEs, they are becoming a possible alternative in various activities, such as Concept, Development and Experimentation (CD&E), analysis, training, tactic development, and mission rehearsal. The predictable nature of many current CGFs behaviour is one of their biggest problems, making it easy for the trainee to distinguish between human-controlled and computer-controlled entities in the simulation environment. This can result in negative or ineffective training as the trainee quickly learns to predict the behaviour of the CGF entity and easily defeats it in a way that would not happen with a human opponent. This results in a requirement for humans to control synthetic entities, thus limiting simulation exercises by the availability of operators. If instead the Artificial Intelligence (AI) of these entities could be improved, the number of operators required will, thus, be reduced. The first step in such an effort is evaluating the AI capabilities commonly available in CGFs. Such an analysis was performed at the Defence Research & Development Canada (DRDC), revealing the common strengths and weaknesses of available CGFs, and suggesting which might be most useful as a platform for further AI research. This document presents the methods and results of this analysis.		
15. SUBJECT TERMS		

16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT SAR	18. NUMBER OF PAGES 12	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

human (e.g. correcting strange or incorrect behaviour). The human operator provides the knowledge and skills to ensure that synthetic entities perform in a realistic manner so the training can be effective or the experimental results valid. The number of critical entities is thus limited by the number of available operators. The lack of realism and full autonomy of synthetic entities thus limits CGF ability to replace human operators. Obviously, a more realistic AI modeling is needed: something that mimics human behaviour, including plausible mistakes and correct decisions.

Often, current CGF systems do not adequately model such complex human behaviour because their entities are governed by static scripts. By design they behave predictably, which makes them unable to respond to unexpected events. Also, scripting is time consuming and cumbersome way to capture possible behaviours. If something not anticipated by the script developer happens, the script may have no proper response[12]. Furthermore, scripted entities will not learn from or leverage previous experiences [3]. We seek an alternative system that reacts appropriately to the situation without having to anticipate and script everything at the design phase. Ideally, learning can be combined with this to develop new responses when something previously unseen happens, whether through offline training [9] or during the simulation [6].

1.2 Purpose

The goal of this research is to develop a reliable, realistic, and robust human behaviour modelling capability; and by ricochet, to reduce the staffing needed to operate and manage complex simulations by improving the autonomy and realism of synthetic entities' behaviour. This can be accomplished by improving artificial intelligence (AI) in computer generated forces (CGF). This will be achieved by developing an AI module that acts either as a federate in a distributed simulation, or as an integrated plug-in with selected CGFs. The AI module will take full control of the constructive entities, improving their behaviour and reducing human interventions [7].

2.0 COMPARATIVE ANALYSIS

The initial step in this process is the selection of tools that can be used to develop, and demonstrate a CGF AI module. This is accomplished by conducting an evaluation of existing CGF tools. There are two key questions to be answered by the comparison:

1. What are the AI-features missing from the considered products? This deficiency identification will tell us the status of synthetic entity's autonomy and realism. The end goal is the development of an AI module designed to address these gaps, where the missing capabilities will be met by academia, industry, and other fields such as games.
2. Which tools are convenient platforms for our AI research? Tools are required that can be compatible with an external AI component and with enough modularity to be able to add new features. Existing features should be leveraged to the full extent possible.

The analysis was based on a requirements wish list, which includes things such as realism, autonomy, and learning capability. Evaluation criteria were then elaborated based on these requirements. These criteria were used to classify the CGFs relative to each other as well as to the overall requirements list.

The tools selected for evaluation were based primarily on those CGFs that are used or can be made available to the project's client, the Canadian Forces Aerospace Warfare Centre (CFAWC). The evaluation candidate list contains Government Off-The-Shelf (GOTS) and Commercial Off-The-Shelf (COTS) simulation products as well as some serious games. The list was created starting with a catch-all list of possible CGFs and AI products found via web survey. The list was shortened in consultation with our client. The short list candidates are given in Table 1.

GOTS CGFs	Commercial CGFs	Serious Gaming
JSAF	MÄK Technologies - VR-Forces with Kynapse	BIA – Virtual Battlespace (VBS) 2
ONESAF		
XCITE	Presagis – STAGE Scenario with AI.Implant	Sonalyst Combat Simulations - Dangerous Waters

Table 1: List of Candidate Products

3.0 AI MODULE REQUIREMENTS

AI module requirements are based on addressing AI anomalies observed in simulations [3], as well as ideas for improving entity behaviour. The requirements are classified into five categories: autonomous operation, learning, organization, realism, and architectural requirements. The categories are broken into subsets of assessment criteria. Figure 1 shows a mindmap of the requirements.

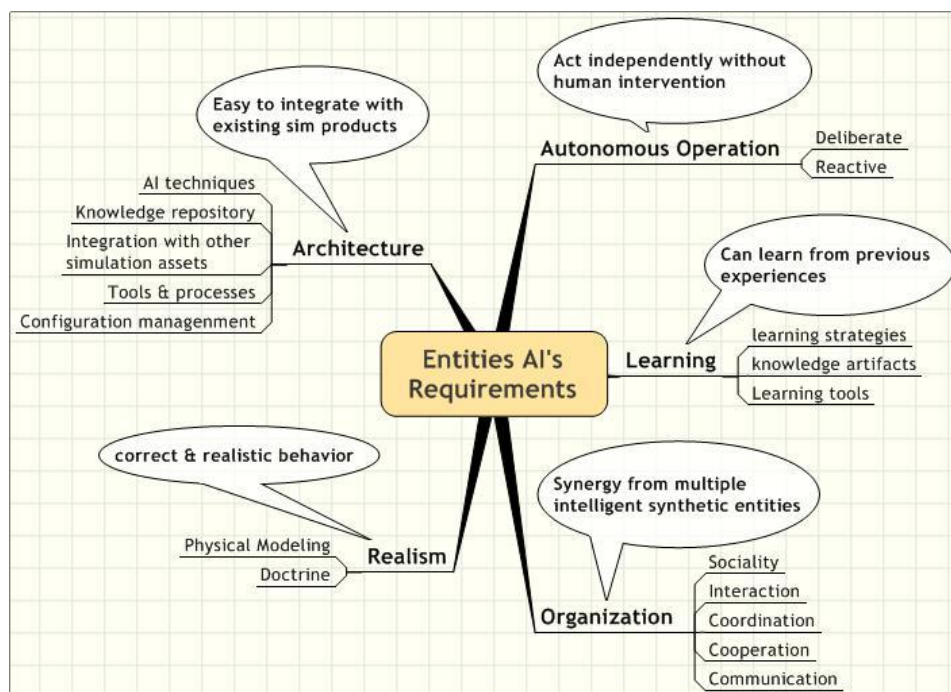


Figure 1: AI Module Requirements

3.1 Autonomy

Autonomy is the ability of a synthetic entity to act reasonably without a human's direct intervention. Increased autonomy reduces the need for human supervision and control. This is demonstrated by evidence of a "sense, think, do" loop. The used criteria are identified in accordance with the client's needs and seek to answer questions like:

- Does the AI have access to its own state, sensors, and effectors?
- Does it have decision rules, production rules, and a goal hierarchy (or more advanced features)? Can it behave unexpectedly? Can it perform target acquisition, fire weapons, and dispense countermeasures?
- Can AI predict expected actions of adversaries?
- Does the AI prevent obvious mistakes from happening, like ships grounding themselves?
- Are collisions handled realistically – do they model damage and velocity changes? Does this work when entities are of different types?
- Can the AI follow a planned route without human intervention? Does it avoid dynamic obstacles?

3.2 Learning and Adaptation

A learning process could make it easier to achieve correct AI behaviour without having to program it directly. In doctrine development scenarios, the entities could, learn and adapt to behave appropriately through human-directed training. For concept development and experimentation, the AI could discover effective behaviours on its own, possibly of interest to doctrine developers.

Learning can be accomplished in reflection after the scenario (e.g. processing event logs). The after-fact learning requires event recording, or monitoring tools for post-simulation analysis. Online learning requires specific capabilities to be built in to the AI.

The learning assessment is based on:

- availability of learning strategies
- automatic generation of doctrine based on experience and performance
- a good interoperable mechanism for import/export of AI knowledge (e.g. open format)
- knowledge artefacts organized in a non-proprietary database

3.3 Organization

In this survey, organization means a method of generating sought group behaviour. The organization is created from relationships between entities resulting in a team with qualities not found at the individual level. This could be demonstrated by the ability to give team-level orders that a group of units can carry out without further intervention. Organization also covers realistic behaviour of clutter targets [13].

Organization is rated by the following criteria:

- sociality, or the ability to communicate with multiple other entities
- the existence of a protocol for representing and transmitting information, goals, and decisions
- coordination, or the ability to perform an activity without conflicts with other entities
- cooperation, or the ability to work with other entities to achieve a common purpose together
- competition, or the ability to work towards a goal where its achievement implies the failure of other entities

- negotiation, or the ability of synthetic entities to reach an agreement about something
- allegiance alteration, or the ability to leave and join groups

3.4 Realism

Realism is a very subjective characteristic. In this context realism means that the autonomous entity behaves as if it is controlled by a human. Ideally a human playing against an entity with adequate realism wouldn't be able to tell whether their opponent was human or AI [14]. Realism can also be defined as behaving correctly given the situation; in other words, having a plausible doctrine [8]. Realism also covers the perception of information at the entity level – do sensors sense the SE or are they provided with “ground-truth”?

When looking at realism, we also consider the following:

- availability and variety of doctrine for CGF entities,
- ability to modify doctrine, and
- speed of decision making – is it similar to a human?

AI computational performance is also considered; it must be able to perform in real-time, which depends on the scenario's scope and power of computer(s) running the simulation. The approach taken was to gradually increase the number of entities, while running a simple scenario; whenever the simulation falls behind real-time, the present number of entities is the maximum for that specific CGF.

3.5 Architecture

Architecture is a broad category, covering the arrangement of the AI entities, external interfaces, support for different modes of operation, and technical support and documentation.

The architecture assessment consists of the following criteria:

- Built-in capability:
 - the availability of built-in AI models (e.g. finite state machines, neural networks, etc.) and the flexibility for adding new ones
 - the ability to modify entity behaviour, at building-time as well as at run-time
 - the diversity of entities' types (land, sea and air)
 - the modularly structured entity behaviours database and its interoperability with other formats for import and export
- External Interface: availability of a programming interface that allows entity control, simulation events' passage and sensor information transfer.
- Modes of operation: availability of recording and playback tools, configuration management, and support for human in loop and Monte Carlo simulations.
- Technical Support: mundane details like supported operating systems, quality of documentation and availability of technical support, ease of installation, and robustness of the software.

4.0 DESCRIPTION OF EVALUATION PROCESS

To assure a fair evaluation a consistent process was followed for each product. The process consisted of eight steps: product installation, documentation review, performance analysis (CPU load), scenario configuration, baseline scenario evaluation (without AI), AI configuration (using CGF tools), intelligent scenario evaluation, and summary of results. A common scenario was built in each tool as a basis for the evaluation, and scored criteria were developed based on the AI requirements.

4.1 Scenario description

The evaluation scenario was developed in consultation with CFAWC. The scenario's theme is smuggling detection and prosecution in a littoral environment. Smugglers using small aircraft, helicopters, cigarette boats, and/or unmanned vehicles transport contraband assets from offshore vessels into Canadian or US territory or across the Canada/US border. Using air, land, and marine assets, the friendly forces will detect and if possible prosecute smugglers. The scenario takes place off the south coast of Nova Scotia, encompassing the coastline from St. John, N.B. to Halifax, N.S.

The scenario was designed as such to highlight the problems identified in this research, such as unrealistic background traffic (creeping over ground of surface platforms, damage-less collisions, etc). Taking place in a peacetime environment means there is a high volume of air and marine traffic to provide the necessary volume for performance evaluation of a high number of entities. Also smugglers are highly adaptive and reactive, a combination well suited to the characteristics of AI we wish to investigate.

4.2 Scoring

The detailed evaluation criteria are derived from the AI requirements, and scores were awarded according to test methodology, assessment type, and priority weighting [7]. The test methodologies are demonstration, inspection, test, or analysis. The two assessment types are binary or subjective.

Binary assessment (Yes/No) is assigned to criteria requiring the existence of a capability. Two points are awarded for a "Yes", and zero points for a "No". An example of this type is the ability to support both deterministic and stochastic behavioural modelling. If scenarios can be configured both to produce repeatable results and accept a level of randomness that would allow variations in the outcome of the scenario, the assessment would then be "Yes".

Subjective assessment is used for requirements allowing variation in coverage. The subjective assessment allows for four degrees of compliance (abbreviated NDME):

- N for "Not Met" (CGF performance did not meet the requirement), 0 points
- D for "Deficient" (CGF performance was less than the requirement), 1 point
- M for "Met" (CGF performance met the requirement), 2 points
- E for "Exceeded" (CGF performance significantly exceeded the requirement), 3 points.

The basis for awarding an NDME score is specific to the evaluation criterion. For example, assessing available learning strategies will score as follows: N – no learning strategies are implemented, D – evidence of a partial implementation, M – if one learning strategy implemented, and E – if more than one learning strategy implemented.

Each scored criteria is multiplied by a weighting factor according to whether it was: Key (x3), Important (x2), or General (x1). The key criteria are defined as those most critical to the effective completion of this research. The idea is that at the end of the program all key criteria should be met by the selected tool or

tools. If they are not met by available products, they will be developed within this research project. It is worth noting that this approach does not preclude a different assessment according to a different weighting, for example for a different end-user's priorities.

Because categories have different numbers of requirements and different proportions of binary and subjective assessment criteria (which offer different numbers of points), this procedure led to categories having very different point totals. Consequently, the overall score is not necessarily representative of a product's overall standing and a meaningful product comparison must be performed across categories. To emphasize this, category scores are reported as a percentage of the total possible points. This focuses attention on the category score and avoids bias due to total point differences.

5.0 RESULTS

Table 2 shows each product's score per category, expressed as a percentage of the maximum possible score. The "Standard" score located adjacent to the category name shows the met threshold. This is the score for a hypothetical product earning a "Yes" for every binary criteria and "Met" for every subjective criteria. The overall product score represents the average score over all weighted requirements – in other words, the total points received divided by the total number of possible points. The red lettering of the Virtual BattleSpace (VBS2)[2] score indicates that it is based on an incomplete evaluation. Where evaluation criteria were incomplete, to ensure an unbiased evaluation, the missing criteria were given the average score of all other products. As a result VBS2 scores for those categories are not meaningful. Nevertheless, this approach makes it possible to include the partial evaluation about VBS2.

5.1 Common strengths and deficiencies across the products

The background colour of each table entry shows how close its score is to achieving the "met standard" score. Arbitrary divisions have been applied for the sake of visual aid, where white indicates that the product scored within 5% of the standard, light-gray that the product scored between 25% and 5% below the standard, and dark-gray that the product is more than 25% below the standard. The colours reveal some patterns of compliance, where nearly half of the squares are dark-gray (lower than 25% of the standard), a third are light-gray (lower than 5% of standard), and only eight of thirty-five are white (above 5% of standard).

CATEGORY	Standard	GOTS			Commercial		Serious Games	
		JSAF 2007	OneSAF	Xcite	VR-Forces /Kynapse	Stage / AI.Implant	Virtual Battlespace 2	Dangerous Waters
Autonomous Operations	85%	77%	71%	47%	82%	86%	77%	56%
Learning	67%	25%	33%	0%	33%	25%	33%	25%
Organization	62%	24%	55%	28%	55%	52%	24%	24%
Realism	67%	75%	83%	54%	74%	83%	58%	49%
Architecture	76%	60%	71%	32%	71%	63%	57%	51%
OVERALL PRODUCT SCORE	71%	61%	69%	38%	70%	70%	65%	47%

Table 2: Candidate Products Compliance

The top scores were COTS and GOTS, but serious games did not produce a single white score in any category. The commercial products had 5 out of the 8 white scores (commercial products are catching up on the legacy GOTS).

This analysis distinguished the met and unmet requirements. Figure 2 shows how many requirements were met by how many products. Of the 52 scored requirements, 47 were achievable to an acceptable level by at least one of the candidate products. In other words, 90% of the desired capability could be met by using a hypothetical integration of all the candidate products.

Nine requirements were adequately addressed across all products [7]. These are from Autonomous Operation, Realism, and Architecture categories. Met autonomy requirements included the use of logical rules to control entity behaviour, route following, autonomous piloting of own ship, and the ability to perform target acquisition, fire weapons, and dispense countermeasures. Met realism requirements included realistic physical and motion models. Met Architecture requirements included HLA and DIS interoperability and the ability to conduct adequate training. These capabilities are safely out of the realm of current research.

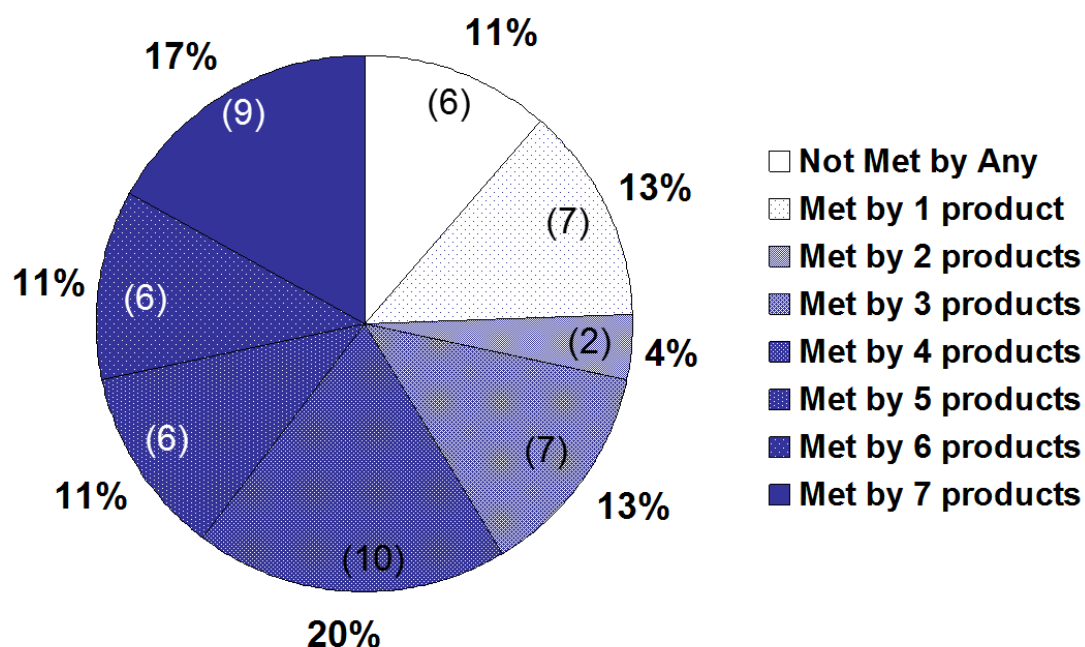


Figure 2: Depth of Requirements Coverage

Six requirements were not addressed by any candidates [7]. These belong to learning, organization, and autonomous operation categories. From the learning category, the learning strategies as well as the procedures were lacking coverage. Learning strategies are methods for learning, such as chunking or reinforcement learning, where learning procedures are processes for developing new behaviours. To meet the procedure's requirement there must be a well-defined or automated process for generating AI behaviour based on recorded performance. From Organization category, the deficient criteria were behaviours for competition and negotiation. Competition describes interactions where different entities seek the same goal, but one entity's success implies the failure of others. A fight is one example, but

competition also includes non-lethal interactions where the entities are aware of the competition and it affects their decision making. Negotiation is a process of cooperative decision-making between concerned parties regarding the resolution of a conflict. The goal of negotiation is to develop a settlement that is acceptable to both parties. A typical example would be a hostage-taking situation. From Autonomous Operation the deficient criteria were unexpectedness and initiative prediction [5]. Unexpectedness is defined as emergent behaviour that is not explicitly specified by doctrine [9]. This was not a rated criterion, but it was not observed in any of the CGFs. Initiative prediction is the ability of AI to assess the intent and expected actions of adversaries. As a result of this analysis initiative prediction, the learning, and the organization are the areas that will be considered for future research.

5.2 Specific product evaluation and platform selection

This section compares products within their product categories. Overall, COTS and GOTS categories did well, while serious games scored below expectation. This was generally because the scope and purpose of the GOTS and COTS products was better suited to our requirements than that of games.

Serious games are designed for a specific user domain, such as training vs. what-if gaming. Those we evaluated were aimed at a specific service (e.g. army or navy) instead of supporting all forces. They had limited or no Monte Carlo support and limited entity count for adequate performance. The commercial Dangerous Waters has a weak AI capability, limited external API, limited doctrine available, and lack of configurability, low entity count, limited entity database. Dangerous Waters doesn't support Monte Carlo simulation out of the box, but DRDC has previously commissioned a custom version in which this capability was added [10][11]. VBS2 has low entity count, limited sensor support, no Monte Carlo simulation, limited programming interface, and limited entity database. As mentioned above the evaluation of VBS2 was incomplete; however it was sufficient to exclude it from our consideration. These factors all hurt the serious games performance in our rated criteria.

The COTS candidates scored very well, each with no key requirement deficiencies. They are both integrated with professional standalone AI engines (Kynapse for VR-Forces, AI.Implant for STAGE). These AI engines focus on obstacle avoidance and path-finding, which made the scoring of their companion CGFs high in those areas. They both had good Monte Carlo simulation support. VR-Forces generally had an excellent AI capability built-in, supported by Kynapse/B-Have. It scored among highest for Architecture, with very good documentation and technical support, and support for data logger export to SQL, Matlab, and Excel. Stage scored high in realism and architecture, had excellent AI, good support, rich documentation, and appropriate external API. The difference in scores (about 10%) between Stage and VR-Forces is due solely to the weighting of obstacle avoidance and Monte Carlo simulation. VR-Forces with Kynapse and Stage Scenario were both fully compliant and are good candidates as platforms for follow-up development.

Other than Xcite [1], the GOTS products (JSAF [4], [17] and OneSAF [15], [16]) also did well. The version of Xcite available to DRDC did not have any AI capability, which led to low scores in most of our criteria. Other GOTS products are mainly focused on training and human-in-the-loop simulation. The AI in JSAF & OneSAF is based on built-in scripting as opposed to the external tools used by the COTS candidates. JSAF had just one key deficiency, in its complexity to create, manage, and modify entities. It also scored low on documentation. However, it has been used in the past with an external AI integration (e.g. Soar). OneSAF has exceptional entity AI, good realism and architecture, but poor documentation and support. OneSAF met the requirements, though its limited documentation means training would be helpful to maximize use of the product. Both OneSAF and JSAF are suitable for this research, with OneSAF coming out ahead.

The preceding analysis offers enough tools and equipment for recommending products that will fit this research's purpose. Basic requirements include a basic level of configurability, AI performance, and the



ability to run Monte Carlo simulations. This need eliminates DW and VBS2. Xcite is also not suitable because of the lack of AI capability. The remaining candidates are more or less evenly matched. Overall, Stage, VR-Forces, and OneSAF all scored within 5% of the “met standard” score, with VR-Forces having the best score. JSAF was within 10% of the met standard; its lower score was due entirely to minor architectural factors. Because of the tight scoring differences among the succeeding products, any of these products can be suitable for the rest of the research. As a result JSAF, OneSAF, VR-Forces, and Stage CGFs were all judged to be suitable as development platforms.

6.0 CONCLUSION

This study evaluated a list of candidate CGFs to measure their AI capabilities. One goal of this process was to identify capability deficiencies common to currently-available products. This study revealed that 90% of the capabilities sought were available across the candidate products, but the best fit product addressed 70% of those capabilities. Few requirements were fulfilled by none of the available candidates. The major gap was the absence of any learning process that can automatically generate behaviour based on experience. From an organizational perspective, there was no evidence for negotiation or competition between entities. From an autonomy viewpoint, no CGF has any sort of prediction capability (i.e. entities predicting others' intent) nor demonstrated emergent behaviour. These areas represent research directions for follow-up work.

The second goal was to evaluate each candidate independently and as a platform for this research. The serious games evaluated were not suitable for this research; however, all the Off-The-Shelf candidates were found to be suitable. VR-Forces, Stage Scenario, OneSAF, and JSAF are satisfactory, with VR-Forces as the overall winner.

In the next stage of work is the design of an AI module that addresses the missing requirements. It will interface with one or more of the selected CGFs. This evaluation helped identify the basic capability gaps that will form the basis of the AI module design. These are based on the desire for entities to operate autonomously in a synthetic environment, learn from experiences, participate as part of a larger organization, and perform realistically. Improving the autonomy and realism of synthetic entity behaviour will make the CGF supply those needed and rarely available wingmen, ground control, and other support personnel. By offering convincing synthetic entities we can reduce the level of staffing required for simulation-based training and concept development. Such simulations can offer an alternative to expensive live training exercises and provide opportunities for new concept development. By making this capability more available to our military forces, we will contribute to their success in current and future missions.

7.0 REFERENCES

- [1] Air Force Research Laboratory, Expert Common Immersive Theater Environment – Research and Development (XCITE^{R&D}) User's Manual Version 1.0, Mesa, Arizona, USA, April 2007.
- [2] Bohemia Interactive Australia, VBS2 VTK Application Scripting Interface (ASI) 1.0 Interface Control Document, 2008.
- [3] Fletcher, M., "A Cognitive Agent-based Approach to Varying Behaviours in Computer Generated Forces Systems to Model Scenarios like Coalitions", Proceedings of the IEEE Workshop on Distributed Intelligent Systems: Collective Intelligence and its Applications, 2006.
- [4] Hassaine, F., et al, Effectiveness of JSAF as an Open Architecture, Open Source Synthetic Environment in Defence Experimentation, Meeting Proceedings RTO-MP-MSG-045, Paper 11. Neuilly-sur-Seine, France, September 2006.
- [5] Laird, J.E., .It Knows What You.re Going To Do: Adding Anticipation to a Quakebot., AAAI 2000 Spring Symposium on Artificial Intelligence and Interactive Entertainment, March 2000.
- [6] Montemerlo, M. et al "Winning the DARPA Grand Challenge with an AI Robot", AAAI, 2006.
- [7] Parkinson, G., Abdellaoui, N. Scientific Authority, Artificial Intelligence (AI) in Computer Generated Forces (CGFs) Comparative Analysis- Summary Report, DRDC Ottawa Contract Report, October 2009.
- [8] Sandercock, J., Padgham, L., Zambetta, F., "Creating Adaptive and Individual Personalities in Many Characters Without Hand-Crafting Behaviors" Springer Berlin / Heidelberg, ISBN 978-3-540-37593-7,.2006,
- [9] Sandercock, J, Papasimeon, M., Heinze, C., "An Agent, a Bot and a CGF Walk Into a Bar..." , Proceedings of SimTecT 2004 Conference, Canberra, Australia, May 2004.
- [10] Sonalysts Combat Simulations, Dangerous Waters MALO Database Editor, date unknown.
- [11] Sonalysts Combat Simulation, Dangerous Waters MALO Monograph on NSE doctrine language, date unknown.
- [12] Spronck, P. et al., "Adaptive Game AI with Dynamic Scripting", Kluwer Academic Publishers, The Netherlands, 2005.
- [13] Sycara, K., Lewis, M., "Agent-based Approaches to Dynamic Team Simulation" Millington, TN 38055-1000 NPRST-TN-08-9, September 2008.
- [14] Turing, A.M., Computing and Machine Intelligence, Mind, 59:433-460, 1950.
- [15] US Government, OneSAF Users Manual and Help for OneSAF International V1.0a, 2008.
- [16] US Government, One Semi Automated Forces (OneSAF) International Installation.
- [17] US Government, Joint Semi-Automated Forces (JSAF) User Manual, 2007.

