

Effects of Combat Simulation Variance on Course of Action Development

**Ms. Janet O'May, Mr. Eric Heilman
Dr. Barry Bodt, Mr. Richard Kaste**

U.S. Army Research Laboratory
ATTN: AMSRL-CI-CT
Building 321
Aberdeen Proving Ground, MD 21005

Abstract

The U.S. Army Research Laboratory (ARL) is exploring the use of combat simulations in the development of military courses of action (COAs). We currently have the capacity to accept an automatically generated COA and produce a set of simulated results to compare using measures of effectiveness (MOEs). In a previous study, we evaluated a prototype COA generator by playing out its recommendations in the combat simulation Modular Semi-Automated Forces (ModSAF). Two large scenarios were played several times within ModSAF running on ARL's high performance computers to establish an empirical distribution of outcomes. An unexpected finding was the high variability, the so-called pure error variance, observed in the ModSAF results. Some variability is expected; but even after transformation to mitigate the instability of a ratio measure, the magnitude was surprising. In this paper we investigate sources of variability within ModSAF's successor, One Semi-Automated Forces (OneSAF), which provides flexibility to the user regarding simulated characteristics for units, terrain, weather, rules of engagement, and others. Changing input parameter settings introduces another form of variability expressed in signal effects formed from possible combinations of those parameters. The direction and magnitude of those effects is explored in consideration of pure error variance for representative cases.

1.0 Introduction

The U.S. Army Research Laboratory (ARL) is exploring the use of combat simulations in the creation of military courses of action (COAs). We have developed the research capability to produce a set of simulated results from automatically generated COAs and compare these using measures of effectiveness (MOEs). In a previous study, we evaluated a prototype COA generator by playing out its recommendations in the well-known combat simulation Modular Semi-Automated Forces (ModSAF).¹ Two high-entity-count scenarios were played several times, each with ModSAF running on ARL's high performance computers (HPC), to establish an empirical distribution of ModSAF

¹ Bodt, Barry, et. al., "OBJECTIVE FORCE COMMAND AND CONTROL: COURSE OF ACTION TOOL ANALYSIS", Proceedings of the 2000 US Army Science Conference, 2000.

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 2001		2. REPORT TYPE		3. DATES COVERED 00-00-2001 to 00-00-2001	
4. TITLE AND SUBTITLE Effects of Combat Simulation Variance on Course of Action Development				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) Arm Research Laboratory,ATTN: AMSRL-CI-CT,Bldg 321,Aberdeen Proving Ground,MD,21005				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 The original document contains color images.					
14. ABSTRACT					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT	18. NUMBER OF PAGES 14	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

outcomes. The MOEs for this study were based on the occupation of positional objectives and the conventional loss exchange ratio. An unexpected finding was the high variability, the so-called pure error variance, observed in the ModSAF results. Some variability is expected; but even after transformation to mitigate the instability of a ratio measure, the variability was surprising.

To better understand the viability of using combat simulation as a COA evaluation tool, we are studying the causes and implications of variances observed in our previous experimentation. The question of variability in stochastic simulation is of critical importance. Doctrinal adjustments and materiel decisions may hang in the balance if natural variation manifested in a small number of runs is confused with the effect of a factor in a study.

In the work discussed in this paper we investigate sources of variability within ModSAF's successor, One Semi-Automated Forces (OneSAF)², which affords great flexibility to the user regarding characteristics of units, terrain, weather, supplies, and rules of engagement, among others. Changing the level of individual input parameters introduces another form of variability expressed in terms of signal effects formed from possible combinations of those parameters. The direction and magnitude of those effects is explored in consideration of pure error variance for representative cases.

2.0 Challenge

For centuries, man has studied the art of war to better understand the implications of battlefield maneuver. The results have manifested themselves in the form of devices, processes, doctrine, models, simulations and games. While useful in determining the tactics employed in some situations, these methods have fallen short of expectations; a method has not been developed to reliably predict combat results. A number of reasons explain this, but the most prevalent of these is the inclusion of man as a variable in the equation. Man can only in the general sense predict what man might do in any given specific situation.

The key to models, simulation, and gaming is abstraction. Abstraction enables us to gain general knowledge from inexact or generic formats. As we have progressed in the study of war, we have developed a thirst for a more exact depiction of warfare. Great efforts have been made to incorporate actual tactics and techniques, weapon physics, and human behaviors into current simulations. However, the cost for improved combat realism is increased overhead and complexity of simulation operation.

The more realistic a combat simulation, the more computation is required to offset the complexity. In reality, combat is an extremely complex interaction that involves many entities, each with independent freedom of action. While the encapsulation of reality within simulation has progressed, combat remains as an abstract concept. The

² OneSAF Testbed Baseline, Version 1.0, developed for U.S. Army Simulation, Training and Instrumentation Command (STRICOM) by Science Applications International Corporation and Lockheed Martin Information Systems Company.

adjudication of combat abstractions has progressed from absolute rules, to human judges, to physical model representations, to sophisticated computer codes. Yet throughout all of this development, our best wargames and simulations today still rely on abstraction, human intervention and subject matter experts.

On a compared scale of playability and realism, the general rule of thumb is an increase in realism, and thus the resources necessary for simulation operation, results in a decrease in playability. While military science has progressed toward the ability to better simulate combat, few studies attempt to specify that progress experimentally. We are attempting to fill this knowledge gap.

3.0 Objectives

The goal of this research is to apply the principles of scientific examination to a modern, sophisticated computerized wargame in an attempt to gage its effectiveness in the prediction of combat. Specifically, we have created a statistically based experimental design to examine the variance in combat results produced by changing several different user-controlled parameters within the OneSAF simulation. Understanding this variance will enable us to better incorporate OneSAF into our ongoing efforts to create a COA generation-evaluation suite for the future U.S. Army.

The experiment described herein is the culmination of Phase One in our COA Technology Integration (COATI) project. We are attempting to create a system that accepts COAs generated from any source and evaluates those COAs using the combination of combat simulation and statistical methodologies.³

We will use the outcome of this experiment to assist us in simulation choice for inclusion in a prototype system to be created in COATI's Phase Two. This prototype is expected to operate on a PC and contain COA generation ability, combat simulation software, and statistical analysis capability. A well understood combat simulation outcome interval would enable a commander to better judge the applicability of COAs to an actual combat situation.

4.0 Experimentation

4.1 Experimental Design

The experimental design in this study was a 2⁴ factorial design in the four factors: competency, formation, movement, and strength. We included average (A) and expert (E) forces to represent the two levels of this factor with respect to competency. With respect to formation, we chose column (C) and wedge (W). Movement consisted of march (M) or traveling overwatch (T) and strength played as even (1:1) [E] or uneven (3:1) [U] with advantage to Blue. Nine replications supported each of the sixteen treatment

³ For an explanation of COATI Phase One, readers are directed to: Bodt, Barry, et al., "An Experimental Testbed for Battle Planning", Proceedings of the 2000 Command and Control Research and Technology Symposium, 2000.

combinations, yielding 144 runs in all. We repeated this basic design for a host of potential responses collected simultaneously during each run.

Responses for Red and Blue forces included measures of fuel use, ammunition use (combined and specific munitions), whether or not the objective was taken, and remaining force strength. Remaining force strength considered the remaining function of each vehicle using the damage assessment categories of kill, mobility and firepower kill, firepower kill, mobility kill, and no damage. Each vehicle, examined in its end state, received a score in accordance with a commercial combat simulation game point scoring metric⁴ (Avalon Hill). The scoring metric was created by subject matter experts, employed by an independent source, to design a realistic and non-arbitrary force scoring system. Fuel use is expressed as a record of the percentage of fuel remaining at battle's end. We expressed this value as the average percentage of remaining fuel among mobile vehicles. The percentage of remaining fuel with respect to all Blue or Red units in battle was also recorded. Similarly, total ammunition and specific munitions were measured for the percent remaining with respect to all vehicle stockpiles at the beginning of battle and with respect to only the vehicles with firepower at the end of the battle. In addition to these responses, we saved the damage end-state of each vehicle in each run to provide a drill-down capability into the data for each of the 144 runs. Logistical characteristics with respect to fuel and ammunition remaining were also saved for each vehicle.

4.2 Scenario Development

The experimental scenario set was created using several different criteria including vehicles and unit formations supported by OneSAF, force equivalence as based on the Avalon Hill system, and limited time available for scenario execution. The set consisted of two basic scenario types predicated on the commercial point system. The first scenario represented an even fight at a 1:1 numerical strength ratio and the second represented an uneven fight at a 3:1 combat strength. Both scenario types featured a friendly force attack of an in-place enemy force. A diagram of the 1:1 scenario is shown in Figure 1.

Commanders cross attach infantry and armor in combat to take advantage of combined arms doctrine; therefore, we founded the scenario set on a friendly company-sized taskforce supported within OneSAF consisted of two tank platoons and one mechanized infantry platoon. In the even strength scenario, the threat force equivalent to the chosen friendly taskforce consists of a tank company and an infantry company. The single taskforce was easily augmented to create the second scenario force structure by replicating the friendly force twice more to create a ratio of 3:1.

⁴ The point system is derived from the MBT Game Rules printed by the Avalon Hill Game Company, USA, copyright 1989.

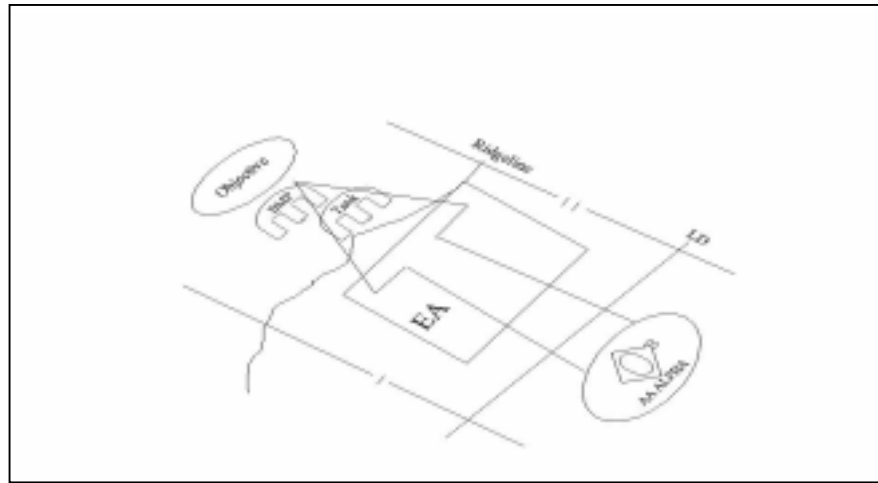


Figure 1: Scenario diagram for the even strength scenario.

4.3 Execution

The 144 scenario executions were completed over the course of nine days. The computers utilized included HPC systems such as SGI Origins and Sun 1000s, and smaller systems consisting of an SGI Onyx and an Sun UltraSparc 60. The average completion time of each scenario was thirty to forty-five minutes. Controls were in place for a data collector to reserve one of the sixteen specific experimental cell types while running the scenario. This system prevented confused data during intermediate collection points due to system data file overlaps. Each user kept a log containing the scenario type, date of execution, file name for the data, system of execution, user name, and information on the status of the objective. At the end of the 144 runs, the individual logs were merged.

5.0 Data Collection

5.1 Collection Mechanisms

In previous experiments, we collected the data by hand. This method was error-prone and time consuming. In order to allow more computer scenario executions, we recognized the need for automatic data collection. Our search yielded two distinct OneSAF data collection systems, neither fully satisfying our data collection and network environment requirements. Creating an in-house capability for data collection became our focus. In future work we will be able to capture killer information, providing the type of ammunition and the originating vehicle whenever an entity is destroyed.

The OneSAF release provides source code, enabling us to modify OneSAF to meet our needs. The simulation provides a *Unit Status* function that allows the user to select vehicles and obtain user-specified status information for the selected vehicle or unit. We enhanced this capability by writing out a data file for every vehicle in the simulation whenever the status function is enabled. In addition to the standard status information, we

also included each vehicle's unique object identification and position. Our team observed each scenario execution from start to finish in order to collect data. A scenario was deemed ended when all surviving Blue entities achieved the mission objective by arriving within the designated terrain area. For each scenario run, we cycled the status capability prior to starting the simulation and at the end of the simulated battle. The status capability remained functional throughout the execution. The data file produced by this sequence is uniquely named for each run by using the system clock time.

5.2 Post Processing

After the 144 runs were completed, the data was analyzed and tabulated. We wrote seven programs to perform these tasks. The first task was to create individual data files for each iteration of the status function. Data for vehicle status, fuel and each type of ammunition was then tabulated for every vehicle at the beginning and the end of each run. The next step was to create "rollup" files. One large rollup file contained a record for each scenario run and data consisting of competency level, vehicle formation, mode of travel, strength, percentage of Blue and Red fuel remaining, percentage of each ammunition type remaining, and total vehicle final score for each side. Rollup files for each run contained specific data for each vehicle. The data files were then transformed into spreadsheets for easy portability into statistical tools.

6.0 Statistical Analysis

The data collected in this study offer many intriguing relationships, but we restrict our attention to a small subset for this presentation. We treat the final strength for Blue forces and their TOWs remaining, the latter with respect only to vehicles with firepower at the end of the battle. Responses are considered individually rather than in a multivariate context. We make observations on experimental design effects, the role of variance, and their tie to the battlefield as represented by OneSAF. Analytical detail of our results is only exploratory at this juncture. We eschew formal statistical inference until a more rigorous analysis can be performed and presented in a more extensive report to follow.

Figure 2 summarizes all 144 runs with regard to final strength of the friendly, or Blue, forces. The y-axis is the final Blue strength percentage, with the x-axis representing the run numbers. The color key and centerline combine to distinguish among the 16 experimental conditions. Within a level of competency, a common color for clustered points represents identical experimental conditions. For example, red clustered points share the experimental conditions (column formation, march movement, uneven strength) and are denoted (A)CMU to the left and (X)CMU to the right.

A few observations can be made easily from this figure. Force strength effects are seen clearly by viewing, from left to right, successive pairs of clusters differing only in strength. Each pair shows an increase in Blue final strength for even initial strength relative to uneven initial strength (advantage Blue) conditions. Competency, on the other hand, does not appear to have a strong effect for most conditions. To see this, compare like-colored clusters (e.g., ACME and XCME). Two exceptions, with TE in common,

suggest average competency forces have greater difficulty in executing traveling overwatch in an even strength conflict. Formation also fails to show a strong effect. Within-cluster variability is influenced by experimental condition. Runs with even strength generally show greater variability than their uneven strength counterparts (e.g., ACTE and ACTU).

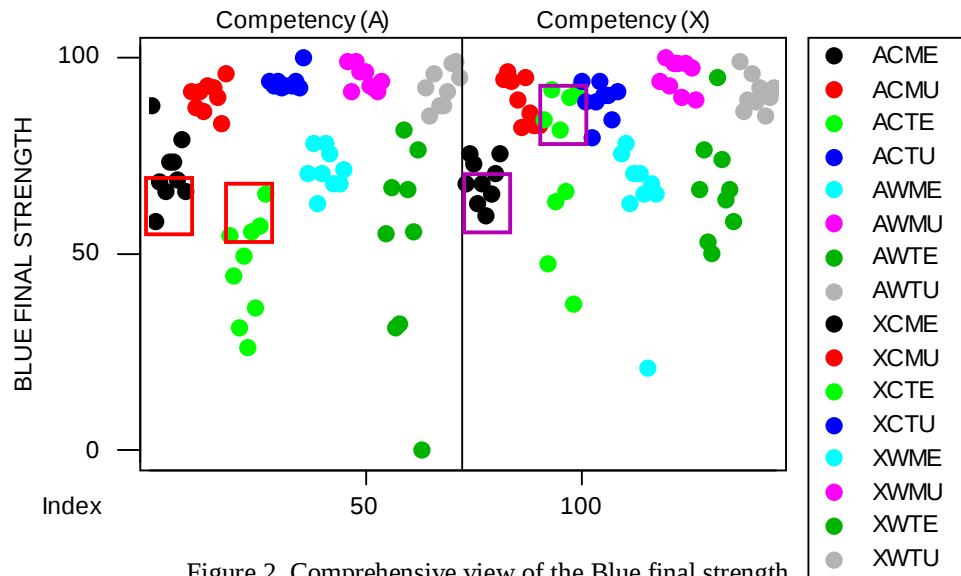


Figure 2. Comprehensive view of the Blue final strength.

On the battlefield, the rule of thumb is to attack with a 3:1 advantage or better. The data shown here support this rule. The traveling overwatch movement technique is considered prudent for advancing into battle. The surprising amount of casualties within OneSAF reflects the exposure of weaker vehicle flank armor to threat fire. Expert crews were able to maneuver their vehicles more efficiently and thus reduce the amount of time their flanks were turned toward the threat. As the data show, when traveling overwatch is performed expediently, the results are fewer casualties. However, the high casualty rates even in scenarios with expert crews suggest the maneuver should not be considered for direct fire combat.

There are potential consequences of differing variability for the investigators who use OneSAF for training or development exercises. If the experimental design fails to recognize, with increased sampling, the variability present, too few samples can lead to true effects missed and false effects claimed. For example, if the boxed subsets of points in Figure 2 had been observed with respect to movement under column formation and even strength, we would find no difference between march and traveling overwatch for average competency forces, but would find a difference for expert forces -- exactly the opposite of what the whole data show. This illustrates the importance of taking adequate samples to reveal signals amidst the noise.

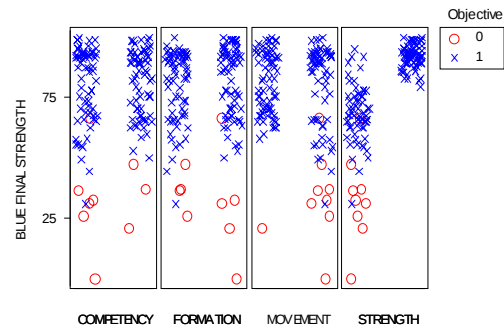


Figure 3. Main effects plot for blue final strength.

Figure 3 shows Blue final strength on the y-axis plotted in separate but adjacent graphs with x-axes given, respectively, as competency [A(left), X(right)], formation (CW), movement (MT), and strength (EU). Color is used to identify whether the objective was taken (blue) or was not taken (red).

From this graph the general claims hold regarding the lack of a competency or formation effect. Strength is perhaps more clearly seen as a strong effect, and we note the failures to achieve the objective occur only when the initial strengths were even. This figure brings to light new information not easily seen in Figure 2, namely that traveling overwatch movement resulted in lower Blue end strengths, owing primarily to the fact that eight of nine failures to achieve the objective occurred under this condition.

Since the best form of attack presents the least vulnerable portion of the vehicle to threat fire, the march maneuver that we chose as the lesser valued condition for attack turned out to be better in the even fight. By not turning vehicles to positions required for traveling overwatch, the march formations proved the value of shock by closing faster on the enemy and surviving long enough behind their thicker armor to destroy defending units.

Figure 4 reveals the disparity of within-condition variances. The horizontal lines on the left half of the figure represent 95% confidence intervals for the standard deviation based on the nine observations taken under each experimental condition. Units on the x-axis are the percentage of Blue strength remaining. On the right half of the graph, factor levels are identified using a (0,1) nomenclature for the levels of the factors. This nomenclature applied to Figure 3 would show level 0 (left) and level 1 (right).

From this graph we see several cells with large variation. These cells share the even strength condition, corroborating the interpretation of Figure 2, and the largest of these have traveling overwatch in common. The graph suggests changing the experimental conditions may have as much or more impact on distribution variation than distribution center. In addition to the implications of variance mentioned previously, attention to variation is a requirement when we move beyond data exploration to formal analysis of variance (ANOVA). ANOVA requires homogeneity of variance across cells.

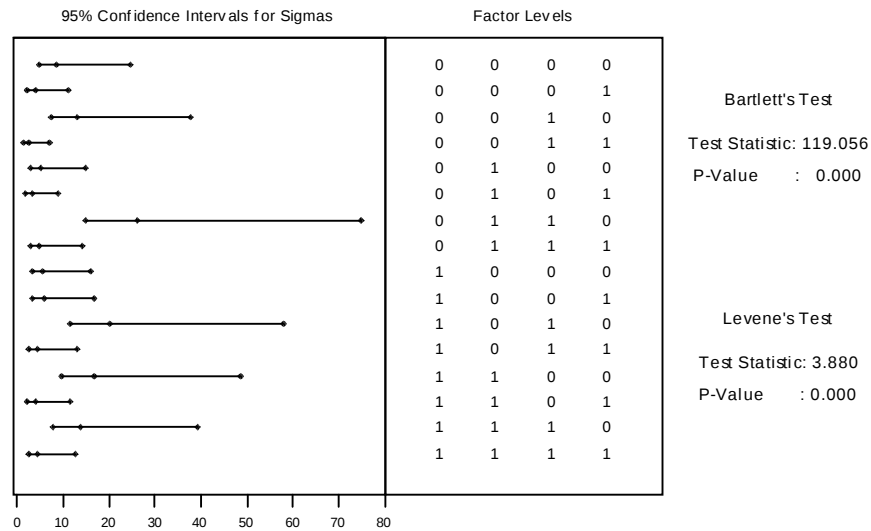


Figure 4. Homogeneity of variance test for Blue Final Strength.

Engaging in attacks with a force similar to the defender's is considered dubious without some sort of combat advantage. Commanders understand that force performance on any given day may be different. Variables, ranging from fatigue to logistical support to weather conditions, affect performance. Commanders will use combat multipliers, such as tactical surprise or added artillery support, to mitigate variable risks and improve possible battle results. But in the final analysis, as OneSAF shows, combat between similar forces will produce a variable outcome that is likely to be harder to predict.

Figure 5 represents the interactions between design factors in their influence of the Blue strength response. An interaction suggests the effect (change in means) associated with moving from one level of a factor to another is different depending on what level of a second factor is present. The matrix display format considers all possible two-way interactions, with column factor level appearing on the secondary (upper) x-axis and the Blue strength response appearing on the secondary (right) y-axis. The row factor level is indicated by line color. Graphical interpretation of a single plot cell in the matrix focuses on line slope of the line formed between the Blue strength means at the two levels of the column factor. Zero slope suggests no main effect for the column factor, moving from level 0 to level 1; a positive slope suggests a positive main effect; a negative slope suggests a negative main effect. Interactions are indicated according to the departure of the red and blue lines from parallel.

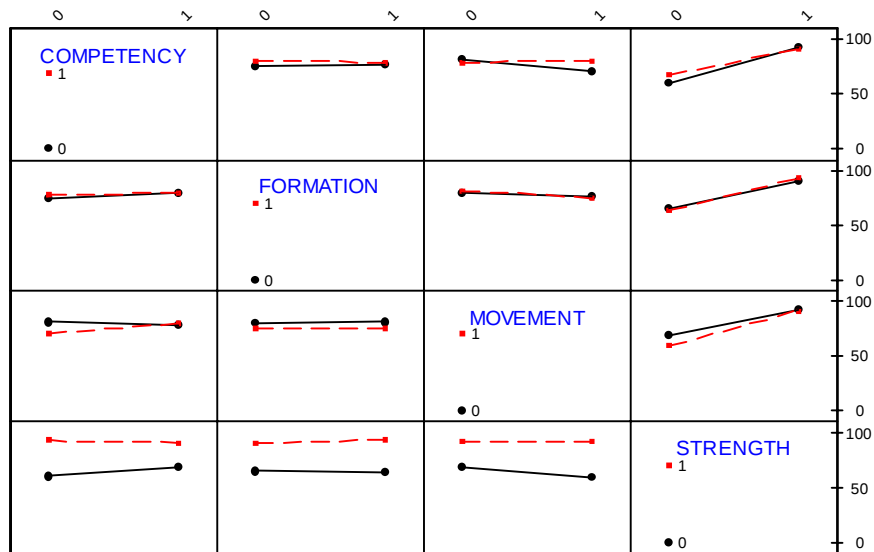


Figure 5. Interaction plot for Blue final strength.

Note the interaction between movement and strength in matrix cell (row 4, column 3). For uneven strength (red), the zero slope of the line indicates no difference in the Blue final strength between the march and traveling overwatch movement conditions. Conversely, for even strengths (black) the traveling overwatch condition shows a reduced Blue final strength compared with that under march. That the movement effect is different depending on strength means there is an interaction between the factors.

A second example in matrix cell (row 1, column 3) illustrates the interaction between competency and movement, with only the traveling overwatch condition adversely affecting Blue final strength.

When traveling overwatch was used, the preponderance of force in the uneven condition scenarios mitigated the disadvantages found in the even condition scenarios. The principle of battlefield mass is upheld by these observations in that a commander can improve battle results and overcome disadvantages by using more forces to accomplish an objective.

Figure 6 moves us to a second response measure, the percentage of tube launched, optically tracked, wire guided missiles (TOW) remaining among firepower capable vehicles at the end of battle. The structure of the graph is the same as that of Figure 2, except for the response. Stark differences in TOW use are seen between levels of force competency (e.g., CME, CTE). The reversal of the competency effect is interesting for these two. Slight differences are seen between levels of force competency for some other experimental conditions (e.g., CTU, WMU, CMU), again with competency differences sometimes increasing and sometimes decreasing the remaining percentage of TOWs.

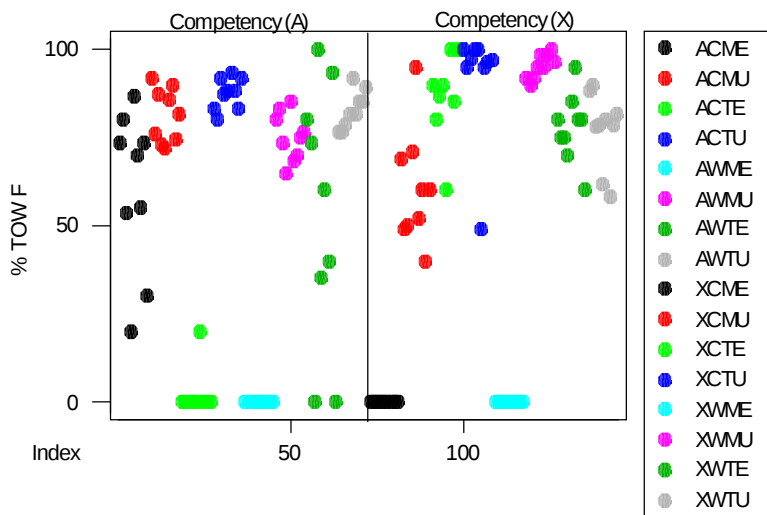


Figure 6. Comprehensive view of the percentage of TOW munitions remaining.

TOW usage analysis caused our investigation to focus on OneSAF procedure. Specifically, performance of OneSAF in the formation condition, i.e., traveling overwatch and march, explains our results. In the case of company sized units performing overwatch, the first platoon is assigned the duties of overwatch. This causes those entities to retreat out of position and follow the main force from behind to engage the enemy while the other platoons advance. During a march, all entities maintain position during the advance.

We learned that when planning a combat formation using column movement, the ordering of units, and not the crew expertise, is critical to success. OneSAF placed company formations in what seemed to be a random platoon ordering. Each friendly company in our scenarios had two armor platoons and one infantry platoon. The platoon order for ACME and ACTE is armor followed by infantry, followed by armor. The platoon order for XCME and XCTE is infantry, followed by armor, followed by armor. Keep these orderings in mind when reading the next two examples.

When the traveling overwatch procedure is used for ACTE, the first armor platoon moves to the rear, causing the more vulnerable infantry platoon to lead into combat. As is shown in Figure 6, few or no TOWs remain after combat because vehicles having that system are apt to be eliminated in combat. Using the XCTE ordering, the initial platoon (infantry) moves to the rear, allowing an armor platoon to lead into combat. The armor platoon serves to shield the infantry platoon, resulting in infantry entity survival and a number of remaining TOWs.

When the march procedure is used for ACME, the first armor platoon leads into combat. The armor platoon shields the infantry platoon next in line, causing some infantry entities to survive with TOWs remaining. Using the XCTE ordering, the initial infantry platoon

leads the formation and is in every case eliminated, leaving no TOWs at the end of combat, as shown in Figure 6.

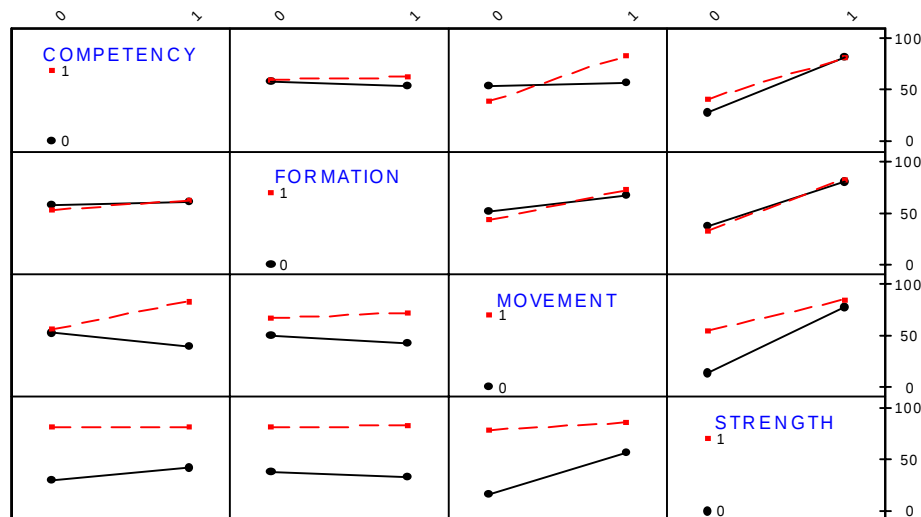


Figure 7. Interaction plot for the percentage of TOW munitions remaining.

Figure 7 reveals the two-way interactions involving the percentage of TOWs remaining. Interpretation and structure of the graph is the same as for Figure 5. Of the interactions possible, four appear strong: competency by movement, competency by strength, formation by movement, and movement by strength. For example, consider competency by movement in cell (row 1, column 3). For forces with average competency, there is no apparent change in the percentage of TOWs remaining between march and traveling overwatch movement, both leaving approximately 50%; for expert competency, TOWs are used far more sparingly in the traveling overwatch mode (approximately 80% remaining) than in the marching mode (approximately 40% remaining). Here again, we must be careful of variability. To understand this effect better, a three-way interaction is appropriate. Note, Figure 8 clearly shows the four conditions, ACTE, AWME, XCME, and XWME, exhaust all TOWs available. The impact of formation together with movement in ACTE complicates the interpretation of the interaction when we pool observations together to form the means for the two-way interaction.

The behavior of OneSAF with regard to our two WME conditions points out a fallacy in entity representations. The placement of the company put the entities within contact drill range at the beginning of the scenario, effectively causing all of the platoons to immediately assume a wedge formation. Although the friendly infantry platoon begins farther away from the threat in both scenarios, a combination of an inherent infantry advantage and threat ordinance choice causes the exhaustion of all TOWs. Specifically, infantry vehicles are faster than tanks in forming the wedge (an advantage), but subsequently move ahead of friendly tanks toward the threat (a disadvantage), coming within range of the threat tanks. As the infantry entities pause to fire TOWs at the threat tanks, fire is returned; however, the threat tanks choose to use tank gun ordnance that

appear to consistently miss at the engagement range. Once most of the TOWs have been expended, the infantry entities continue to advance toward the threat position ahead of the friendly tanks and are subsequently destroyed, leaving no remaining TOWs.

Movement strength also shows indication of interaction. The interaction in cell (row 3, column 4) shows that under traveling overwatch (red), TOWs remaining were approximately 55% for even strengths and 85% for uneven strengths, a difference of 30%; under march the change in percent TOWs remaining between even and uneven strengths is 78% to 15% respectively, a difference of 63%. Thus we can say that the change in the percentage of TOWs remaining subject to force strength conditions was more keenly felt under march than under traveling overwatch movement.

In the uneven cases, mass overcame force placement in march as well as infantry platoon formation speed, causing increased damage to threat forces before friendly entities could use a significant number of TOWs. Mass, coupled with the above explanation of infantry losses under WME conditions, causes the extreme difference in TOW usage to become clear.

Figure 8 clarifies the two-way interactions described in the previous figure in terms of a three-way interaction involving competency, movement, and strength. Each of the eight points on the graph represents the mean of 18 observations pooled over formation. The percentage of remaining TOWs is on the y-axis. Movement is on the x-axis. Points are jittered in the x direction (i.e., random variation is added) to make distinct points more visible. Lines representing average (A) and expert (X) competency are drawn for even strengths (red) and uneven strengths (blue).

The graph shows clearly the interaction among the three factors. For the uneven strength condition, there is no appreciable difference between expert and average forces in how they respond to march and traveling overwatch movement conditions. In all cases, approximately 20% of the TOWs were expended. However, for the even strength condition, average competency forces respond vastly differently than expert forces. Average forces expended approximately 70% of TOWs for march and traveling overwatch; expert forces expended 100% of the TOWs for march and only 20% for traveling overwatch. Thus the interaction between competency and movement depends on the level of strength.

Most of the indications for TOWs remaining were given above. The three-way chart in Figure 8 reveals an increased instability found at more even combat ratios. When conditions are considered even, adjusting the formation or movement of a force can have serious impact on battle outcome; whereas, most of those same choices do not impact the fight when attacking with superior odds. The mitigation of losses and logistics use can be realized through superior numbers.

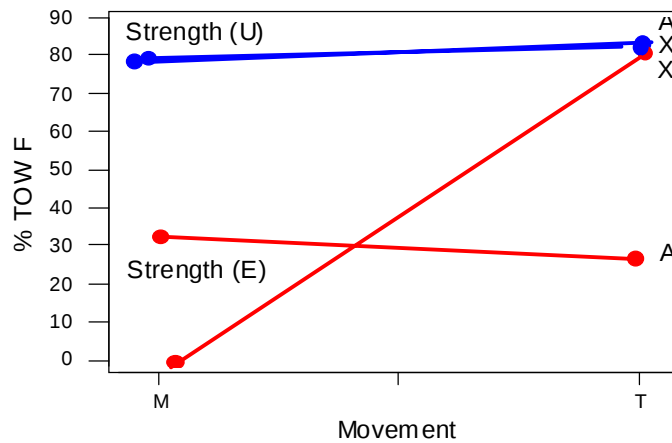


Figure 8. Three-way interaction plot for the percentage of TOW munitions remaining.

7.0 Conclusion

The findings in this report represent an effort to understand a combat simulation using a small data subset collected according to sound experimental design. So closely related is physical representation to military interaction that we chose to present our findings by intermixing the two through narratives. These narratives of simulation operation describe several interesting trends and behaviors, making the potential for feedback into the simulation development process high.

With the tremendous number of independent variables involved in combat, a true simulation is beyond the current capability of military science. Yet combat simulations do provide important training and intelligence tools. Our endeavor is to increase the understanding of limitations and advantages found in combat simulations. Simulation outcomes are applicable to real world military operations only when the means to those outcomes are clearly understood.