

12th ICCRTS

“Adapting C2 to the 21st Century”

Title: Putting the science back in C2: What do the buzzwords really mean?

Topics: C2 Concepts, Theory and Policy, Cognitive and Social Issues, C2 Metrics and Assessment

Authors: Paul Havig, Denise Aleva, George Reis, and John McIntire

Contact: Paul Havig

AFRL/HECV

2255 H St

Wright-Patterson AFB, Ohio 45433

Phone: 937-255-3951

Fax: 937-255-8366

e-mail: Paul.Havig@wpafb.af.mil

Adresses:

Denise Aleva

AFRL/HECV

2255 H St

Wright-Patterson AFB, Ohio 45433

George Reis

AFRL/HECV

2255 H St

Wright-Patterson AFB, Ohio 45433

John McIntire

AFRL/HECV

2255 H St

Wright-Patterson AFB, Ohio 45433

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 2007		2. REPORT TYPE		3. DATES COVERED 00-00-2007 to 00-00-2007	
4. TITLE AND SUBTITLE Putting the science back in C2: What do the buzzwords really mean?				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) AFRL/HECV,2255 H St,Wright-Patterson AFB,OH,45433				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 Twelfth International Command and Control Research and Technology Symposium (12th ICCRTS), 19-21 June 2007, Newport, RI					
14. ABSTRACT Current command and control (C2) terminology is laden with buzzwords that may, or may not, be useful to helping advance the science of C2 (e.g., effects-based operations (EBO) or sensemaking). In theory each term was devised for a reason, however, more often than not the reasons are lost and the terms are bantered about as "proof" of a good system, experiment, etc. We review some of the major terms and their history, as well as the potential evidence for their scientific integrity. Next we discuss how best to understand these terms by investigating their psychological (e.g., cognitive and social) as well as decision making roots. Finally we show how one may develop experiments and then eventually systems that either test or use these terms as they were originally defined. We give examples of how we are attempting to test these ideas in the laboratory as well as how others may test them in the future. We conclude with some discussion about the usefulness of buzzwords in the C2 realm as well as ways to keep them effective exemplars of their original meanings thus helping to advance the theory as well as knowledge of C2 systems.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 36	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

Putting the science back in C2: What do the buzzwords really mean?

Paul Havig
Denise Aleva
George Reis
John McIntire*

*Air Force Research Laboratory
Wright-Patterson Air Force Base, OH
Consortium Research Fellows Program

Abstract

Current command and control (C2) terminology is laden with buzzwords that may, or may not, be useful to helping advance the science of C2 (e.g., effects-based operations (EBO) or sensemaking). In theory each term was devised for a reason, however, more often than not the reasons are lost and the terms are bantered about as “proof” of a good system, experiment, etc. We review some of the major terms and their history, as well as the potential evidence for their scientific integrity. Next we discuss how best to understand these terms by investigating their psychological (e.g., cognitive and social) as well as decision making roots. Finally we show how one may develop experiments and then eventually systems that either test or use these terms as they were originally defined. We give examples of how we are attempting to test these ideas in the laboratory as well as how others may test them in the future. We conclude with some discussion about the usefulness of buzzwords in the C2 realm as well as ways to keep them effective exemplars of their original meanings thus helping to advance the theory as well as knowledge of C2 systems.

Introduction

The impetus for this paper spawned from the fact that certain words, “buzzwords,” get picked up by both upper management and at times the scientific community in general and we, as scientists, are told to defend our work based on them. Three simple examples will underscore the frustration this comes about when attempting to conduct scientific research yet told to defend in terms of a buzzword. The first example was a situation in which the current zeitgeist in the laboratory was to do “cognitive science” and we were told it was the “latest thing.”

Unfortunately, cognitive science has been around for decades as noted when one peruses Collins and Smith (1988). This well put together collection of readings in cognitive science starts with Turing’s (1950) seminal paper on testing intelligence in machinery. Arguments aside over the start date of the founding of cognitive science, the real frustration follows from two facts. First, cognitive science as a discipline was turned into a buzzword without understanding as to what cognitive science was, and secondly it was obvious that no one had noted that many of us were already tackling our research with multidisciplinary teams, a foundation of cognitive science. Luckily this era was short-lived and it was not necessary for any of us to ever need to add an obligatory “cognitive science” bullet to any of our viewgraphs.

A second example arose when one of the authors was defending some proposed research and was asked if it was "...cognitive enough?" Here the word cognitive is not even used properly. The Merriam-Webster Online dictionary (www.m-w.com) defines cognitive as:

- "1: of, relating to, being, or involving conscious intellectual activity (as thinking, reasoning, or remembering) <cognitive impairment>
- 2: based on or capable of being reduced to empirical factual knowledge"

Thus the implication would be that either our research must not seem to have any "...conscious intellectual activity" or that there was a fear that it could not be "...reduced to empirical factual knowledge." After a brief discussion, the truth was that the word cognitive was misused and the intention was to understand if issues relating to cognitive psychology (i.e., memory load, attentional issues, etc.) were being addressed. In this case, simple communication resolved any semantic ambiguity that might give rise to an erroneous evaluation of a program.

The third example highlights an issue that has been turning up at many conferences lately. Namely, after an author presents their data there is at least one person who is bound to ask "Did you check your data and then correct for violations of sphericity?" Statistically, a violation of sphericity means that the variances in one's data set are heterogeneous which violates the basic assumptions of the analysis of variance. However, if one looks deeper into the issue, having heterogeneous variances is not a problem if your significance value is large. In fact, correcting for violations of sphericity in this case does very little to effect the outcome. The problem is that most people do not know this fact. Thus if one says they did not need to correct for violations of sphericity, then some may assume the analysis is faulty (for an excellent review of sphericity see: Field, 1998). However, if one understands the issue at hand, a correction may, or may not, be called for in the analysis.

The point of these examples is that there is a very real potential for buzzwords to potentially hinder scientific research as we are forced to take time to prove our research worthy in terms of the word of the day. Or, they confuse the issue and detract from the research when they are not use properly in the first place. Further, these examples use relatively simple words; we in this paper will delve into the realm of buzzwords for Command and Control (C2) applications which raises its own set of issues. Moreover, the burden of understanding the buzzwords falls upon all involved. If one group uses the buzzword without understanding and the other group either doesn't know what it means, or has an incorrect understanding of the meaning, no intelligent communication can take place. This point serves to underscore the fact that to keep a word or phrase from becoming a buzzword careful definitions must be set forth so that confusion is kept to a minimum.

Thus, the question that may be raised is "Why use buzzwords?" To lay the groundwork, let us first define what exactly a buzzword is and determine why they are used. The Merriam-Webster Online dictionary (www.m-w.com) defines buzzword as:

- "1 : an important-sounding usually technical word or phrase often of little meaning used chiefly to impress laymen
- 2 : a vogueish word or phrase -- called also *buzz phrase*"

A similar definition is found on Wikipedia (www.wikipedia.com) which also adds the caveat that these words often have unclear meanings which in and of itself can lead to arguments, as we will see when we get to network-centric warfare (NCW). Some may view that Wikipedia, the online encyclopedia that is open for edit by anyone who may want to help with the topic, may not be a proper reference. However, we include this definition as it shows a consensus of the meaning of the word and more importantly it cites reasons as to why one may (or may not) want to use a buzzword. According to Wikipedia, buzzwords define new concepts, they are intentionally vague so that no one can question them, or they are intentionally vague so as to give rise to new ideas and concepts by forcing discussion of their meaning. This paper attempts the later and further pushes to provide examples of how these buzzwords may be tested in a scientific manner as the authors are themselves scientists who wish to provide evidence that their research is indeed relevant in the current and future states of military operations. In that vein, it is important to also understand that we are not here to “attack” buzzwords nor even their meaning in this paper. As we will see, that has been done elsewhere. However, we choose to attempt to show either current research, or proposed ideas, that may or may not, show the viability of certain concepts, as summed by Joseph Joubert “Le but de la dispute ou de la discussion ne doit pas être la victoire, mais l'amélioration” (The aim of argument, or of discussion, should not be victory, but amelioration.) (Raynal and Joubert, <http://visualiseur.bnf.fr/CadresFenetre?O=NUMM-88671&M=notice>)

The course of the rest of this paper will be to discuss three specific buzzwords, their background, definitions, and how they may be investigated in the course of a scientific experiment. The three we will look at, effects-based operations, sensemaking, and network-centric warfare are all very “hot” or buzzworthy topics today. Whereas we feel that their use as buzzwords has at times been misused or used without understanding. There are some basic principles of these words that can indeed be tested in a rigorous scientific fashion. Moreover, we hope to show not necessarily the “correctness” of the words but a way forward in which they may be tested in the laboratory.

That being said, we need to address the issue of experimentation. This term itself may cause confusion if not defined at the outset. Our definition of experimentation is in the classical experimental psychology approach in which one develops a hypothesis to test, develops an experiment to test the hypothesis, conducts the experiment, and then performs statistical analysis to assess the results of the originally stated hypothesis. In Alberts and Hayes (2002) these are termed “hypothesis testing experiments,” as opposed to “demonstration experiments” (or sometimes termed technology demonstrations). Our approach is that investigating small portions of the issue will only serve to enhance the bigger picture. In fact, there are times that portions of technology demonstrations do not go as planned and these must be investigated on a simpler level (sometimes even as a hypothesis testing experiment). We feel that both are needed, and if done properly, form a continuous loop from hypothesis testing experiment to demonstration experiment and back again to further refine new and innovative technologies.

Example One – Effects-Based Operations (EBO)

The first buzzword we discuss is EBO defined by Smith (2002) as “Effects-based operations are coordinated sets of actions directed at shaping the behavior of friends, foes, and neutrals in peace, crisis, and war.” This concept is not new with many historic examples provided by Smith

(2002) and well defended with more recent historical examples by Deptula (2001). In fact, based on the definitions above it might be argued that EBO does not fit as a buzzword. However, there have been times when one of us was asked if our research “supports EBO.” Given this type of question one could argue for the “en vogue” use of EBO and thus we tackle it as an “easy” example.

As previously mentioned, we are in the vein of doing hypothesis testing experiments (from here on out experiments for ease of discussion). First, it must be pointed out that we do not agree with Smith’s (2002) assertion that we (as psychologists) would say that EBO is nothing more than a series of stimulus-response (S-R) chains. Nor do we agree that this is even an acceptable way to explain EBO. This type of response points out exactly the buzzword problem when either: one party does not understand the phraseology being used or when the two parties are not on the same page as to what is meant by the page. S-R psychology is largely the realm of the behaviorists of the 1930’s through the 1960’s. Much interesting work was done in that time but in this day and age we realize that there is more to human perception, cognition, emotion, etc. than simple S-R chains with a lack of cognition per se. For an interesting overview, one could read Hebb’s (1958) textbook on psychology. Comparing that work to what we know about perception, cognition, and especially neuroscience shows how far the field of psychology has come in the last 50+ years in understanding the human brain, and with that, how an outdated concept such as an S-R chain of response is woefully inadequate to explain even the simplest human behavior. The main problem with this line of explanation is that a strict S-R account would say for example that dropping ordinance on a certain building (stimulus) and a group of people getting upset (response) does not take into account the multitude of issues that may surround the response (e.g., economic, political, social, cultural, etc.). Thus, almost taking away from the strength of EBO planning and execution in the first place. While it may be an abstraction, a very good way to view EBO is Smith’s (2002) “operations in the cognitive domain.” This is the true heart of EBO and may lead to the most fruitful attempts when designing experiments.

The first question that arises when one wants to design an experiment is “What hypothesis is being tested?” It is important to note that throughout this paper we are only discussing those types of experiments that are conducted with people. In other words, we leave the world of modeling and simulation to the experts. Many examples of EBO, and simulations thereof, have been done and presented elsewhere (for two interesting papers see: Wagenhals and Levis, 2002 and Wagenhals and Wentz 2003). However we will express our ideas in terms of hypothesis testing questions that we will discuss in terms of proposed experiments.

This brings about an interesting choice of three examples which we will explore in turn. In the first, one may ask “Is EBO better than non-EBO?” Granted this may seem an extreme example, but the question that arises is one of whether or not an experiment can be designed to test this situation in which one group uses EBO planning and execution and the other does not. The first issue would be to define EBO (and we can take the definition above) and then attempt to develop a system for planning and execution that does not use EBO (and this itself may be a heroic feat). A possible outcome would be that planners in each condition (with and without EBO) could come up with the exact same plan. If that is the case, how then would we be able to determine how and why identical plans arose when the two groups ostensibly used differing techniques?

One might take this argument further and say that planning and execution is always about shaping behavior as the end result can run a gamut from winning the hearts and minds of the enemy forces to submission to complete annihilation. Likewise there is also the possibility that the two planners come up with different plans yet the outcomes are identical. The question eventually in both cases comes down to how one measures the differences between the two groups in both their planning and execution processes.

A second example would be to leave the “Does EBO work?” question alone and focus on “How can we develop processes to enhance EBO?” Here we have decided to buy into EBO and assume it is the optimal way to plan and execute operations and then tailor our research to help support the method. In fact, this tact is what several of us are attempting to do in terms of developing new visualizations for C2 applications. In time condensed situations, it is often not optimal to have to read long paragraphs explaining situations so that a decision may be made. However, much information can be conveyed by a well conceived visualization. The problem with much of today’s visualizations are that they are directly taken (or descended from) Mil-Std 2525 Common Warfighting Symbology (1999) which is intended for showing entity (e.g., tanks) or group (e.g., battalion) information. We are attempting to look at visualizations that convey meaning to the user to help in the planning process understand the potential outcomes (effects) of differing operations. Take as an example a situation in which there is a river with a bridge, blue forces on one side, red on the other, and neutrals on both sides. The commander’s intent is to keep the red forces from coming over the bridge. The question is what are the options and how do we portray not only the tactic (e.g., destroying the bridge, or sending more forces to protect the bridge) but the outcome of the tactic. Let us assume the decision is to take out the bridge. Using EBO we know there is an effect of that action and we can portray those outcomes visually to help enhance the decision making process. The difficult portion here is making a visualization (say of neutral forces angry because you have destroyed their main transportation route), while the easy test is how well others can interpret what you meant. A good example is shown by Mayhorn, Wogalter, and Bell (2004) in which they evaluated a series of safety symbols designed for homeland security. In their study, they showed participants a safety sign and asked them to say what it meant by the sign. Results went from 7 – 100% correct for all the signs tested. We plan on developing a series of visualizations and testing them in this same way, and then refining and testing them again until we get good consensus that all who view the visualizations “see” the same thing. Next we will then use them to show various outcomes of a certain course of action (e.g., destroying a bridge) and asking participants what each of the various outcomes mean in terms of EBO. The benefit of these studies is that we arrive at a specific number (the percent correct) and can further learn from the mistakes of the participants so as to refine our visualizations.

A third example is much less tenable at the current time but is getting close to a possibility. What we may be able to do in the future is simulate operations and effects given the growing ability to make “intelligent” agents. For example, there are many popular video games in which agents learn and adapt to player strategies. If we could simulate to a high enough fidelity, one could imagine being able to simulate EBO and then analyze the results of the tests. Perhaps a first test would be a simulation of past historic campaigns where we know the operations and the effect.

A final question is that of the real world validity of any of these example experiments. Done properly, any experiment should be able to generalize from the sample to the population. The issue that may be raised though is the fact that these would be done in tightly controlled laboratory settings. While this is true, that does not preclude their generalizability into the “real world.” In fact, we see Alberts and Hayes (2002) “demonstration experiments” as the perfect venue for testing real world validity. The continuum as mentioned previously is naturally from hypothesis testing experiments to demonstration experiments and back until an acceptable solution is found.

Example Two – Sensemaking

The second buzzword we discuss is sensemaking which has been defined in many different but related ways. For example, Russell, Stefik, Pirolli, & Card (1993) define sensemaking as “...the process of searching for a representation and encoding data in that representation to answer task-specific questions.” For the C2 realm, this is a concise summary which might be roughly translated as understanding the data (search for representation and encode data) to perform mission planning (answer task-specific questions. Similarly, Alberts, Garstka, Hayes, & Signori (2001), define the way organizations engage in sensemaking as “(the way in which)...they relate their understanding of the situation to their mental models of how it can evolve over time, their ability to control that development, and the values that drive their choices of action.” Further, sensemaking involves three steps:

- 1) Generating alternative actions intended to control selected aspects of the situation.
- 2) Identifying the criteria by which those alternatives are to be compared.
- 3) Conducting the assessment of alternatives.

The main difference is that Russell et. al., (1993) talk in terms of representations and encoding data as they come from a more algorithm and computing approach, whereas Alberts et. al. (2001) describe mental models as they are attempting to describe mission planning and other military functions. Weik (1995) also adds that there are seven properties of sensemaking: it is grounded in identity construction, is retrospective, is enactive of sensible environments, is social in nature, is ongoing, is focused on and by extracted cues, and finally, it is driven by plausibility rather than accuracy. A good way to understand the complexity of the differing viewpoints on sensemaking is seen by the conceptual models of sensemaking developed in Leedom (2002) which run the gamut from a simple idea of the data fitting a story to complex ideas taking into account story, awareness, action, information, etc. in a variety of feed-forward and feed-backward loops. The common thread in all of these definitions from a psychologist point of view is that they all involve definitions and words that have been used in the fields of cognitive psychology, decision making psychology, artificial intelligence, and cognitive science. Attempting to discredit the word sensemaking would be tantamount to discrediting hundreds of years of research on the human! As such, perhaps it is better to call sensemaking not a buzzword but a theory. Note that it can still be used in a buzzworthy fashion, but as shown in our very first buzzword example, almost any word can be abused as such.

As such, to attempt to prove or disprove sensemaking would be even harder than the example we discussed for proving EBO. There is no way to control the representations, schemas, mental

models, etc. that humans create when problem solving and so it would be impossible to test, for example, planning and execution with and without sensemaking. We are not saying that humans cannot be trained to perform certain tasks, nor that they would not perform these tasks quite well. The arising problem is that we are attempting to understand if a task can be done without sensemaking. As we see it, as stated above, human cognition and problem solving is always about sensemaking. One does not go through their daily life trying to “not” understand the world.

However, there is another direction that we feel is a viable course of action, namely, developing ways to “enhance” sensemaking abilities. Specifically, one of the jobs to be done when performing sensemaking is to gather, integrate, analyze and assess data. In terms of the definitions, this includes incorporating this information into one’s own mental model of not only the problem space but also one’s world view. Research in one of our laboratories is looking at how best to portray center of gravity and other data rich information in a graphical format to help users gather, understand, interpret, and use this information more quickly and efficiently. The tool is called Visualization of the Operational Environment for Understanding & Response (VOEUR). The point of this tool is not to give the user something they do not have but rather to enhance their ability (we hope) to process this information. Any piece of software can fail and one is left to doing their task the “old way,” however giving the user a new way to look at and understand their world should help to bring about new ways of processing necessary information.

How then might one test this tool in an experimental setting? Several options are available. Let us use as an example that we are assessing intelligence information about red force communications. We can assess how quickly participants can understand a communications network based on the tool they are given to understand the information. For example, they may be given a situation that is purely text, a combination of text and simple graphics, or the VOEUR tool with enhanced graphics and interaction ability. One could see how long it takes to reach an answer, the correctness of the answer, what types of mistakes are made in interpretation, etc. In general, one could argue that any tool that allows the user to do their job more efficiently, more accurately, and with less mistakes has enhanced the user’s sensemaking abilities. One should be cautious in setting up a straw man so that the newly devised tool is given an unfair advantage over current tools. Overall, it can be concluded that any tool developed that enhances decision making abilities and helps users to understand not only their environment as well as the problem space can be said to enhance sensemaking.

An interesting question arises in terms of how this tool could be evaluated in terms of Doctrinal, Organizational, Training, Materiel, Leadership, Personnel and Facilities (DOTMLPF) frameworks. Indeed, if this tool worked in several experiments and was proven to be worthwhile, the next logical step would be further testing in terms of DOTMLPF. However, in this paper we are addressing the buzzwords question at the “bench scientist” level. As such experiments (at any level) serve to help enhance either a tool (as above) or understanding (as in what is EBO?). We could expound for pages on developing a formal acquisition line for any of these examples but shall not do that here. This is not to say it would not be a worthy endeavor but rather that it is very much beyond the scope of this paper and indeed the scope of a “simple” experiment that we can run in the laboratory.

Example Three – Network-centric Warfare

The final buzzword we discuss is network-centric warfare (NCW) which is most easily defined by its tenets as outlined in Alberts (2002):

- 1) A robustly networked force improves information sharing.
- 2) Information sharing and collaboration enhance the quality of information and shared situation awareness.
- 3) Shared situation awareness enables self-synchronization.
- 4) These, in turn dramatically increase mission effectiveness.

Although there are earlier papers on NCW (Cebrowski and Garstka, 1998) we use the Alberts (2002) paper due to the simple way it outlines NCW as shown above. We felt this was a good way to keep the argument as simple as possible and set the definitions up front.

As stated in the beginning of this paper, we have decided not to attack the buzzwords but rather assess whether it is in the realm of possibility to test them in hypothesis testing experiments. In the case of NCW, the criticism has been done elsewhere and need not be reiterated here (Bolia, 2005; Bolia, Vidulich, Nelson, & Cook, 2006; Griffin & Reid, 2003a, b). The question then becomes how does one go about testing NCW? Whereas hypothesis testing experiments have been criticized for lack of real world validity, conversely demonstration experiments lack the necessary control to precisely know why certain outcomes occurred. For example, one may argue that as in tenet one, we are a superior armed force because we are robustly networked. Yet, that robust connectivity (when working and we won't even go down that rabbit hole here) is not the only change in the last 5, 10, however many years one chooses. Furthermore, can it be proven that a robustly networked force improves information sharing? One may point to the internet and the myriad of topics one can discover as "evidence" of improved information sharing. But, even with that robust connectivity there are times when even the most precise search leads to null or near null results. Or even more importantly searches turn up too much information and one must somehow determine the validity of the information on any given web page.

As an example, we could have multiple groups with the same or a similar set of goals (e.g., destroy buildings A, B, then C) but with varying levels of network interconnectivity. The levels of "network-centrality" could be varied in each group by permitting or restricting communication between nodes within the network or by restricting information flow to a single direction through some nodes. For example, in Figure 1, group 1 might only have one-way communication with their commander, who hands down the orders that are dutifully followed with little or no communication thereafter. In this case, there are few connections between nodes and information is primarily flowing in only one direction.

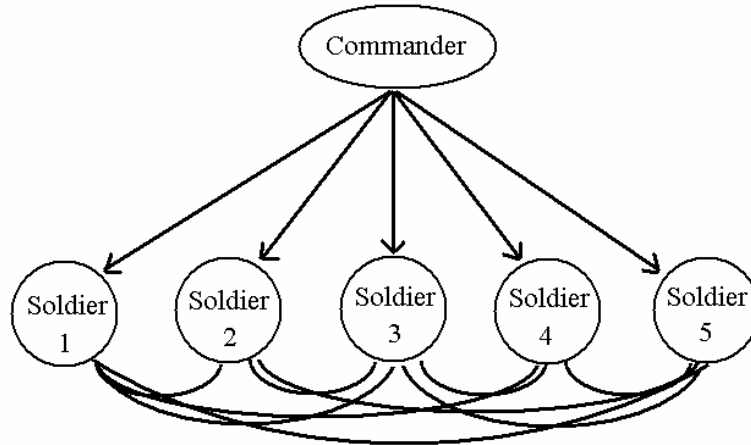


Figure 1 - Group 1. Information flow is primarily one-way, from the commander to the group, although there is unrestrained connectivity between the group members themselves. Notice that there are 5 one-way links and 10 two-way links.

Likewise, a second group might allow two-way communication between a single foot soldier and his or her commander. In this instance, the speed and efficacy of communication between the commander and the foot-soldiers is limited by the go-between who relays the messages of the commander to the remaining soldiers (and vice-versa). Information can easily flow to and from the central node (the go-to soldier) but again, the limited number of connections between nodes could hamper the flow of communication as shown in Figure 2.

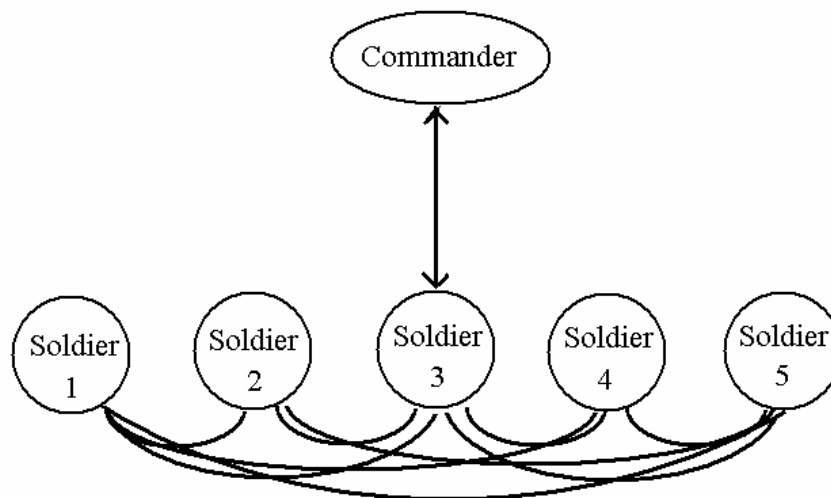


Figure 2 - Group 2. Information between the commander and the group flows (two-way) through a single group member (Soldier 3). Again, there is unrestrained connectivity between members of the group. Notice that there are a total of 11 two-way links.

A third group might allow for simultaneous communication between all nodes of the network, i.e., all soldiers can communicate directly with each other and with the commander. In this case, there is high connectivity and no limitation on the flow of information between nodes of the network as shown in Figure 3.

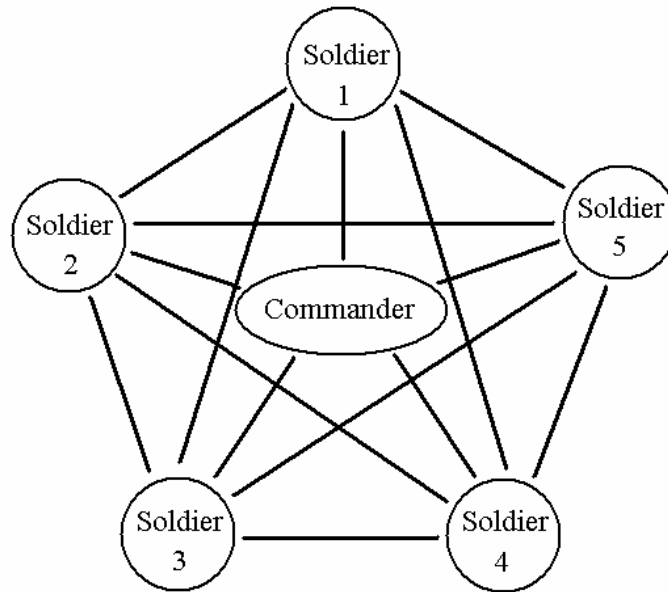


Figure 3 - Group 3. Information flows unrestrained through the entire network, which is fully inter-connected. In this group, there are 15 two-way links.

While setting up these types of networks may not be new and indeed may have been tested at various levels from simple to communication to information transfer, we posit to use them only as independent variables that are manipulated to assess higher-level functioning of group behavior. By assessing the performance (e.g., situation awareness, goal effectiveness, etc.) of each of these groups, one might be able to quantify the effects of numbers and types (i.e., one-way or two-way) of connections between nodes. Thus, we could directly test the advantages gained by using NCW, especially tenet number one. More independent variables that might be beneficially studied using such techniques are: varying the forms of communication that link the nodes (voice + video communications, voice communications, instant-messaging, etc.). Additional dependent variables that might be beneficially studied are: the speed of information flow through the network, the quality of information flow (message degradation), etc. One might even argue that parts of tenet two could be tested as a certain message could be passed up (or down) the chain of command and then tested as to its accuracy thus looking at the effect of quality of information due to collaboration. However, the other components of tenets two and three, namely shared situation awareness are much harder to test and even define (see Bolia and Nelson, in press; Vidulich, Bolia, and Nelson, 2004). However, using techniques outlined by these papers one may be able to use the example above as a way of controlling the information sharing and then testing situation awareness.

Conclusions

As shown we have proposed several simple techniques that may be turned into full-fledged hypothesis testing experiments to help better understand some of today's buzzwords. The question may still remain if buzzwords are needed or not, or if the terms above are even to be considered buzzwords. Nonetheless, we have shown that it may be possible to test experimentally some of these issues in order that they are better understood.

Acknowledgement

We would like to thank Robert Bolia for many fruitful discussions that contributed to this paper as well as the example of "Is EBO better than non-EBO".

References

- Alberts, D.S., (2002). *Information Age Transformation: Getting to a 21st Century Military*. Washington D.C.: Command & Control Research Program.
- Alberts, D.S., Garstka, J.J., Hayes, R.E., & Signori, D.A. (2001). *Understanding Information Age Warfare*. Washington D.C.: Command & Control Research Program.
- Alberts, D.S., and Hayes, R.E. (2002). *Code of Best Practice Experimentation*. Washington D.C.: Command & Control Research Program.
- Bibliothèque Nationale de France. (n.d.). "Pensées, essais, maximes et correspondance de J. Joubert", <http://visualiseur.bnf.fr/CadresFenetre?O=NUMM-88671&M=notice>.
- Bolia, R. S. (2005). Intelligent decision support systems in network-centric military operations. *Intelligent Decisions? Intelligent Support? Pre-proceedings for the International Workshop on Intelligent Decision Support Systems: Retrospects and Prospects*, 3-7.
- Bolia, R. S., & Nelson, W. T. (in press). Characterizing team performance in network-centric operations: Philosophical and methodological issues. *Aviation, Space, and Environmental Medicine*.
- Bolia, R. S., Vidulich, M. A., Nelson, W. T., & Cook, M. J. (2006). The use of technology to support military decision-making and command & control: A historical perspective. In Cook, M. J., Noyes, J. M., & Masakowski, Y. (Eds.), *Decision Making in Complex Systems*. Aldershot, UK: Ashgate Publishing Ltd.
- Cebrowski, VAdm Arthur K., USN, and John J. Garstka. (1998) Network Centric Warfare: Its Origin and Future." *In Proceedings of the Naval Institute* 124:1, 28–35.
- Collins A., and Smith, E.E. (1988). *Readings in Cognitive Science: A Perspective from Psychology and Artificial Intelligence*. San Mateo, California: Morgan Kaufman Publishers, Inc.

Deptula, David, A. (2001). *Effects-Based Operations: Change in the Nature of Warfare*. Arlington, Virginia: Aerospace Education Foundation.

Field, A. P. (1998). A bluffer's guide to sphericity. *Newsletter of the Mathematical, Statistical and computing section of the British Psychological Society*, 6 (1), pp. 13-22.

Giffin, R. E., & Reid, D. J. (2003a). A woven web of guesses, canto one: Network-centric warfare and the myth of the new economy. *Proceedings of the 8th International Command and Control Research and Technology Symposium*. Washington: Command and Control Research Project.

Giffin, R. E., & Reid, D. J. (2003b). A woven web of guesses, canto two: Network-centric warfare and the myth of inductivism. *Proceedings of the 8th International Command and Control Research and Technology Symposium*. Washington: Command and Control Research Project.

Hebb, D.O. (1958). *A Textbook of Psychology*. Philadelphia, Pennsylvania: W.B. Saunders Company.

Leedom, D. K. (2002). *Final Report: Sensemaking Experts Panel Meeting 25 – 26 June*. Washington D.C.: Command & Control Research Program.

Mayhorn, C.B., Wogalter, M.S., & Bell, J.L. (2004). Homeland Security Safety Symbols: Are We Ready?, *Ergonomics in Design*, Fall, 6 – 14.

Merriam-Webster Online Dictionary. (n.d.) <http://www.m-w.com/dictionary/cognitive>.

Merriam-Webster Online Dictionary. (n.d.) <http://www.m-w.com/dictionary/buzzword>.

Mil-Std 2525B (1999). *Common Warfighting Symbolology*. Washington D.C.: Department of Defense.

Russell, D.M., Stefik, M.J., Pirolli, P., & Card, S.K. (1993). The Cost Structure of Sensemaking, *In Proceedings of INTERCHI: Conference on Human Factors in Computing Systems*, 269 – 276.

Smith, E.A. (2002). *Effects Based Operations: Applying Network Centric Warfare in Peace, Crisis, and War*. Washington: Command & Control Research Program.

Turing, A.M. (1950). Computing Machinery and Intelligence. *Mind*, 59, 433-460.

Vidulich, M.A., Bolia, R.S., & Nelson, W.T. (2004). Technology, Organization, and Collaborative Situation Awareness in Air Battle Management: Historical and Theoretical Perspectives, *In Branbury, S. and Tremblay, S. (eds.) A Cognitive Approach to Situation Awareness: Theory, Measures, and Application*, Aldershot, UK: Ashgate Publishing Ltd.; 233 – 253.

Wagenhals, L.W., and Levis, A.H. (2002). Modeling support of effects-based operations in war games. *Proceedings of the 7th International Command and Control Research and Technology Symposium*. Washington: Command and Control Research Project.

Wagenhals, L.W., and Wentz, L.K. (2003). New effects-based operations models in war games. *Proceedings of the 8th International Command and Control Research and Technology Symposium*. Washington: Command and Control Research Project.

Weick, K. E. (1995) *Sensemaking in Organizations*. USA: Sage Publications, Inc.

Wikipedia: The Free Encyclopedia. (n.d.) <http://en.wikipedia.org/wiki/Buzzword>.

Putting the science back in C2: What do the buzzwords really mean?



Paul Havig

AFRL/HECV

Wright-Patterson AFB, Ohio



Co-authors



- **Denise Aleva - AFRL/HECV**
- **George Reis - AFRL/HECV**
- **John McIntire - AFRL/HECV (Consortium Research Fellows Program)**



Overview



- **Why the buzzwords?**
- **An easy example – EBO**
- **A harder example – Sensemaking**
- **The hardest example – NCW**
- **What good are buzzwords?**
- **The way ahead.**



Into the Lion's Den!



- Not here to attack buzzwords

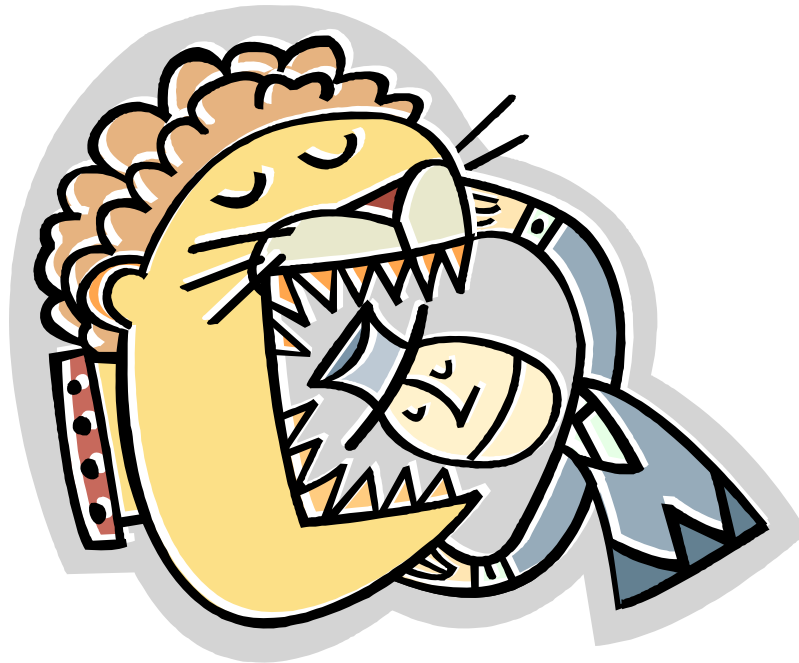




Into the Lion's Den!



- But to generate discussion





Buzzworthy annoyances



- **Where this all got started**
 - “Cognitive Science”
 - “Cognitive enough?”
 - Statistical conundrums
- **Bottom line: Must understand definitions**
 - At all levels of the research food chain



What is a Buzzword?



Merriam-Webster Online dictionary (www.m-w.com)

Buzzword:

- 1 : an important-sounding usually technical word or phrase often of little meaning used chiefly to impress laymen**
- 2 : a voguish word or phrase -- called also *buzz phrase***



What is a Buzzword?



Wikipedia (www.wikipedia.com):

Buzzwords:

- 1) Define new concepts**
- 2) Are intentionally vague so no one can question them**
- 3) Are intentionally vague so as to give rise to new ideas and concepts by forcing discussion of their meaning**



What is an experiment?



- **Hypothesis testing versus demonstration experiments**
 - **What is the difference**
 - **Pros and cons of each**
 - **What is our approach**
 - **Is it a continuous cycle of hypothesis to demonstration and back again?**



Effects-Based Operations (EBO)



- **Smith (2002) - “Effects-based operations are coordinated sets of actions directed at shaping the behavior of friends, foes, and neutrals in peace, crisis, and war.”**
- **Three examples of “EBO experiments”**



Effects-Based Operations (EBO)



- **Example 1**
 - **Is EBO better than non-EBO?**
 - **Potential outcomes**
 - **Both groups same plan same outcome**
 - **Both groups different plan same outcome**
 - **Can planning be done without EBO?**
 - **If so, how does one do/measure it?**



Effects-Based Operations (EBO)



- **Example 2**
 - **How do we develop a process to enhance EBO?**
 - **Assume EBO works**
 - **Develop a way to enhance the decision making process**
 - **Well designed visualizations?**
 - **Mayhorn, Wogalter, and Bell (2004)**
 - **Homeland security safety symbols**



Effects-Based Operations (EBO)



- **Example 3**
 - **Simulate operations and effects with intelligent agents**
 - **Current video games getting “smarter”**
 - **e.g., 1st person shooters in which the bad guys run away**
 - **If one can create and believe the fidelity**
 - **Could simulate and analyze results**
 - **How would we get there?**
 - **Simulate historical campaigns with known outcomes (both expected and surprising)**



Sensemaking



- **Russell, Stefik, Pirolli, & Card (1993) - “...the process of searching for a representation and encoding data in that representation to answer task-specific questions.”**
- **Alberts, Garstka, Hayes, & Signori (2001) - “(the way in which)...they (organizations) relate their understanding of the situation to their mental models of how it can evolve over time, their ability to control that development, and the values that drive their choices of action.”**
 - **Generating alternative actions intended to control selected aspects of the situation.**
 - **Identifying the criteria by which those alternatives are to be compared.**
 - **Conducting the assessment of alternatives.**



Sensemaking



- **Leedom (2002) several conceptual models**
 - **Run the gamut from a simple idea of the data fitting a story to complex ideas taking into account story, awareness, action, information, etc. in a variety of feed-forward and feed-backward loops**
- **Common thread in all is that they all involve definitions and words that have been used in the fields of cognitive psychology, decision making, artificial intelligence, and cognitive science**
- **Attempting to discredit the word sensemaking would be tantamount to discrediting hundreds of years of research on the human**



Sensemaking



- **No way to control the representations, schemas, mental models, etc. that humans create when problem solving**
 - **It would be impossible to test, for example, planning and execution with and without sensemaking.**
- **The problem (as in EBO) is that we are attempting to understand if a task can be done with or without sensemaking.**
 - **But human cognition and problem solving is always about sensemaking**
 - **One does not go through their daily life trying to “not” understand the world**



Sensemaking



- **But can you test it?**
 - Yes by evaluating your widget, program, etc. against the standard of the day
- **How do you know if it enhanced sensemaking?**
 - Is learning/training faster?
 - Is memory better? (e.g., prospective memory)
 - Is decision making better? (e.g., expected value judgments)
 - Is perception better? (e.g., minimization of attentional capture, change blindness, etc.)



Network-Centric Warfare (NCW)



- **Alberts (2002):**
 - **A robustly networked force improves information sharing.**
 - **Information sharing and collaboration enhance the quality of information and shared situation awareness.**
 - **Shared situation awareness enables self-synchronization.**
 - **These, in turn dramatically increase mission effectiveness.**
- **Arguments against can be seen elsewhere...so...**



Network-Centric Warfare (NCW)



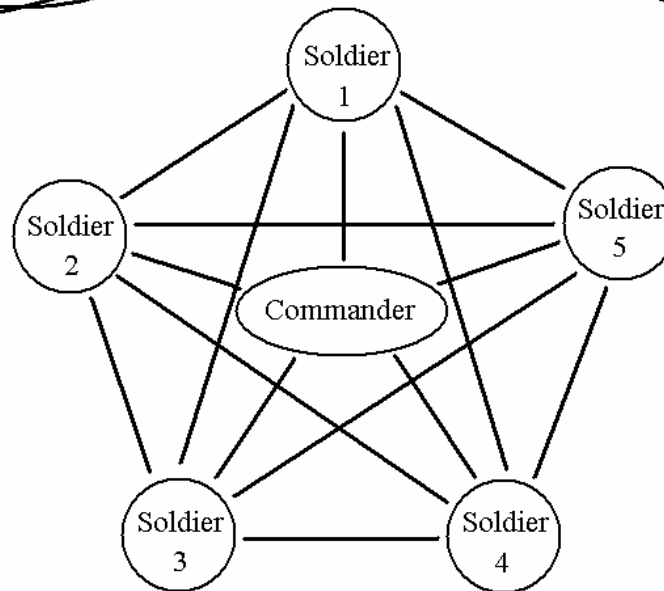
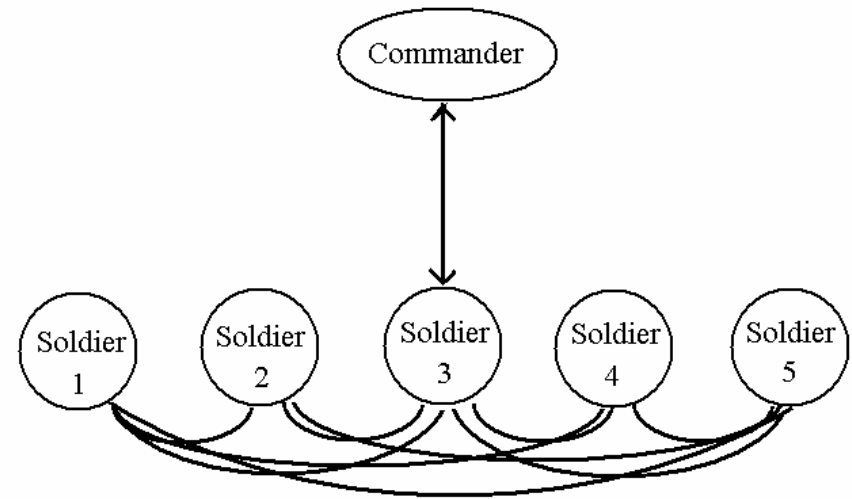
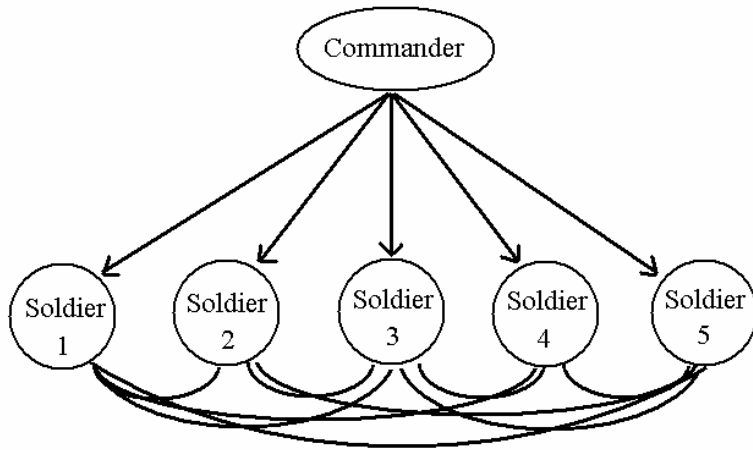
- **Can we test these tenants?**
 - **A robustly networked force improves information sharing. (Yes)**
 - **Information sharing and collaboration enhance the quality of information and shared situation awareness. (Yes)**
 - **Shared situation awareness enables self-synchronization. (Yes)**
 - **These, in turn dramatically increase mission effectiveness. (Well...more of a conclusion)**
- **So how?**



Network-Centric Warfare (NCW)



- Three groups same goal, different communications and connectivity

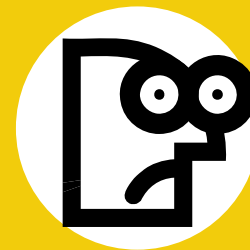
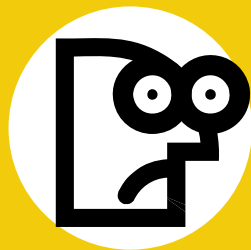




Conclusions



- **So...buzzwords may not be that bad if:**
 - The definition is known and understood
 - The ideas can be tested
 - Those using the words understand the above
- **Are they needed?**
 - Perhaps
- **Useful?**
 - Again perhaps
- **Will we ever be rid of them?**
 - No
- **Is that a bad thing?**



It's QUESTION TIME !!

Paul.Havig@wpafb.af.mil