

15th ICCRTS
“The Evolution of C2”

Development of Metrics for Trust in Automation

Topic 3: Information Sharing and Collaboration Processes and Behavior

Topic 8: C2 Assessment Metrics and Tools

Dr. Janet E. Miller

Air Force Research Laboratory

Sensors Directorate

Dayton, Ohio 45433

Janet.Miller3@wpafb.af.mil

937.255.2713 X4304

LeeAnn Perkins

711th Human Performance Wing

Human Effectiveness Directorate

Dayton, Ohio 45433

LeeAnn.Perkins@wpafb.af.mil

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE JUN 2010		2. REPORT TYPE		3. DATES COVERED 00-00-2010 to 00-00-2010	
4. TITLE AND SUBTITLE Development of Metrics for Trust in Automation				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) Air Force Research Laboratory,Sensors Directorate,Dayton,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 Proceedings of the 15th International Command and Control Research and Technology Symposium (ICCRTS '10), Santa Monica, CA, June 22-24, 2010					
14. ABSTRACT Research in ?trust? in automation has gained momentum and ?trust? has been identified as playing an essential role for implementing effective work-centered computer systems (Hoffman, Lee, Woods, Shadbolt, Miller & Bradshaw, 2009). In a socio-technical work system, the automation handles the majority of an algorithmically-intense workload but the human is generally a final decision-maker. Therefore, the human?s acceptance of the automation?s output is required for a successful result. Some researchers believe that system failures are connected to the human nature of trust, which is based on experiences with other humans, acting as the foundation for reliance on automated systems. However, using a common word as ?trust? allows for misunderstandings when used in multiple contexts. While all have some overtures of similarity, there are important unstated differences. Additionally, if trust is critical, then a method to accurately measure its goodness or level during active interaction between a human and automation would be beneficial. This paper will discuss three qualifiers for a trust evaluation such that measures can be developed to gauge a user?s trust perception over time; will lay out five components to specifically evaluate trust in automation, and propose a technique for measuring and monitoring trust in automation.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 18	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

Development of Metrics for Trust in Automation

Dr. Janet E. Miller
Air Force Research Laboratory
Sensors Directorate
Dayton, Ohio 45433
Janet.Miller3@wpafb.af.mil
937.255.2713 X4304

LeeAnn Perkins
711th Human Performance Wing
Human Effectiveness Directorate
Dayton, Ohio 45433
LeeAnn.Perkins@wpafb.af.mil

Abstract

Research in ‘trust’ in automation has gained momentum and ‘trust’ has been identified as playing an essential role for implementing effective work-centered computer systems (Hoffman, Lee, Woods, Shadbolt, Miller & Bradshaw, 2009). In a socio-technical work system, the automation handles the majority of an algorithmically-intense workload but the human is generally a final decision-maker. Therefore, the human’s acceptance of the automation’s output is required for a successful result. Some researchers believe that system failures are connected to the human nature of trust, which is based on experiences with other humans, acting as the foundation for reliance on automated systems. However, using a common word as ‘trust’ allows for misunderstandings when used in multiple contexts. While all have some overtures of similarity, there are important unstated differences. Additionally, if trust is critical, then a method to accurately measure its goodness or level during active interaction between a human and automation would be beneficial. This paper will discuss three qualifiers for a trust evaluation such that measures can be developed to gauge a user’s trust perception over time; will lay out five components to specifically evaluate trust in automation, and propose a technique for measuring and monitoring trust in automation.

Introduction

For many concepts, the English language lacks exquisite distinctions in words allowing for misunderstandings in communication and causing frustrations in research as people talk past

each other. The word 'trust' is just such a case. A simple search on Amazon.com for books that include the word 'trust' returns 710,654 titles as relevant. Perusing the returned list, you would note that the books run the gamut to include legal guidance (i.e., "Casenote Legal Briefs: Wills, Trusts & Estates"); art (i.e., "Trust the Process: an Artists' Guide to Letting Go") and religion ("Can We Trust the Gospels"). Going to SpringerLink, and performing the same simple search, 69,590 book, chapter, and journal article titles are returned with a more scientific bent but just as broad a usage ranging from computing science ("Aspects of General Security & Trust") to medicine ("Trust, benefit, satisfaction, and burden") to psychology (Foraging for Trust: Exploring Rationality and the Stag Hunt Game"). For comparison with another word often used in a variety of ways and contexts, in Amazon.com, 957,838 titles are returned for a 'love' search while in SpringerLink, only 41,687 titles are returned with a search on 'love.' One might expect a great number of titles for such a multi-faceted, broadly encompassing term as 'love' but surely most are surprised that 'trust' is used about as frequently, ambiguously and in so many contexts. Even in the United States' National Security Agency's Information Assurance Technical Framework document (2000), the term 'trust' is used 352 times, ranging from reference to the trustworthiness of technology, to a trusted human relationship to a Trusted Third Party. Interestingly, in the "Handbook of Trust Research," (2006) the word 'trust' itself is defined in only 9 of the 22 book chapters and invariably each accepts a somewhat different definition and reference.

Truly the word 'trust' is one of everyday parlance greatly sprinkled in everyone's daily communications, allowing much potential confusion especially when various scientific communities research the term and try to apply their findings to critical problems. To make progress with regard to understanding and researching 'trust,' the concept needs to be disambiguated for particular contexts and then a method to support capturing levels or degrees of trust at a time snapshot can be developed. This paper will discuss three qualifiers (context, components and object) for a trust situation such that measures can be developed to gauge a user's trust perception over time. Specifically focusing on trust in automation, five components were identified as relevant through a literature evaluation. An experiment was run to test the hypothesis that the five components are positively correlated to an overall evaluation of trust in

the experimental condition. The experiment method and results will be described and then a technique to actively measure and monitor the trust a human has in a system is proposed.

Background

Trust was a topic for discussion even in Socrates' day. Socrates never wrote his philosophies down and confined his viewpoints to spoken debate as he was concerned that portability and staleness of the written word could alter the author's intentions especially over time. Think today of the similar potential gap between the intentions of a computer programmer or system developer and the end user. For example, consider the global positioning system (GPS). Many are the tales of the end user who over-relied on a GPS to their detriment. For example, in December 2009, an elderly couple traveling from Grants Pass, Oregon to Reno, Nevada relied on their GPS for directions and got stuck in snow for three days when their GPS unit sent them down a remote forest road. Was the shortfall in the technology developers' viewpoint on how the technology would be used or with the user in not understanding what the system developers intended? Parasuraman and Riley (1997) discuss such types of technology usage issues as misuse, abuse, and disuse. They define *use* as the voluntary employment of an automation technology and discuss the factors that influence the decision to use, misuse, disuse or abuse a specific technology. *Disuse* is defined as the discontinuation or underutilization of technology; *misuse* is described as overreliance on a specific technology, and *abuse* is defined as inappropriate application of technology by designers or managers. (For extensive background on trust, refer to Lee and See (2004), Adams et al (2003) or Artz and Gil (2007) among many others.)

Trust itself has many definitions but most have some overtures of similarity. The Mayer et al (1995) definition is the most widely accepted definition of trust, "A willingness to be vulnerable to another party when that party cannot be controlled or monitored." If 'agent' is exchanged for 'party,' a somewhat more encompassing definition arises and this will be the general definition used for this paper. However, the definition still begs questions. Vulnerable to what extent? Vulnerable to what outcome? How willing? What are the ramifications of being vulnerable? Does the context matter? Monitored or controlled to what extent?

For such a broadly used word to be the basis for evaluating the performance of socio-technical systems, the authors propose that three qualifiers are required to focus an operational definition for trust. **One**, that the **context of interest** be sufficiently defined. For example, trust in the context of corporate financial dealings would be quite different in detail from trust with respect to internet chat rooms. Also, various contexts can entail differing levels of attributes such as vulnerability, risk, and reward all of which affect levels of trust. Levels of automation should also be included in the description of context and domain of interest. Parasuraman, Sheridan, and Wickens, C. D. (2000) proposed four broad classes of functions: 1) information acquisition; 2) information analysis; 3) decision and action selection; and 4) action implementation. Within each of these types, automation can be applied across a continuum of levels from low to high, i.e., from fully manual to fully automatic. For this paper, the context is a human interacting with decision support technology where the automation does not perform until the human permits the automation.

Two, that the term ‘trust’ be broken down into **lower level components** that allow measurement for the context of interest. This may in part address the concern in Dekker and Hollnagel (2004) that a generally used concept such as ‘trust’ be decomposed from a large construct into more measurable specifics. Additionally, decomposition can ensure a proper definition of ‘trust’ is used in the particular context of interest. For example, ‘loyalty’ was one of the keywords identified by Adams et al (2003) in their literature review but there may be disagreement on the word’s applicability for trust in automation.

For this paper, the components of trust are **competence, predictability, dependability, consistency and confidence**, which are five attributes often cited as contributors to trust in automation. These five attributes were chosen by an in-depth literature search of trust factors. A list of all the trust factors discussed was made, and a tally of occurrences was taken. The factors that had the most repeats were used to form the above list.

Competence

Competence is the ability to do the task at hand. The human’s perception of automation’s competence is critical when making decisions using and trusting in technology. The ability to do the task at hand is a vital component, and the user must be aware of how to judge the competency of the automation and place appropriate use. Several researchers have cited and

stated competence as a major influential factor in trust. See and Lee (2004) identified the following researchers being associated with tagging competence as a factor of trust in automation: (Barber, 1983; Butler & Cantrell, 1984; Kee & Know, 1970; Mishra, 1996).

Predictability

Predictability is the matching of performance with expectations. Predictability of automation plays an important role in trust. If the user can predict what the automation should do, then the user can adequately assess when the technology fails and how to perform the task without the automation. Muir (1994) discusses predictability as an important factor in a trust equation: trust = predictability + dependability + faith.

Dependability

Dependability, or always being there to perform, is important as the basis of trust is being vulnerable to another party (Lee & See, 2004). Rempel, et al (1985) offers dependability in automation to be essential in trust. The user must first be able to rely and depend on the automation to perform appropriately in building trust. Muir (1994) also includes dependability in the trust equation. If a user cannot depend on the automation, then the sole purpose of having automation is irrelevant and performance is ceased.

Consistency

Automation's ability to be consistent in performance is imperative for a user to build trust. Consistency is being free from variation or contradiction. If the automation does not produce similar outcomes to identical tasks, the user's trust can be skewed. Butler and Cantrell (1984) cite consistency as being one of the most influential factors of trust alongside competency. Inconsistency in automation is the first clue for a user to distrust and question the automation's validity (Lee & See, 2004).

Confidence

The final component of trust is confidence which is a user attribute toward the automation. Confidence is the user's certainty that the automation will perform appropriately. If a user has no confidence then the automation will not be used for the advantages. If the user is confident in the automation, trust can be built over time. If the user is not confident in the automation, then trust will not be able to be built over time. Going hand in hand with dependability, the user's confidence will be crucial in relying on automation to adequately perform (Moray, Inagaki, &

Itoh, 2000). If a user is too confident, he/she may abuse automation and cause damage (Lee & See, 2004).

The Table 1 is the survey used in the experiment that displays how the factors were presented and asked to be rated by the participant.

Table 1. Trust Factor Survey

Read each item and then circle the number of the response that best describes the extent to which you would rate the Route Planner's performance. <u>Indicate to what extent you generally feel this way.</u>	Not At All	A little	Sometimes	Frequently	All the Time
1. To what extent is the Route Planner competent in mapping out the routes?	1	2	3	4	5
2. To what extent can the Route Planner's routes be predicted?	1	2	3	4	5
3. To what extent can you rely on the Route Planner to plan the routes?	1	2	3	4	5
4. To what extent is the Route Planner consistent in planning the routes?	1	2	3	4	5
5. To what extent are you confident in the Route Planner's performance?	1	2	3	4	5

The experiment described in the next section was to investigate whether these five components are reasonable for defining the qualifier of 'lower level components' in a trust in automation situation. The hypothesis for the experiment was that the five components listed above are positively correlated to an overall evaluation of trust in automation in the experimental condition.

Three, that the **object of trust** be defined as research has noted the importance of the object of trust and the ability of individuals to discriminate trustworthiness of different targets. Trust is a relationship but just as love is a relationship and there is an object of love, similarly there is an object of trust. The object of trust in the experimental condition for trust in automation is a simulated global positioning system and the trust is one way.

Method

Participants

Ninety-five undergraduate students ($M = 20$, $SD = 3.96$) from a medium Midwestern university participated in the GPS simulation experiment. A within subjects experimental design was adapted where all participants completed all the GPS conditions.

Platform

The automated tool used for the experiment was a Route Planner that resembles a GPS in that it assists in determining directions to a destination (figures 1-3). However, the Route Planner only displayed the entire map for an area of interest on the screen while a standard GPS could displayed either the current intersection or the entire area map. In addition, the Route Planner had the following simulated wireless updating capabilities for use in different experiments in this research: traffic jams, car accidents, burning buildings, unsafe neighborhoods, riot outbreaks, and drive by shootings.

The algorithm underlying the platform provided an interesting aspect to this research. The automation determined the best course of action by first determining the next set of potential intersections from its current location. It then calculated which of those possible intersection points were closer to the final destination without considering intersections beyond that point. When the area map was complicated enough, this algorithm may have suggested the best next intersection was one which headed into a cul-de-sac or eventual dead-ends. This aspect was interesting as it can be used to stress all five factors of interest: competence, predictability, dependability, consistency and confidence.

In the experimental scenarios, the user was given a navigation goal to get from point “A” on the map to point “B” within a specific constraint described below. The user had the choice to either use the automated suggested route or create a manual route. The participant would map out the

entire route of his/her choice, but updated information would occur as the scenario proceeded, and the user would have the option of changing decisions based on the new information. Backtracking was not an allowed option. The following figures represent a typical decision pattern a user would go through with the Route Planner

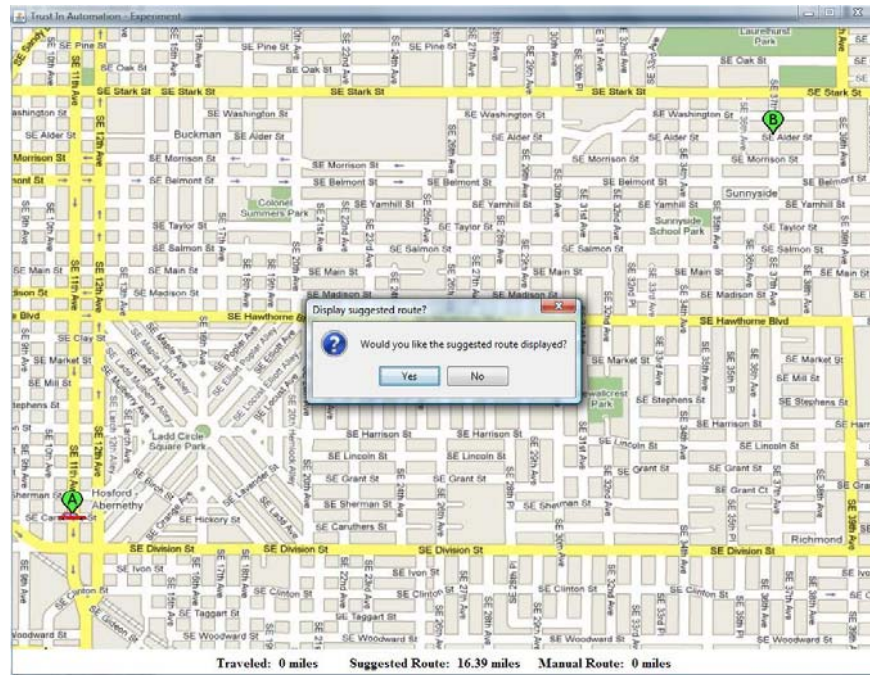


Figure 1: Representative Display Screen: Asking user if wanted suggested route displayed

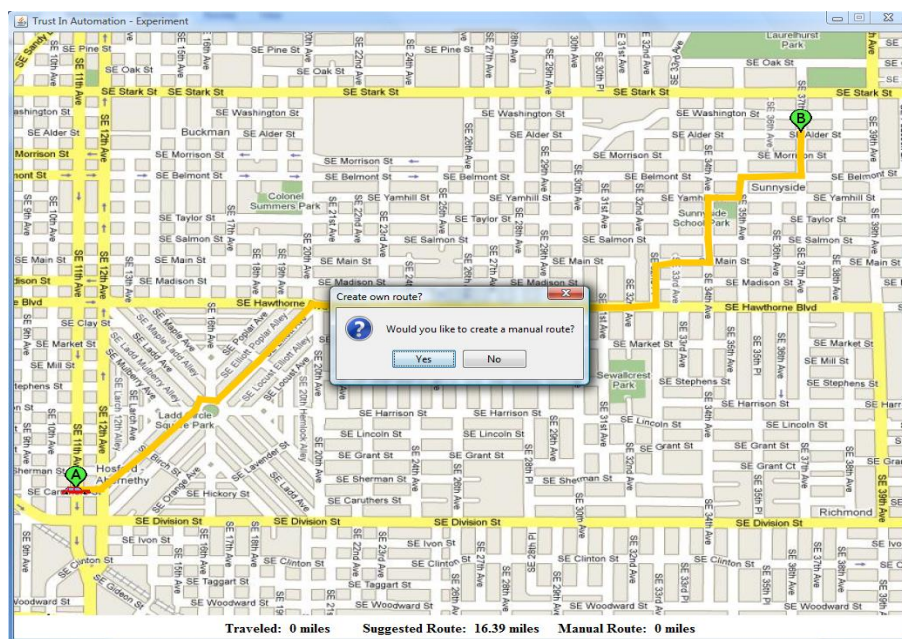


Figure 2: Representative Display Screen: Showing Suggested Route (Yellow): Asking user which route to select

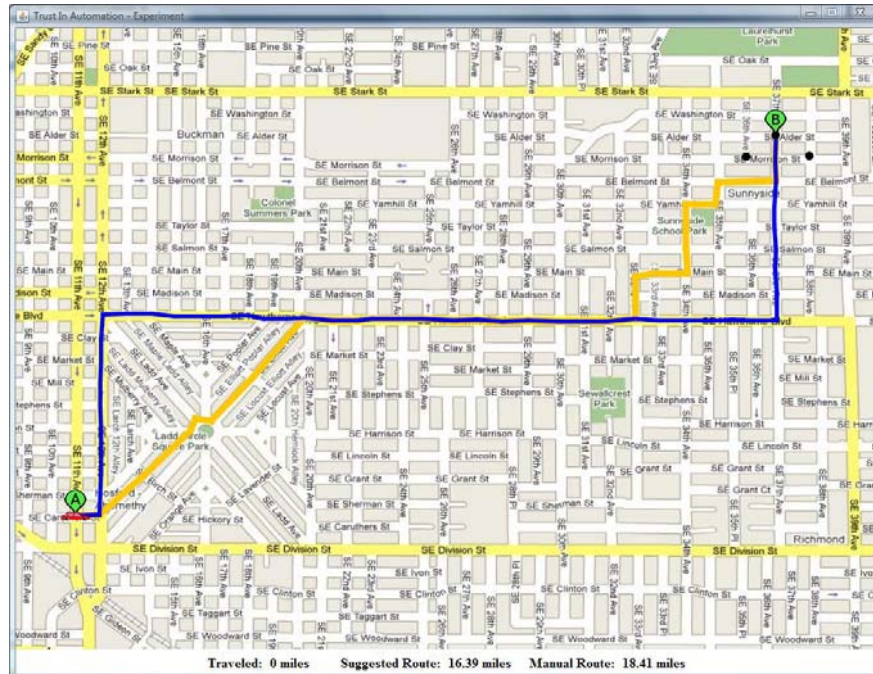


Figure 3: Representative Display Screen: Showing user choose to create a manual route (blue)

Procedure

Experiment One: Control Scenario

The objective given to the participant for the first run in the study using the Route Planner was simply:

- Travel from point A to B using the shortest distance.

For experiment one there was no context of time pressure manipulation or hazards in the testing. The purpose of this experiment was to observe human-automation interaction in a neutral setting. A distance ticker was provided to the participant to track miles used. A timer not visible to the participant calculated the time for each decision made by the user.

Experiment Two: Low Risk: Time Pressure Scenario

For the second experiment time pressure was added. The same objective was given to the subject except that the subject was given a time limit of thirty minutes for reaching the destination.

The same experimental setup from experiment one was implemented, and a distance ticker was provided. The timer was noted in this testing as well. The purpose of this experiment was to build on to the general experiment of observing human-automation interaction to see how time pressure affects the trust components.

Experiment Three: Medium Risk: Common Hazards Scenario

For the third experiment, a risk context was added. The same objective was given to the subject except that the subject was to avoid hazards and map out the shortest distance. Common hazards that included traffic jams and car accidents were randomly positioned in the scenario, and the subject was told that the suggested route would update to avoid all hazards. A distance ticker was provided to track the miles.

Experiment Four: High Risk: Uncommon Hazards Scenario

For the fourth experiment, there was a combination of a risk context and a time constraint added to the initial navigational goal. In this scenario, participants were asked to avoid all hazards and to get to destination “B” in twenty minutes or less. Uncommon hazards included burning buildings, unsafe neighborhoods, riot outbreaks, and drive by shootings that were randomly positioned in the map. Participants were again told that the suggested route would update to avoid all hazards. A time ticker was provided to track the minutes.

Randomization of Maps and Scenarios

Four maps and four scenarios were used to create sixteen conditions that would be randomized in order to counterbalance for the actual map display and the scenario order. Each map was of a location not in the Midwest and was stripped of identification so the participant would have no familiarity bias. In addition, each map would have all of the four experimental scenarios to create a counterbalance effect to ensure that the map did not influence route decisions. Lastly, a pseudo randomized route order list was computerized so that participants would not receive the scenarios in any particular order and would eliminate the possibility of bias responses based on risk building.

Data

The subjects were first given a demographic survey which included gender, age, class year, and familiarity and experience with GPS systems and computers in general. The subjects were video-taped using Cam Studio, with permission, during the experiments and the think-aloud protocol was used. During this procedure, participants were asked to talk aloud freely during the scenarios about the tasks and asked to give reasons for route selections. At the end of the experimental runs, the subjects were given questionnaires using a five point Likert scale addressing their level of competence, predictability, dependability, consistency and confidence after their experience with the system.

Results

The purpose of this section of the paper is to discuss the determination of whether there was sufficient correlation between the responses to the five point Likert scale on factors of competence, predictability, dependability, consistency and confidence and the responses for the overall degree of trust.

Table 2. Levels of Five Components of Trust

	Mean	Range	SD
1 - Competence	3.91	2-5	.65
2 - Predictability	3.44	2-5	.86
3 - Dependability	3.94	2-5	.77
4 - Consistency	3.99	2-5	.78
5 - Confidence	3.81	2-5	.87
Mean of Five Factors	3.81	2.2-4.8	.71

Table 2 shows the levels of the five components from the experiment. To determine the correlation, the mean of each participant's response for each of the five factors was averaged and considered as x. Each participant's response to the overall degree of trust was considered y. The sample correlation coefficient was used to estimate the population Pearson correlation between X and Y. For this experiment, the Pearson correlation was .67 which indicates highly moderate

correlation between the participant assigned Likert scale value of the factors of competence, predictability, dependability, consistency and confidence and the participant assigned value for overall trust in the GPS system.

General Discussion

For the purposes of this paper, the above correlation value supports that the five factors identified through the literature evaluation can be reasonably used as the lower level components to evaluate the level of trust a human has in an automated decision system with the context and object as described. The next challenge is how to use this information to actively and sufficiently measure and monitor trust in a socio-technical system.

Futuristically, the state-of-the-art will be adequately advanced to have identified what and how physiological measurements can be used to gauge the amount of trust a human has in an object. Although the science is not there yet, some efforts are being made. For example, the US's Intelligence Advanced Research Projects Activity (IARPA) has a new program entitled "Tools for Recognizing Useful Signals of Trustworthiness (TRUST)" in which sensing and validated protocols will be brought together to provide tools for assessing trustworthiness. Also, at George Mason University, research has been done on the neural correlates of trust (Kreuger, 2007). These efforts are still at the very basic level of research and usable results may be a long time coming. They may very well still need to be bounded by the qualifiers of context, components and object of trust. In the meantime, conscious-action trust management can be used to monitor the level of trust a human has for an object.

Now, trust is not a static attitude. As time progresses and aspects of a situation change, the level of trust will change. A basic principle of today's environment where adversaries seek to dominate networks to their advantage is that vulnerabilities and conditions need to be continuously evaluated and experimented. A more subtle and more dangerous situation would be one in which operators may continue to trust safe operation of systems that are no longer trustworthy. Therefore, trust must be constantly, consciously and actively monitored.

One solution would be to have a radar, or spider, diagram of the germane lower trust components relevant for the context and object of trust as a display on a user's desktop or other workspace. Depending on the level of potential hazards and risks of the automation and of the situation, the

user would intermittently note his values for the components of trust by clicking on the gridlines of the display. The entries would be time-stamped and logged. Evaluation of the radar diagrams and logs for trends would occur in near-real time especially looking for situations such as over-trust or for a degradation of trust. Thresholds could be set, most likely for either a certain level for any component or for a certain cumulative level or for a volume of the diagram to be filled.

An illustration can be made using the previously described experiment. Figure 4 shows a radar diagram from the experiment participant whose overall trust was low (mean of the five components of 2.2). Figure 5 shows a radar diagram from a participant whose overall trust level was high (mean of 4.4). At a glance, the difference in trust stance is shown. If these were of one person over time, the change in trust stance could easily be observed and the user queried to determine the cause. Was there an unpredicted response from the system? Did the user encounter inconsistent activity? Did something change the user's confidence in some way? The feedback mechanism would encourage the user and the monitor of the system, if there was such a monitor, to be alert for changes.

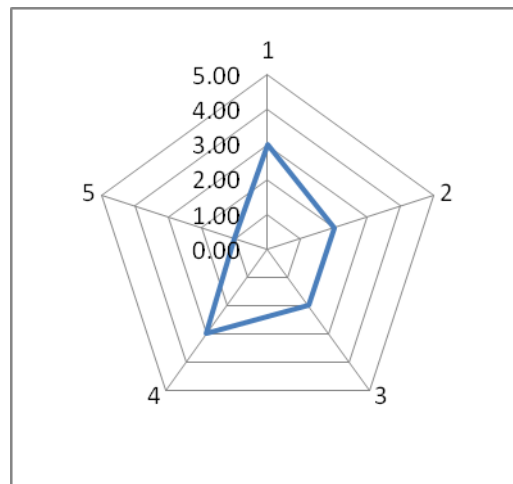


Figure 4. Low Degree of Trust (Participant was 2.2 of 5)

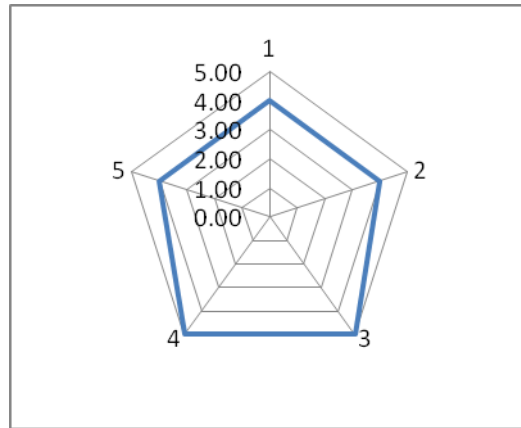


Figure 5 High Level of Trust (Participant was 4.4 of 5)

As with many research projects, there are implementation issues to address on both the monitoring and the input side. A first issue would be determining who would have the most benefit or insight into monitoring the trust stance. An information technology office's network control center (NCC) would probably balk at monitoring both the automated system's trustworthiness as well as the level of interactions the human perceives. However, the NCC would be the place likely to know to what extent a networked system should be trusted. Another potential trust stance monitor would be the security operations department of an organization. They would be the ones who would be versed in the adversary's intent and capability to determine if there is a threat and hence a need to change in a trust stance. A third place might be at the functional level, especially if the automation being monitored is a complex, analytic system with intricate algorithms. Only those versed in the specific technology might know when the loss in trust is because the user is not particularly familiar with the technology rather than the technology itself losing competence or other attribute. However, each of these areas within an organization should be interested in the level of trust a human has in particular automation as trust affects how the technology is used. The answer on how to monitor the trust stance, then, is the same as what a university professor often says. "It depends."

On the input side, questions exist on the balance of intrusiveness and the work situation in which this would be useful. If the radar display is too disruptive of day to day work, the user may submit inaccurate data just to be done with the task or not perform the task at all. Then the question arises on which automated activities need to be monitored. If the radar displays were implemented in my office environment, for example, I would consider being asked to evaluate

the network as represented by Outlook and Internet Explorer, as they are my networked tools, at the beginning and the ending of a logged session to be a reasonable request. Unless the operational environment was under a threat situation, I would not be happy to respond as frequently as one time per hour. If I used a complex, analytic technology, I would deem it reasonable to occasionally give feedback on my trust stance, again perhaps at the beginning and ending of a work session. However, if the technology were being used in a conflict environment and receiving live feeds from potentially infiltrated sources, such as an NCC monitoring the networks, it would be reasonable to have the user question their stance more often. Finally, if the situation is extremely fast-paced, such as a jet pilot in combat, having yet another widget to interfere with his main tasks would be unrealistic. A work and task analysis would need to be done to determine what trust issues there actually are in some given work domain or subdomain which would then identify the system, time and frequency for requesting a user to perform active monitoring.

Conclusion

This paper has considered the issue of trust, especially trust in automation, by proposing three qualifiers that need addressed to focus understanding in the broad topic of trust. The paper also presented the results of an experiment that support the reasonableness of using competence, predictability, dependability, consistency and confidence as the five lower level components of human trust in automation as identified by a trust literature evaluation as an example. Finally presented was a concept on how to use these factors in actively monitoring a user's stance toward trust in automation. If a networked socio-technical system is being evaluated for trust, the network and the socio-technical system as well as the human-automation trust stance must be evaluated.

Other published research on trust suggests lower level components of trust for different contexts and objects of trust. Erickson (2009) lists the key attributes of the US Air Force Research Laboratory's Sensors Directorate effort in pursuing trustworthiness in layered sensing as secure, safe and reliable. Hoffman et al (2009) discuss the need to have trust in macrocognitive work, or resilient, systems by addressing directability, responsiveness, reciprocity and responsibility. These attributes may be appropriate for measuring trust using a radar diagram with scale for their particular context and objects of trust once appropriately defined.

As discussed above, additional research is needed on how to implement an active trust monitoring system to include the frequency of elicitation, which systems need to be actively monitored, and the office of responsibility. However, as trust has risen to be a psychological phenomenon of importance and as development of physiological methods of measuring trust are still in their infancy, pursuing the radar graph methodology using lower level components as attributes in a well defined context and with a specific object of interest would be a feasible, expedient solution.

References

- Adams, B., Bruyn, L. Houde, S., & Angelopoulos, P. (2003) Trust in Automated Systems Literature Review Defence Research and Development Canada Toronto No. CR-2003-096
- Ajzen, I., & Fishbein, M. (1980). *Understanding attitudes and predicting social behavior*. Englewood Cliffs, NJ: Prentice-Hall
- Artz, D. and Gil, Y. (2007) A survey of trust in computer science and the Semantic Web. *Journal of Web Semantics: Science, Services and Agents on the World Wide Web*, Vol 5, Issue 2, 58-71
- Bachmann, R. and Zaheer (2006) *Handbook of Trust Research*. Northampton, MA; Edward Elgar
- Barber, B. (1983). *The logic and limits of trust*. New Brunswick, NJ: Rutgers University Press.
- Butler, J. K., & Cantrell, R. S. (1984). A behavioral decision theory approach to modeling trust in superiors and subordinates. *Psychological Reports*, 55, 19–28.
- Dekker, S. W. A., & Hollnagel, E. (2004). Human factors and folk models. *Cognition, Technology, and Work*, 6, 79–86.
- Hoffman, R, Lee J, Woods, D, Shadbolt, N, Miller, J. & Bradshaw J. (2009) The Dynamics of Trust in Cyberdomains. *IEE Intelligent Systems*, Volume 24, No. 6, pp 5-11
- Jian, J., Bisantz, A., Drury, C. & Llinas, J. (1998). *Foundations for an Empirically Determined Scale of Trust in Automated Systems*. Air Force Research Laboratory Report No. AFRL-HE-WP-TR-2000-0102. Wright Patterson AFB, OH: AFRL/HE.
- Kee, H. W., & Knox, R. E. (1970). Conceptual and methodological considerations in the study of trust and suspicion. *Conflict Resolution*, 14, 357–366.

- Krueger, F., McCabe, K., Moll, J., Kriegeskorte, N., Zahn, R., Strensiok, M., Heinecke, A. & Grafman, J. (2007). Neural Correlates of Trust. Proceedings of the National Academy of Sciences of the United States of America. October 25, 2007.
- Lee, J. D. & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46 (1), 50-80.
- Mayer R. C., Davis J. H., and Schoorman F. D. (1995). An integrative model of organizational trust, *Academy of Management Review*, Vol. 20, No. 3, pp. 709 - 734.
- Mishra, A. K. (1996). Organizational response to crisis. In R. M. Kramer & T. R. Tyler (Eds.), *Trust in organizations: Frontiers of theory and research* (pp. 261–287). Thousand Oaks, CA: Sage.
- Moray, N., Inagaki, T., Itoh, M., 2000. Adaptive automation, trust, and self-confidence in fault management of time-critical tasks. *Journal of Experimental Psychology: Applied* 6, 44–58.
- Muir, B. M. (1994). Trust in automation: 1. Theoretical issues in the study of trust and human intervention in automated systems. *Ergonomics*, 37, 1905–1922.
- National Security Agency. (2000). *Information Assurance Technical Framework, Release 3*. Washington, DC: NSA.
- Parasuraman, R., & Riley, V. A. (1997). Humans and automation: Use, misuse, disuse, abuse. *Human Factors*, 39, 230–253.
- Parasuraman, R., Sheridan, T. B., & Wickens, C. D. (2000). A model of types and levels of human interaction with automation. *IEEE Transactions on Systems, Man, and Cybernetics – Part A: Systems and Humans*, 30, 286–297.
- Parasuraman, R., Sheridan, T. & Wickens, C. (2008) Humans: Still Vital After All these Years of Automation. *Human Factors*, 50 (3), 511-520).
- Rempel, J. K., Holmes, J. G., & Zanna, M. P. (1985). Trust in close relationships. *Journal of Personality and Social Psychology*, 49(1), 95–112.