



NAVAL POSTGRADUATE SCHOOL

MONTEREY, CALIFORNIA

THESIS

UNDERSTANDING OPTIMAL DECISION-MAKING

by

John W. Critz

June 2015

Thesis Advisor:
Co-Advisor:

Quinn Kennedy
Jon Alt

**This thesis was performed at the MOVES Institute
Approved for public release; distribution is unlimited**

THIS PAGE INTENTIONALLY LEFT BLANK

REPORT DOCUMENTATION PAGE			Form Approved OMB No. 0704-0188	
Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instruction, 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, and to the Office of Management and Budget, Paperwork Reduction Project (0704-0188) Washington DC 20503.				
1. AGENCY USE ONLY (Leave blank)		2. REPORT DATE June 2015	3. REPORT TYPE AND DATES COVERED Master's Thesis	
4. TITLE AND SUBTITLE UNDERSTANDING OPTIMAL DECISION-MAKING			5. FUNDING NUMBERS	
6. AUTHOR(S) John W. Critz				
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Naval Postgraduate School Monterey, CA 93943-5000			8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING /MONITORING AGENCY NAME(S) AND ADDRESS(ES) N/A			10. SPONSORING/MONITORING AGENCY REPORT NUMBER	
11. SUPPLEMENTARY NOTES The views expressed in this thesis are those of the author and do not reflect the official policy or position of the Department of Defense or the U.S. Government. IRB Protocol number ____ N/A ____.				
12a. DISTRIBUTION / AVAILABILITY STATEMENT Approved for public release; distribution is unlimited			12b. DISTRIBUTION CODE	
13. ABSTRACT (maximum 200 words) The military has realized that their most valuable and adaptable assets are its leaders. Understanding optimal decision-making will allow the military to more effectively train its leaders. The Cognitive Alignment with Performance Targeted Training Intervention Model (CAPTTIM) was developed to aid the training of optimal decision making. CAPTTIM determines when decision performance (categorized as near-optimal or suboptimal) is aligned or misaligned with cognitive state (categorized as exploration or exploitation): when someone thinks they have figured out the task (exploitation cognitive state), is their decision performance actually near optimal? Prior research categorized subjects' cognitive states as exploration or exploitation, but the delineation of decision performance had yet been done. The primary focus of this thesis was to use pre-collected and de-identified data to (1) determine and validate a threshold that delineated near-optimal and suboptimal decision performance with the metric, regret, and (2) categorize the combination of cognitive state and decision performance into CAPTTIM on a trial-by-trial basis. A change point analysis of regret provided an effective threshold delineation of decision performance across all subjects. Visualization techniques were employed to categorize decision and cognitive state data into CAPTTIM on a trial-by-trial basis. Thus, CAPTTIM was validated as a means of understanding decision-making.				
14. SUBJECT TERMS optimal decision-making, regret, Iowa gambling task, exponentially weighted moving average, change point analysis			15. NUMBER OF PAGES 97	
			16. PRICE CODE	
17. SECURITY CLASSIFICATION OF REPORT Unclassified	18. SECURITY CLASSIFICATION OF THIS PAGE Unclassified	19. SECURITY CLASSIFICATION OF ABSTRACT Unclassified	20. LIMITATION OF ABSTRACT UU	

THIS PAGE INTENTIONALLY LEFT BLANK

Approved for public release; distribution is unlimited

UNDERSTANDING OPTIMAL DECISION-MAKING

John W. Critz
Captain, United States Marine Corps
B.S., University of North Carolina at Charlotte, 2008

Submitted in partial fulfillment of the
requirements for the degree of

**MASTER OF SCIENCE IN
MODELING, VIRTUAL ENVIRONMENTS, AND SIMULATION (MOVES)**

from the

**NAVAL POSTGRADUATE SCHOOL
June 2015**

Author: John W. Critz

Approved by: Dr. Quinn Kennedy
Thesis Advisor

LTC Jon Alt
Co-Advisor

Chris Darken
Chair, MOVES Academic Committee

Dr. Peter Denning
Chair, Department of Computer Science

THIS PAGE INTENTIONALLY LEFT BLANK

ABSTRACT

The military has realized that its most valuable and adaptable assets are its leaders. Understanding optimal decision-making will allow the military to more effectively train its leaders. The Cognitive Alignment with Performance Targeted Training Intervention Model (CAPTTIM) was developed to aid the training of optimal decision making. CAPTTIM determines when decision performance (categorized as near-optimal or suboptimal) is aligned or misaligned with cognitive state (categorized as exploration or exploitation): when someone thinks they have figured out the task (exploitation cognitive state), is their decision performance actually near optimal? Prior research categorized subjects' cognitive states as exploration or exploitation, but the delineation of decision performance had yet been done. The primary focus of this thesis was to use pre-collected and de-identified data to (1) determine and validate a threshold that delineated near-optimal and suboptimal decision performance with the metric, regret, and (2) categorize the combination of cognitive state and decision performance into CAPTTIM on a trial-by-trial basis. A change point analysis of regret provided an effective threshold delineation of decision performance across all subjects. Visualization techniques were employed to categorize decision and cognitive state data into CAPTTIM on a trial-by-trial basis. Thus, CAPTTIM was validated as a means of understanding decision-making.

THIS PAGE INTENTIONALLY LEFT BLANK

TABLE OF CONTENTS

I.	INTRODUCTION.....	1
A.	BACKGROUND	1
B.	REINFORCEMENT LEARNING IS NECESSARY TO REACH OPTIMAL DECISION-MAKING	2
1.	The Iowa Gambling Task.....	3
2.	Convoy Task	5
3.	Cognitive Alignment with Performance Targeted Training Intervention	5
C.	REGRET	7
D.	THESIS GOALS.....	10
II.	METHODS	13
A.	STATISTICAL SOFTWARE: R STUDIO	13
B.	METHODS USED TO DELINEATE HIGH AND LOW REGRET	14
1.	Exponentially Weighted Moving Average (EWMA)	14
a.	<i>Explanation of EWMA Equation and Uses.....</i>	14
b.	<i>EWMA of Regret.....</i>	16
2.	Simple Moving Average	19
3.	X-Bar Control Chart	22
4.	Change Point Analysis	25
5.	Final Method: Combination of Control Chart and Change Point Analysis	25
III.	RESULTS	29
A.	OVERVIEW OF COGNITIVE STATE DATA DEVELOPED FROM PRIOR RESEARCH	29
B.	CHANGE POINT ANALYSIS COMBINED WITH COGNITIVE STATE DATA.....	33
1.	Delineating High and Low Regret Using Change Point Analysis	33
2.	Combining Cognitive State and Decision Performance to Categorize Subjects within CAPTTIM	36
C.	VALIDATION OF CHANGE POINT ANALYSIS AND COGNITIVE DATA AS CAPTTIM CATEGORIZATION METRICS	42
IV.	DISCUSSION.....	47
A.	SUMMARY OF METHODS USED TO COMPLETE THESIS GOALS.....	47
B.	IMPLICATIONS.....	49
C.	FUTURE WORK.....	51
D.	CONCLUSION	52
	APPENDIX A. PAYOUT SCHEDULE FOR IGT AND CONVOY TASK	55
	APPENDIX B. R SCRIPTS	63

A.	EWMA OF DECISION LATENCY TIMES R SCRIPT	63
B.	CHANGEPOINT ANALYSIS R SCRIPT	71
C.	CAPTTIM VISUALIZATION R SCRIPT	72
D.	CORRELATION TEST R SCRIPT	76
E.	EXECUTE R SCRIPT	77
LIST OF REFERENCES.....		79
INITIAL DISTRIBUTION LIST		81

LIST OF FIGURES

Figure 1.	The Iowa Gambling Task screenshot (from Sacchi, 2014).	4
Figure 2.	The combination of cognitive state and actual decision performance indicates whether a trainee's cognitive state is aligned or misaligned with actual performance. When misalignment occurs, it indicates the need for a training intervention (from Kennedy, 2015).	6
Figure 3.	Payout schedule excerpt. The blue cell indicates the optimal decision; the yellow cell shows the subject's selection on trial 1; the green cell indicates the subject's selection on trial 2.	10
Figure 4.	EWMA of regret for Subject 14 using mean regret. Mean regret proved to be inappropriate as it was performing a EWMA on an already averaged regret value. This accounted for the much less volatile spikes in regret value. The large red dots are high damage instances. The medium blue dots are medium damages, and the small green dots are low damage instances. The shaded red area is where the EWMA is above the threshold and the shaded green area is where the EWMA is below the threshold. The threshold is calculated as 0.5 times the standard deviation of the mean regret.	17
Figure 5.	EWMA of regret for Subject 14 using regret received at each trial. The volatility in high regret is seen with the sharp red peaks which is where regret reaches values of 1250 for high friendly damage. The red, blue, and green dots are for high, medium, and low damages respectively. These dots are plotted along the mean regret line. Shaded red areas are above the threshold, while shaded green areas are below the threshold. The threshold is defined as the standard deviation of the regret received per trial.	18
Figure 6.	Simple moving average of regret per trial for Subject 1 with a window of 5 trials. The solid blue line shows the averaged regret and how high values in regret influenced the average for the 4 previous and 4 successive trials. Had a simple moving average not been used, high values of regret would be single vertical lines.	20
Figure 7.	Simple moving average of regret per trial for Subject 1 with a window of 10 trials. The solid blue line shows the averaged regret and how high values in regret influenced the average for the 9 previous and 9 successive trials. Comparison to Figure 6 shows how, for the same subject, the spikes in high regret are broadened by utilizing a larger window.	21

Figure 8.	Histogram of regret data for Subject 1. This clearly illustrates that the majority of regret values experienced by Subject 1 are of magnitude 50 and that the high spikes in regret only occurred a handful of times.	23
Figure 9.	X-Bar control chart for Subject 1. The solid blue line is the simple moving average that was previously discussed. The dashed red line is the upper control limit. The upper control limit is defined as the median plus the median absolute deviation and is recalculated every 20 trials.	24
Figure 10.	Change point analysis for Subject 1. The solid black line is the regret per trial data. The solid red lines are the process means returned by the change point analysis—they represent the process mean for that range of trails. The large spikes in regret incurred a change in the process mean that spanned the single trial in which the regret was incurred.	26
Figure 11.	EWMA of latency in decision-making times for Subject 4. The y-axis is latency in decision-making times and the x-axis is the number of trials. The colored dots represent damage incurred and are plotted at the actual latency in decision-making time versus the EWMA. The color and size of the dot is correlated with the level of damage incurred on the preceding trial. Red dots are high damage, blue dots are medium damage, and green dots are low damage. The orange shaded regions are where the EWMA is above the threshold (exploration) and the blue shaded regions are where the EWMA is below the threshold (exploitation).	31
Figure 12.	EWMA of latency in decision-making times for Subject 14. The y-axis is latency in decision-making times and the x-axis is the number of trials. The colored dots represent damage incurred and are plotted at the actual latency in decision-making time versus the EWMA. The color and size of the dot is correlated with the level of damage incurred on the previous trial. Red dots are high damage, blue dots are medium damage, and green dots are low damage. The orange shaded regions are where the EWMA is above the threshold (exploration) and the blue shaded regions are where the EWMA is below the threshold (exploitation).	32
Figure 13.	Change point analysis for Subject 4. The y-axis is the regret per trial value, while the x-axis is the trial number. The red lines are the process means returned from the change point analysis. The spikes in the regret value are a result of the subject receiving heavy friendly damage and incurring high regret. These spikes result in a change point that exists over	

	just one trial. The other, longer red lines are where the process mean did not change for that range of trials.....	34
Figure 14.	Change point analysis for Subject 14. The y-axis is the regret per trial value, while the x-axis is the trial number. The red lines are the process means returned from the change point analysis. The spikes in the regret value are a result of the subject receiving heavy friendly damage and incurring high regret. These spikes result in a change point that exists over just one trial. The other, longer red lines are where the process mean did not change for that range of trials.....	35
Figure 15.	CAPTTIM categorization algorithm. This figure illustrates how each subject is categorized in CAPTTIM based on decision-making performance (measured by regret) and cognitive state (measured by latency in decision-making times).....	37
Figure 16.	CAPTTIM categorization chart for Subject 4. The color-coded bar at the bottom of the chart correlates to the category color found within the CAPTTIM model. Yellow is high regret and exploration. Orange is low regret and exploration. Red is high regret and exploitation. Green is low regret and exploitation.....	38
Figure 17.	CAPTTIM categorization chart for Subject 14. The color-coded bar at the bottom of the chart correlates to the category color found within the CAPTTIM model. Yellow is high regret and exploration. Orange is low regret and exploration. Red is high regret and exploitation. Green is low regret and exploitation.....	39
Figure 18.	EWMA of latency in decision-making times for Subject 11. The x- and y-axis are the same as the previously described graphs. Note that Subject 11's EWMA of latency in decision-making times never falls below his/her threshold (shaded orange region). This subject spent the entire time exploring the environment and never exploited any decisions.....	40
Figure 19.	CAPTTIM categorization chart for Subject 11. Note that the values are coded yellow, orange and red. The only reason that Subject 11 was ever categorized as red (high regret and exploitation) within CAPTTIM was due to the fact that subjects are penalized for choosing routes 1 and 2 after trial 100. Subject 11's final damage score was 2200.	41
Figure 20.	Correlation between final damage score and number of trials spent in the red category of CAPTTIM. The red dots show a strong negative correlation between number of trials spent in the red category and final damage score.	42
Figure 21.	Correlation between advantageous selection bias and number of trials spent in the red category of CAPTTIM. The red dots	

	show a strong negative correlation between number of trials spent in the red category of CAPTTIM and the subject's advantageous selection bias.	43
Figure 22.	Correlation between final damage score and number of trials spent in the green category of CAPTTIM. The green dots show a moderately strong positive correlation between number of trials spent in the green category and final damage score.....	44
Figure 23.	Correlation between advantageous selection bias and number of trials spent in the green category of CAPTTIM. The green dots show a moderately strong positive correlation between number of trials spent in the green category and advantageous selection bias.	45

ACKNOWLEDGMENTS

First and foremost I would like to thank God, my wife Sarah, and my daughters, Eva and Elena, for being by my side and supporting me through the completion of my Master's of Science in Modeling of Virtual Environments and Simulations.

I also owe a huge debt of gratitude to Dr. Quinn Kennedy and LTC Jon Alt for their unwavering support and assistance. Without their keen insight and knowledge this thesis would not have been possible. Thank you for seeing me through to the finish.

I am forever in debt of MAJ Moten Cardy for his invaluable assistance in the scripting and writing of the R Studio code that performed the analysis portion of this thesis. Without his innovation and ability to instruct, I would have been totally lost.

A special thank you is also in order for all the Modeling and Simulations professors for their instrumental instruction and tireless preparation that made this thesis possible. Your dedication to educating military members is unmatched and much appreciated.

THIS PAGE INTENTIONALLY LEFT BLANK

I. INTRODUCTION

A. BACKGROUND

Understanding optimal decision-making is an extremely complex task, but one that the military is currently trying to accomplish. The focus on decision-making is being renewed in an effort to not only understand the processes involved in decision-making, but also improve decision-making among service members. The goal of improving effective decision-making is to increase the combat effectiveness of the military. The last 14 years of combat operations in Afghanistan and Iraq have illustrated the necessity for military leaders to be adaptable, agile, and able to operate in a threat environment that spans irregular and regular warfare, terrorist activity, and at times even governance (Lopez, 2011). The combat environment has always been complex; however, in a non-conventional environment (irregular warfare), that complexity is increased exponentially. The recent and ongoing conflicts in Iraq and Afghanistan illustrate the importance of developing leaders with the cognitive flexibility to learn from feedback from their environment to improve decision performance. In these two conflicts leaders sometimes drew false conclusions about the effectiveness of their operations by attending to historically used measures of performance, such as enemy attrition. From personal experience, a lot of confusion occurred when high enemy body counts were not associated with victory or decreased violence. There was an inability to recognize through trial and error and reinforcement learning that the current approach was not successful. A lot of reinforcement of failure occurred, because of this lack of understanding. Had the military understood optimal decision-making better, this reinforcement of failure could have possibly been avoided by making the decision maker more adaptable, agile, and aware of the complex nuances of the counter-insurgency environment.

The military is in an ideal position to evaluate decision-making among current service members who have spent the last eleven years engaged in combat operations in Iraq and Afghanistan. With this wealth of combat

knowledge contained within current active duty service members, the military can glean decision-making patterns from experienced decision makers. These patterns can then be analyzed in order to better understand how experienced decision makers arrive at optimal or near-optimal decisions. Once this process is understood, then the military can (1) improve combat effectiveness by developing programs to improve decision making among its current leaders and (2) instruct future leaders on optimal decision making to improve their leadership potential.

The primary goal of understanding optimal decision-making is to develop training aids to instruct naïve service members in an effort to shorten the experiential knowledge required to develop effective decision-making practices in combat. Another goal of these training aids is to provide the instructor with insight into the trainee's decision-making process. Such training aids would benefit instructor to trainee interaction and provide insight on timing and type of intervention required by the instructor.

Kennedy, Nesbitt, and Alt (2014) developed a training intervention model called Cognitive Alignment with Performance Targeted Training Intervention (CAPTTIM). This model seeks to determine if a trainee's cognitive state is aligned or misaligned with their actual performance. The model utilizes latency in decision-making to determine the trainee's cognitive state; however, no "generic" metric for determining actual performance has been researched. This thesis seeks to determine an appropriate threshold that delineates between high and low regret. Determining a threshold between high and low regret is an essential step before the model can be tested.

B. REINFORCEMENT LEARNING IS NECESSARY TO REACH OPTIMAL DECISION-MAKING

One cognitive characteristic necessary for military personnel to reach optimal decision-making is reinforcement learning, the ability to learn from trial and error (Sutton & Barto, 1998). Reinforcement learning is necessary when there is a high degree of uncertainty. High levels of uncertainty are associated

with combat operations and environments, in which limited intelligence is known about the situation, but high stake decisions still have to be made. In these situations the military leader makes a “best guess” decision based on experience and training. Current reinforcement learning tests, which are typically computerized laboratory tests, do not completely capture the stressors, uncertainty, and high risk conditions of decisions made in combat (Nesbitt, Kennedy, & Alt, 2015). For example, the Iowa Gambling Task (IGT) (Bechara, Damasio, Damasio, & Anderson, 1994), a very common test of reinforcement learning that has been used in hundreds of psychology studies (Krain, Wilson, Arbuckle, & Castellanos 2006), entails selecting cards from four different decks in a low stress, low stakes, game playing environment. This shortfall has led to the need to create realistic military scenarios and simple wargames that elicit reinforcement learning (Nesbitt et al., 2013). Therefore, Kennedy et al (2014) modified the IGT to mirror a military environment.

1. The Iowa Gambling Task

The IGT is a well-known psychology task that elicits reinforcement learning (Bechara et al., 1994) and has been used in hundreds of studies (Krain et al., 2006). Subjects are given a loan of \$2,000, presented four decks of cards (decks A-D) face down, and asked to make selections that result in maximizing profit. Figure 1 shows a screen shot of the IGT setup.

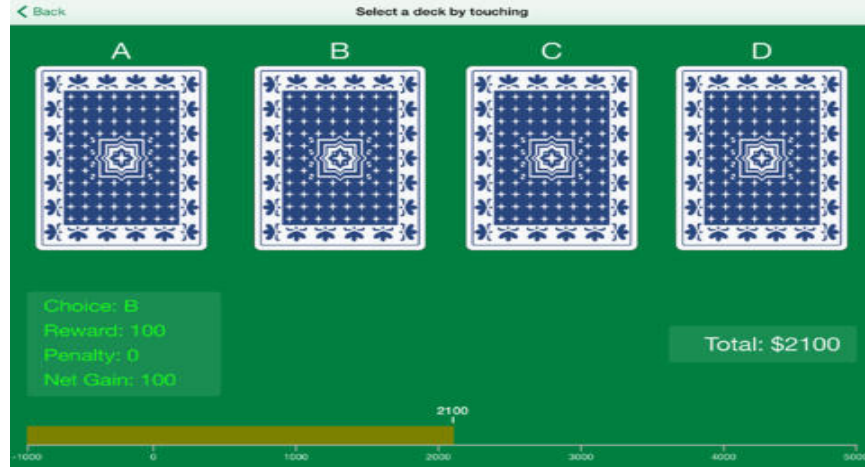


Figure 1. The Iowa Gambling Task screenshot (from Sacchi, 2014).

Each deck has a scheduled dollar payout and penalties that the subject receives depending on their deck selection. The payout amount as well as the severity and frequency of the penalty, differs from deck to deck. Subjects can change the order of their selection at any time and can choose solely from a single deck if they so desire. Through reinforcement learning, healthy subjects eventually discover that decks A and B result in long term losses, despite having higher initial payouts (Bechara et al., 1994). They then realize that, despite lower initial payouts, decks C and D result in long-term gains. Performance is measured by total money won and advantageous selection bias. Advantageous selection bias is calculated by subtracting the number of poor decisions (decks A and B) from the number of good decisions (decks C and D).

Appendix A lists the payout schedule for each deck over the 100 trials. It is important to note that the payout schedule does not reset after each card selection. Until a subject selects a particular deck, the payout for that deck remains the same. For example, Deck B has a negative 1250 penalty every tenth turn but the highest payouts otherwise; the subject cannot game the system by choosing Deck B nine times, but a different deck on the tenth turn, return to Deck B on the 11th turn in an attempt to avoid the negative 1250 penalty.

2. Convoy Task

The IGT was modified into the convoy task to reflect the risks and scenarios faced in a military environment, while mirroring the reinforcement learning elicited by the IGT. In the convoy task each subject selects a route on which to send a convoy and is given a choice between four different convoy routes. The task entails 200 trials of these decisions. At the end of each trial the subject is given immediate feedback with three separate pieces of information: a reward, a penalty, and a running total (Nesbitt et al., 2013). The reward is called *Damage to Enemy Forces*, the penalty is called *Damage to Friendly Forces*, and the running total is called *Total Damage* (Nesbitt et al., 2013). *Damage to Friendly Forces* is analogous to a loss of money in the IGT, while *Damage to Enemy Forces* is analogous to a gain of money. *Total Damage* is analogous to the loan amount and winnings in the IGT. The convoy route selection task's feedback values were adopted from the original IGT payout schedule (see Appendix A). Subjects are instructed that their goal is to maximize the total damage score by minimizing friendly damage and maximizing enemy damage. Like the IGT, subjects should learn through reinforcement learning that routes one and two are bad and routes three and four are good. Data collected from the 34 subjects who participated in the convoy task confirmed that it elicits reinforcement learning (Kennedy et al., 2014).

3. Cognitive Alignment with Performance Targeted Training Intervention

In analyzing data from the 34 subjects that participated in the convoy route task, Kennedy et al. (2015) developed a training intervention model called Cognitive Alignment with Performance Targeted Training Intervention (CAPTTIM) (see Figure 2). This model determines whether a person's cognitive state is aligned or misaligned with actual performance. The model delineates two cognitive states, exploration and exploitation. Exploration is defined as naïve decision-making, in which a person is seeking to further their understanding of the environment by gathering information. Exploitation is defined as experienced

decision-making, in which a person believes that they have attained enough information to begin acting upon that knowledge. The model quantitatively characterizes exploration and exploitation by variability in latency times on making each decision (Fricker, 2010). A standard deviation for each subject was calculated utilizing only the latency times on their decisions that resulted in no damage. Variability greater than twice the subject's standard deviation is considered exploration, whereas variability less than twice the standard deviation is considered exploitation. However, changes in latency time variability provided no measure of actual performance for the individual.

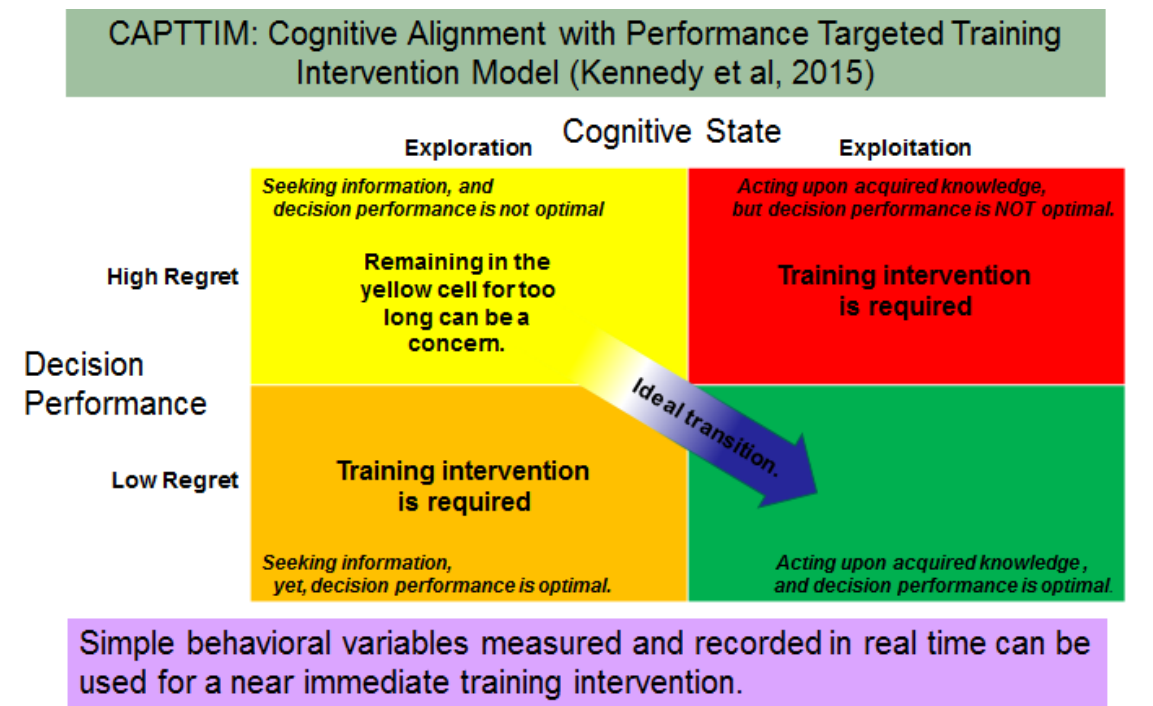


Figure 2. The combination of cognitive state and actual decision performance indicates whether a trainee's cognitive state is aligned or misaligned with actual performance. When misalignment occurs, it indicates the need for a training intervention (from Kennedy, 2015).

Actual performance is measured by regret. Regret is quantified as the difference between the maximum possible payout for a particular trial, and the actual received payout for a particular trial (Agrawal, 1995). Because the payout

schedule is consistent from individual to individual, their deviation from the optimum path can be measured. However, a threshold delineating high from low regret has not been calculated yet.

The convoy route task has a specific sequence of payouts, providing the ability to know at any point in the sequence of trials which route provides the most advantageous reward (Nesbitt et al., 2015). Because the best reward is known, it is possible to calculate the difference between the best reward and the subject's received reward at that specific trial in the convoy route selection task. This difference is defined as regret.

Regret is an absolute performance metric that provides the ability to compare actual performance of the subject with their cognitive state. If the subject's performance is misaligned with their cognitive state then the instructor can intervene and make the appropriate correction. This is very similar to Type I and Type II error from statistics. The subject's performance can be correctly aligned with their cognitive state, which is the ideal transition that is captured in CAPTTIM. Otherwise the subject is making incorrect exploitation decisions believing them to be correct (false positive), or they are making the correct decision, but do not know that they are making the correct decision (false negative). Either of the latter two options requires instructor intervention. The possibility of being able to align a trainee's cognitive state with actual performance is consistent with what the military is trying to accomplish in their pursuit of understanding optimal decision-making.

C. REGRET

Regret is used in numerous fields ranging from computer science, machine learning, and even the medical field. It is very easily applied to scenarios, like the IGT, where the optimum decision is known. For the medical field it is applied retrospectively to describe the diagnosis or misdiagnosis of patients (Djulfbegovic, Elqayam, Reljic, Hozo, Miladinovic, Tsalatsanis, Kumar, Beckstead, Taylor, & Cannon-Bowers, 2014). An interesting application from this

publication that directly relates to the research question of this thesis is how much regret affects future decisions (Djulbegovic et al., 2014).

The defining principle of regret is that if you minimize regret, then you are converging on the correct decision, or for multi-arm bandit scenarios, the correct slot machine (Agrawal, 1995). This principle will be directly applied to this thesis to determine a subject's performance and determine if their performance is aligned or misaligned with their cognitive state. In layman's terms, is the subject making the right decision ignorantly, making the wrong decision thinking it is the correct decision, or do they transition correctly?

Most utilization of the principle of regret has been on analyzing its impact on decision-making or convergence on a decision in a multi arm bandit scenario. No articles could be found that discussed using regret as a method of measuring performance in the way that it is being proposed in this thesis. Other papers use regret as an additional factor in an expected utility function in an attempt to explain behaviors and choices (Bell, 1982).

Bell gives an illustrative anecdotal example of regret. He describes a farmer who has a field of crops that are not yet ready to be harvested. A buyer approaches the farmer and offers him five dollars a bushel for his produce. The farmer knows that, depending on the harvest, his produce could sell for as much as seven dollars a bushel or as little as three dollars a bushel. The farmer is faced with two potential forms of regret: (1) where he accepts the five-dollar-a-bushel offer and the harvest yields a seven-dollar-a-bushel product, (2) he refuses the five-dollar-a-bushel offer and the harvest yields a three-dollar-a-bushel product. Bell then describes how these two forms of regret have very different effects on differing subjects. For some subjects, the fear of losing two dollars per bushel, in the event of an inferior crop, influences their decision much more than the possibility of gaining an extra two dollars per bushel (Bell, 1982). Bell then highlights this phenomenon later on in his paper, when he discusses the utility function. In this example, he discusses how a person might "feel" greater regret between an outcome of \$1,000 and \$2,000 than an outcome of

\$1,000,000 and \$1,001,000, despite the fact that both gained or lost \$1,000 (Bell, 1982). He discusses how the increment is not “felt” the same between both outcomes (Bell, 1982). Bell (1982) additionally made the following comment that is applicable to this thesis and could possibly explain decisions made by subjects: “At an extreme, a decision maker who has severe problems with regret may sometimes prefer to have only a single alternative offered than a choice among two or more” (p. 969). This idea could possibly explain certain subjects’ behavior and their decision to only select certain routes, rather than exploring all options.

Bell additionally looked at regret to explain behaviors and gives anecdotal examples in the realm of insurance and gambling. “The consequence with the largest regret is that in which you choose not to bet, but hear that you would have won” (Bell, 1982, p. 971). If an individual decides not to bet on the horse with long odds, he or she experiences a high amount of regret if that horse wins (Bell, 1982). If you bet on the same lottery number for an extended period of time, the thought of that being the winning number as soon as you stop choosing it could be strong enough to encourage you to continue gambling (Bell, 1982). Bell argues that regret can be used to justify risk-prone behavior (gambling) and risk-averse behavior (purchasing insurance) on the part of the same decision maker (Bell, 1982). For risk-averse behavior, subjects are willing to accept the regret associated with paying for insurance, but never making a claim (Bell, 1982).

Regret is an effective performance metric in tasks in which the payout or reward is known for each decision. For this reason, it is a common performance metric used in gambling scenarios, specifically with multi-arm bandit gambling scenarios (Nesbitt et al., 2015). In these scenarios, the optimum path can be determined. Deviations from this optimum path can be quantified by this notion of regret. We now provide an example of how regret is calculated in a scenario in which the optimum path can be determined—the convoy task payout schedule (Figure 3). In this excerpt, if a subject chooses Route 4 on trial 1, their regret will be $100 - 50 = 50$, because the optimum choice was either Route 1 or Route 2.

If the subject chooses Route 4 again on trial 2, their regret will be $100 - (-250) = 350$, because the optimum choice was still either Route 1 or Route 2. If the subject chooses Route 2 on trial 3, their regret will be $100 - 100 = 0$, because Route 2 was one of the optimum choices. If by trial 9 all routes have been selected exactly twice and the subject chooses Route 2, their regret will be $0 - (-1250) = 1250$, because the optimum choice was Route 4 with a payout of zero. Another key note to make about this payout schedule is that the payout does not redistribute after each selection. The columns can be viewed as a stack where each payout choice remains at the top until chosen. For example, from the schedule below in Figure 3, if a subject does not choose Route 1 until trial 6, their payout would still be 100.

Route 1	Route 2	Route 3	Route 4	Subject's Selection	Regret
100	100	50	50	Trial 1: Route 4	$100 - 50 = 50$
-350	0	-50	-250	Trial 2: Route 4	$100 - (-250) = 350$
-250	-1250	-50	0	Trial 3: Route 2	$100 - 100 = 0$
0	0	0	0		
-200	0	-50	0		
0	0	0	0		
-300	0	-50	0		

Figure 3. Payout schedule excerpt. The blue cell indicates the optimal decision; the yellow cell shows the subject's selection on trial 1; the green cell indicates the subject's selection on trial 2.

D. THESIS GOALS

This thesis has four objectives: (1) find a threshold that delineates between high and low regret (decision performance), (2) combine the decision performance data with the cognitive state data, (3) validate these results and CAPTTIM, and (4) develop a visualization method for displaying a subject's CAPTTIM category on a trial-by-trial basis. A superficial analysis of regret, from the previously collected data, showed that it was consistent with subject's actual performance, as measured by total damage score. Subjects that identified the convoy route with the optimal long term result had a decreasing amount of regret

(Nesbitt et al., 2015). If a threshold for regret is validated, then the utility of CAPTTIM can be tested with other military tasks. CAPTTIM has the potential to provide the instructor with real time guidance on type and timing of intervention in a training scenario.

THIS PAGE INTENTIONALLY LEFT BLANK

II. METHODS

The data used in the analysis portion of this thesis was previously collected from the convoy task and de-identified. This chapter will list in detail the tools and methods used to analyze the regret data in an effort to delineate a threshold between high and low regret. These methods were initially tested (i.e., piloted) on a randomly selected subset of eight of the 34 participants who completed the convoy task. Data from the remaining 26 participants would be used to test the final, selected method. An iterative process was conducted to find an appropriate method, in which initially selected methods informed and directed the subsequent methods. As a result, all the methods described below are more or less in chronological order (exponentially weighted moving average, simple moving average, $\bar{x}$ bar control chart, change point analysis).

A. STATISTICAL SOFTWARE: R STUDIO

The programming language R (R Development Core Team, 2008), which was developed for statistical computing, was utilized for the analysis of the regret data collected from the convoy task (Nesbitt et al., 2013). All the code written for this analysis can be viewed in Appendix B. R-Studio, the integrated development environment (IDE) that was developed for the R language, was used to develop the code that analyzed the regret data. R-Studio is an open source IDE that allows the user to code line by line the exact code for statistics equations. R-Studio varies from a statistics program like JMP in that it requires the user to understand and program every function rather than operating in a drag and drop type fashion like JMP.

B. METHODS USED TO DELINEATE HIGH AND LOW REGRET

Each of the following methods used to research a threshold delineating between high and low regret were coded and calculated in R Studio. Once an analysis was conducted with a specific method, the research team was briefed on the results. This collaboration led to the rejection of three of the four methods utilized to distinguish a regret threshold.

The following sections will chronologically list each of the four methods that were researched. A thorough explanation of each method and how it was used in an attempt to delineate between high and low regret will be given. Additionally, the shortfalls of the first three methods to delineate between high and low regret will be explained.

1. Exponentially Weighted Moving Average (EWMA)

The following section will give a brief introduction of the EWMA equation and its common uses. The next section will discuss how the EWMA was used to analyze the data collected for this thesis. This was the first method explored in an effort to find a threshold to delineate decision performance (high versus low regret).

a. Explanation of EWMA Equation and Uses

“The Exponential Weighted Moving Average (EWMA) chart is used for monitoring process by averaging the data in a way that give less weight to old data as samples are taken and gives more weight to most recent data” (Braithwaite, Osanaiye, Omaku, Saheed, and Eshimokhai, 2014, p. 1). EWMA also is very effective at detecting minor changes in the process mean (Braithwaite et al., 2014). It was originally developed by S. W. Roberts in 1959 as a means of monitoring control/performance charts in industrial processes (Braithwaite et al., 2014). It also has been very useful in time series analysis and forecasting (Braithwaite et al., 2014). The following is how an individual EWMA value is calculated as

$$Z_i = \lambda X_i + (1 - \lambda) Z_{i-1} ,$$

where Z_i is the EWMA control statistic, λ is the weighted parameter, and X_i is the actual observed data value

A key difference between EWMA and a simple moving average is that EWMA considers all previous data points, while a simple moving average only considers data points within a specified window (Braimah et al., 2014). “EWMA weights samples in geometrically decreasing order so that the most recent samples are weighted most highly while the most distant samples contribute very little” (Braimah et al., 2014, p. 2). This weighted parameter, λ ($0 < \lambda \leq 1$), is a mathematical representation of how heavily memory of past data is relied upon (Kalgonda, Koshti, and Ashokan, 2011). As λ increases from zero to one, more weight is placed on recent data points and less weight is placed on distant data points. If $\lambda = 1$, then 100 percent of the weight is placed on the most recent data point and no weight is placed on the past (Kalgonda et al., 2011). The sensitivity of the EWMA to small shifts in the process mean is reliant upon the value of λ (Kalgonda et al., 2011).

The use of EWMA as a means of detecting changes in regret was based on the EWMA’s sensitivity to small shifts and reliance on memory. Because decisions on the convoy task rely heavily upon working memory and the influence of past decisions on future decisions (Kennedy et al., 2013), this method of averaging regret seemed more appropriate than a simple moving average.

Using EWMA to analyze regret was the initial approach taken because it worked exceptionally well in characterizing subject’s cognitive state based on decision time latencies in the convoy task. An effective threshold delineating between the cognitive states of exploration and exploitation was applied to this EWMA and accurately portrayed subject’s transition between these two states.

The threshold that was used was double the standard deviation of each subject's latency times in decisions that resulted in low damage. The EWMA equation for time latency utilized a λ value of 0.1. This λ value means that subjects had a heavy reliance on past decisions, since $(1 - \lambda)$ determines the weight placed on past data points. This code was modified to analyze regret and utilized the same value of λ .

b. EWMA of Regret

The initial EWMA of regret looked at the mean values of regret. This meant that the EWMA was looking at the cumulative regret divided by the number of trials. This analysis produced some interesting results. However, upon further discussion with the research team and additional analysis, the use of the mean regret as the values on which to conduct the EWMA was determined to be incorrect. By using mean regret the values were essentially being smoothed twice. Dividing the cumulative regret by the trial was taking an average after every trial; this average was again being averaged with the EWMA based on the weight placed on past data. This realization led to the decision that the EWMA should be conducted on the regret per trial for each subject.

By using the regret received by the subject at each trial, the EWMA was looking at actual values and not an already averaged value. The result was much more volatile changes in the EWMA.

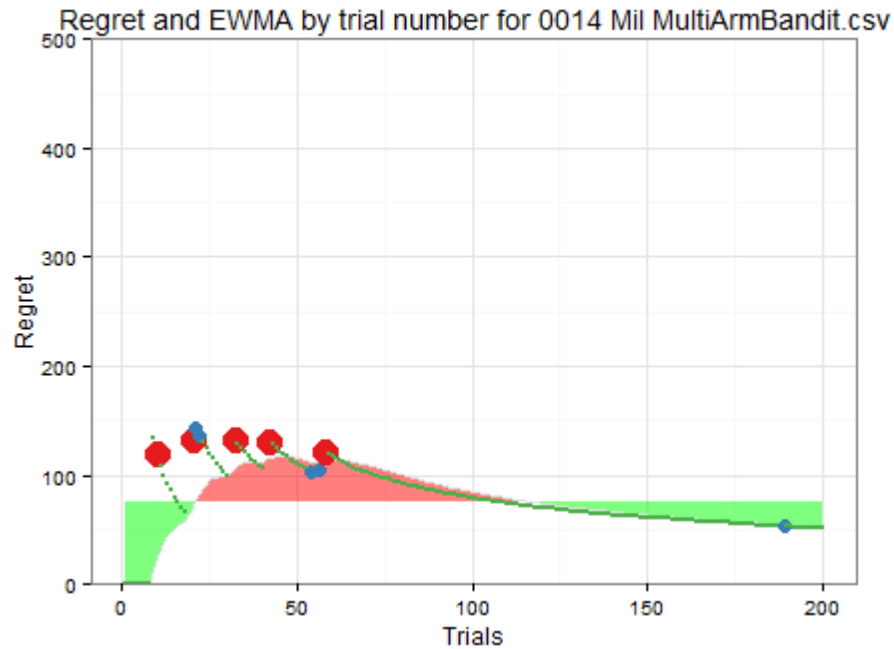


Figure 4. EWMA of regret for Subject 14 using mean regret. Mean regret proved to be inappropriate as it was performing a EWMA on an already averaged regret value. This accounted for the much less volatile spikes in regret value. The large red dots are high damage instances. The medium blue dots are medium damages, and the small green dots are low damage instances. The shaded red area is where the EWMA is above the threshold and the shaded green area is where the EWMA is below the threshold. The threshold is calculated as 0.5 times the standard deviation of the mean regret.

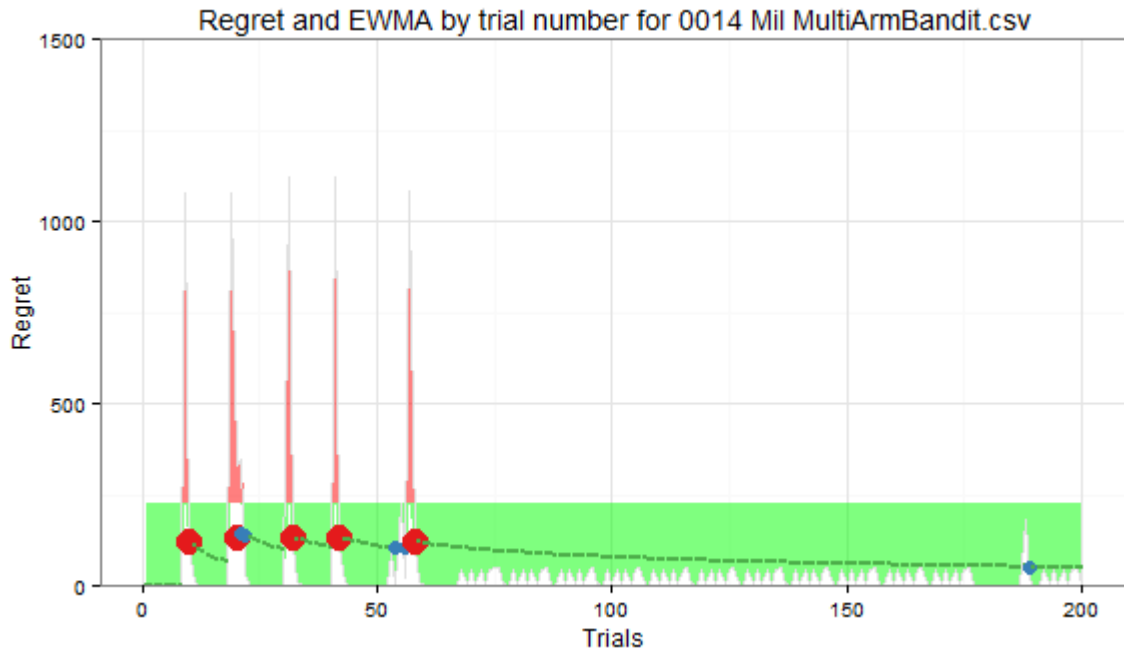


Figure 5. EWMA of regret for Subject 14 using regret received at each trial. The volatility in high regret is seen with the sharp red peaks which is where regret reaches values of 1250 for high friendly damage. The red, blue, and green dots are for high, medium, and low damages respectively. These dots are plotted along the mean regret line. Shaded red areas are above the threshold, while shaded green areas are below the threshold. The threshold is defined as the standard deviation of the regret received per trial.

The threshold value for the EWMA conducted on mean regret had to be adjusted to one half the standard deviation of regret in order to have the EWMA fall above and below the threshold, as can be seen in Figure 4. This adjustment was as a result of averaging an already averaged value. The threshold for the EWMA conducted on regret received per trial was strictly the standard deviation of the regret per trial and did not require any fractional adjustment. After discussion and further analysis with the research team, it was suggested that a sensitivity analysis of λ to the regret per trial data be conducted. Based on the sensitivity analysis the ability to tune λ to the actual data could be achieved.

This sensitivity analysis of regret per trial to λ resulted in the realization of the difficulty of tuning this parameter for this use case. The analysis showed that a λ value of 0.9 achieved the line of best fit for each subject to the actual regret data (this realization is trivial given the EWMA equation). This value of λ illustrated that subjects placed very little weight on past regret and that the immediate results influenced their decision the most. Figure 5 illustrates this point—had Subject 14 weighted past decisions heavily, the spikes in regret would have become less volatile and been spread across future decisions, illustrating that he/she had been influenced by the previous decision.

Thus, this EWMA was fit to the actual regret per trial data and led to highly volatile changes in regret. Despite a defined payout schedule, values of regret are very random across subjects with a wide range of possible values. For example, one subject may have only experienced regret values of 50 if they converged on the optimal path, while another subject may have experienced regret values of 1250 since they did not converge on the optimal path. The high volatility of these values made defining a single threshold difficult, since regret could range from 0 to 1250. This issue made it difficult to classify into which category of the CAPTTIM model a subject should be categorized. Therefore, other approaches were sought. The next method examined was the simple moving average.

2. Simple Moving Average

Rather than looking at a trial by trial analysis of whether regret was increasing or decreasing, a simple moving average was conducted to “block” regret by a specific number of trials. As a reminder, simple moving average differs from EWMA in that it only considers the data within a specific window, whereas the EWMA considers all data points and weights them according to the value of λ . Two approaches were taken: (1) the simple moving average looked at a moving window of five trials throughout the 200 trials of regret data (2) the simple moving average did the exact same calculation with a moving window of

10 trials. The moving window of five trials allowed for more granularity in observing this subject's changes in regret. Utilizing a larger window gives less blocks to analyze changes in regret and thus does not provide as much sensitivity for changes in regret (see Figures 6 and 7). As a result, the simple moving average that utilized a window of 5 trials was used for the follow on analysis of regret.

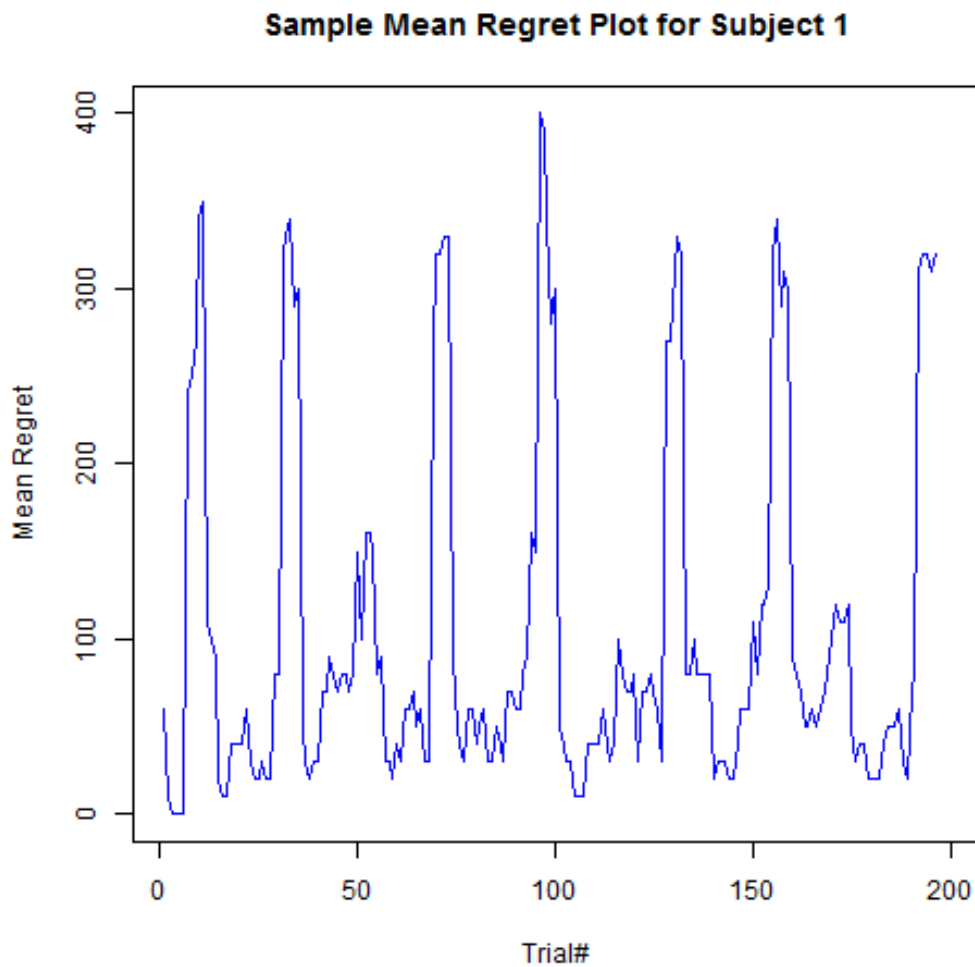


Figure 6. Simple moving average of regret per trial for Subject 1 with a window of 5 trials. The solid blue line shows the averaged regret and how high values in regret influenced the average for the 4 previous and 4 successive trials. Had a simple moving average not been used, high values of regret would be single vertical lines.

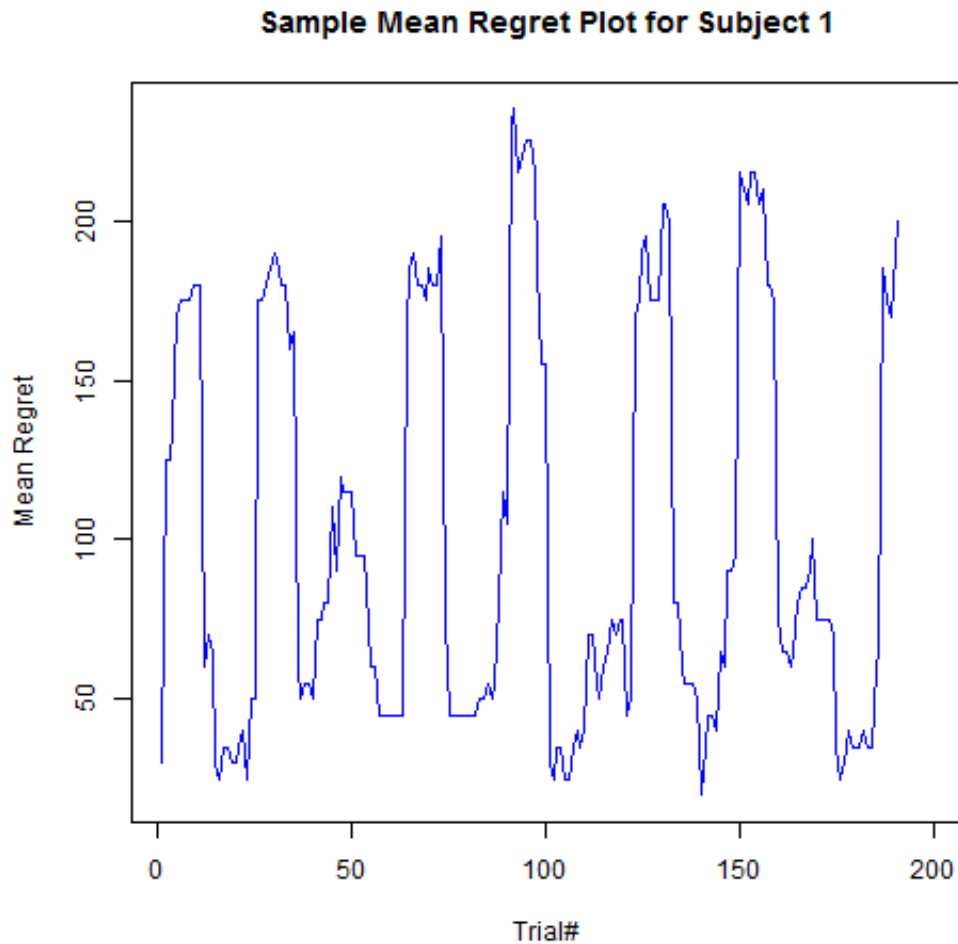


Figure 7. Simple moving average of regret per trial for Subject 1 with a window of 10 trials. The solid blue line shows the averaged regret and how high values in regret influenced the average for the 9 previous and 9 successive trials. Comparison to Figure 6 shows how, for the same subject, the spikes in high regret are broadened by utilizing a larger window.

The use of a simple moving average of regret provided more insight into defining a subject's performance than the EWMA of regret. Because regret for most subjects was extremely random, trying to define a threshold to differentiate between high and low regret using an EWMA was very difficult to do. The simple moving average allowed an analysis of discrete blocks to determine the slope of the line, which in turn showed whether regret was increasing or decreasing at

specific points. However, as described in the section below, it was discovered that the simple moving average method also had drawbacks.

3. X-Bar Control Chart

Instead of looking at a simple moving average of regret and applying a threshold that delineated between high and low regret, a better approach could be to create a control chart that defines a median and an upper control limit. As long as the value falls within the upper control limit, the subject is deemed within tolerance or having low regret. The control chart made it a lot easier to classify subjects into their specific category in CAPTTIM. Originally the control chart looked at using the mean of regret per trial plus the standard deviation of regret to define the upper control limit. This upper control limit adjusted utilizing the same 5 trial window that the simple moving average utilized. However, what the research team found was that the mean was not a useful metric for determining the upper control limit of the control chart. This was due to the fact that regret has possible values ranging from 0 to 1250. With such volatility in values, the mean and standard deviation are skewed due to these high spikes in regret experienced by most subjects. Therefore, the upper control limit was falsely classifying subject performance, and as a result very few subjects were being classified as out of tolerance (high regret). In fact, most subjects were being classified as having low regret despite their actual overall performance (final damage score). A histogram of regret was created, in order to illustrate the unsymmetrical characteristic of the regret data (see Figure 8).

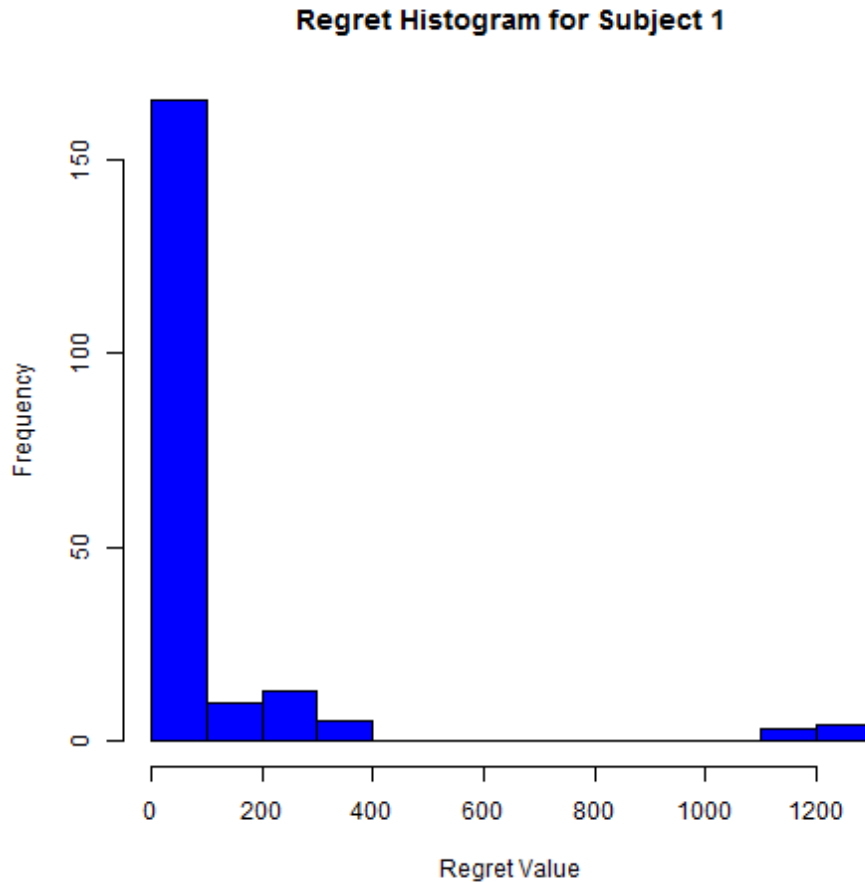


Figure 8. Histogram of regret data for Subject 1. This clearly illustrates that the majority of regret values experienced by Subject 1 are of magnitude 50 and that the high spikes in regret only occurred a handful of times.

Due to the variation in the data for regret, the next approach taken was to look at the median of regret versus the mean. Additionally the research team recommended looking at a window of 20 trials to calculate the median and upper control limit in order to provide a more stable analysis of tolerance. This window of 20 trials was chosen based on the payout schedule and when these large values of regret were incurred. Additionally the window of 20 trials provided an appropriate window in which subjects would be allowed to illustrate reinforcement learning and make mistakes and adjust their course of action. Smaller windows proved to be too restrictive and classify subjects out of tolerance too hastily. The

new upper control limit for the X-Bar chart was then calculated as the median plus the median absolute deviation for the moving window of 20 trials. Figure 9 shows the X Bar control chart for Subject 1. The solid blue line is the simple moving average described before, and the dashed red line is the median plus the median absolute deviation, which is recalculated every 20 trials. Points on the simple moving average that were above the dashed red line are considered out of tolerance (high regret), while points below the red dashed line were considered within tolerance (low regret) (see Figure 9).



Figure 9. X-Bar control chart for Subject 1. The solid blue line is the simple moving average that was previously discussed. The dashed red line is the upper control limit. The upper control limit is defined as the median plus the median absolute deviation and

is recalculated every 20 trials.

4. Change Point Analysis

After discussion with the research team and a recommendation from the team statistician, Dr. Fricker, a change point analysis was conducted to determine the best window size of trials to create the upper control limit for the X-Bar control chart. Change point analysis is useful in determining if a change occurred, how many changes occurred, when the changes occurred, and provides with what confidence the changes occurred (Taylor, 2000). Change point analysis is extremely flexible and can be performed on all types of time ordered data to include, attribute data, non-normal distributions, ill-behaved data, and data with outliers (Wayne, 2000). A key difference between change point analysis and control charts in the context of regret is that control charts can be generated following each individual trial, while a change point analysis can only be generated retrospectively (Wayne, 2000). Change point analysis is typically more sensitive and can often detect changes in the process mean that are missed by the control chart, thus the two methods are best employed in a complimentary fashion (Wayne, 2000).

5. Final Method: Combination of Control Chart and Change Point Analysis

Combining control chart and change point analysis, in this complimentary fashion, is the method being employed in this thesis. The statistical computation language R contains built in packages for conducting change point analysis. The R package utilized in this analysis was the segment neighborhood (SegNeigh) algorithm (Killick, & Eckley, 2014). This algorithm utilizes dynamic programming to calculate the optimal segmentation for $m + 1$ change points and reuses the data calculated for m change points (Killick et al., 2014). This essentially means, that the algorithm searches over all previous change points and chooses the one that results in the optimal segmentation up to that time (Maidstone, Fearnhead, & Letchford, 2013). This package takes a variable Q that specifies the maximum

number of change points to identify. This was useful in the analysis of the non-normal data contained in the data set of regret per trial. Due to the volatility of the regret per trial data, running a change point analysis package that identified every change point was not useful. However, by specifying a smaller number of change points ($Q=15$) the analysis was able to yield results that were useful in delineating between high and low regret. Figure 10 shows the change point analysis performed on Subject 1.

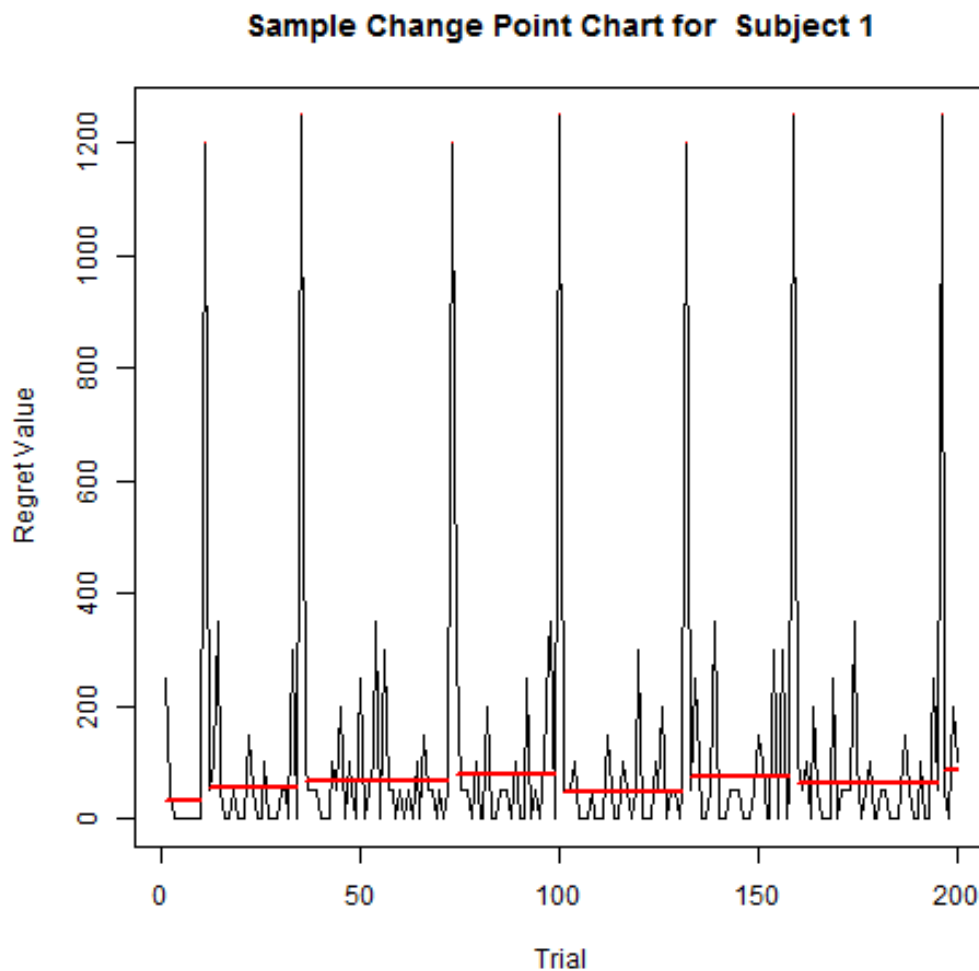


Figure 10. Change point analysis for Subject 1. The solid black line is the regret per trial data. The solid red lines are the process means returned by the change point analysis—they represent the process mean for that range of trails. The large spikes in regret incurred a change in the process mean that spanned the single trial in which the regret was incurred.

After studying the change point analysis and further discussion with the research team, it was decided that, rather than using an X-Bar control chart, creating a box plot of the means associated with each change point and determining if the mean was above or below the median would accurately delineate between high and low regret. Because the change point analysis returns the mean as well as the trial number for each change point, the subject can be accurately categorized in CAPTTIM for a range of trials. This was the final method decided upon for analyzing regret for the subset of 8 subjects along with the subsequent 26 subjects.

In addition to the use of the change point analysis to delineate between high and low regret, the research team decided to add an additional metric for determining decision performance. Subjects that chose route 1 or 2 after trial 100 would be automatically classified as having high regret. This metric took into account the time and duration of the experiment and at which point the optimal performers converged on the ideal decision.

THIS PAGE INTENTIONALLY LEFT BLANK

III. RESULTS

By conducting the change point analysis on all 34 subjects and comparing the resulting means with the median of all change point means, an effective threshold for delineating between high and low regret was established. Once the threshold for delineating between high and low regret was obtained, the data could then be compared with the cognitive state of the subject in order to categorize them in CAPTTIM. This section will detail how each subject's regret was categorized and then compared with the cognitive state data.

A. OVERVIEW OF COGNITIVE STATE DATA DEVELOPED FROM PRIOR RESEARCH

A subject's cognitive state was previously categorized by Maj Pete Nesbitt, who utilized an EWMA of the latency in decision-making times. A threshold was then applied to the EWMA in order to delineate between the cognitive states of exploration and exploitation. The threshold that was utilized was two times the standard deviation of latency in decision-making times immediately following trials that resulted in low damage. It was assumed that decision times after receiving low damage would be relatively fast, and therefore, could be used to determine an individual subject's baseline latency time. In contrast, it was assumed that decision times following trials that resulted in high or medium damage would be longer, because subjects typically reflected on the negative feedback. The threshold was specific to each subject since it was calculated using their baseline. This threshold accurately delineated between exploration and exploitation for all 34 subjects. This prior work allowed the research team to know on a trial-by-trial basis whether the subject was exploring or exploiting (see Figure 11). This knowledge was crucial in the development of the CAPTTIM categorization algorithm.

Most subjects illustrated a pattern of taking longer to make decisions in the beginning of the convoy task when they were exploring and gathering information on the environment (higher latency times between decisions). Most subjects then transitioned to making decisions more rapidly (lower latency times between decisions) once they believed that they had converged on the correct choice and were exploiting that path. This pattern can easily be seen in Figure 11, where Subject 4 spent approximately 45 trials exploring (shaded orange region) and then transitioned to exploitation (shaded blue region) from trial 45 to 200. As can be seen from Figure 11, even though Subject 4 began exploiting the decision that he/she thought was the correct decision, heavy friendly damages (large red dots) were incurred throughout the remainder of the trials. Because Subject 4 incurred heavy and medium friendly damages throughout the 200 trials, his/her final damage score was much lower than those of subjects who converged on the optimal choice. As a reminder, each subject began the experiment with a positive final damage score of 2000. When they received friendly damage this would deduct from their final damage score and when they inflicted damage on the enemy this would increase their score. The average final damage score across all 34 subjects was 2,402.94. Subject 4's final damage score was 2050.

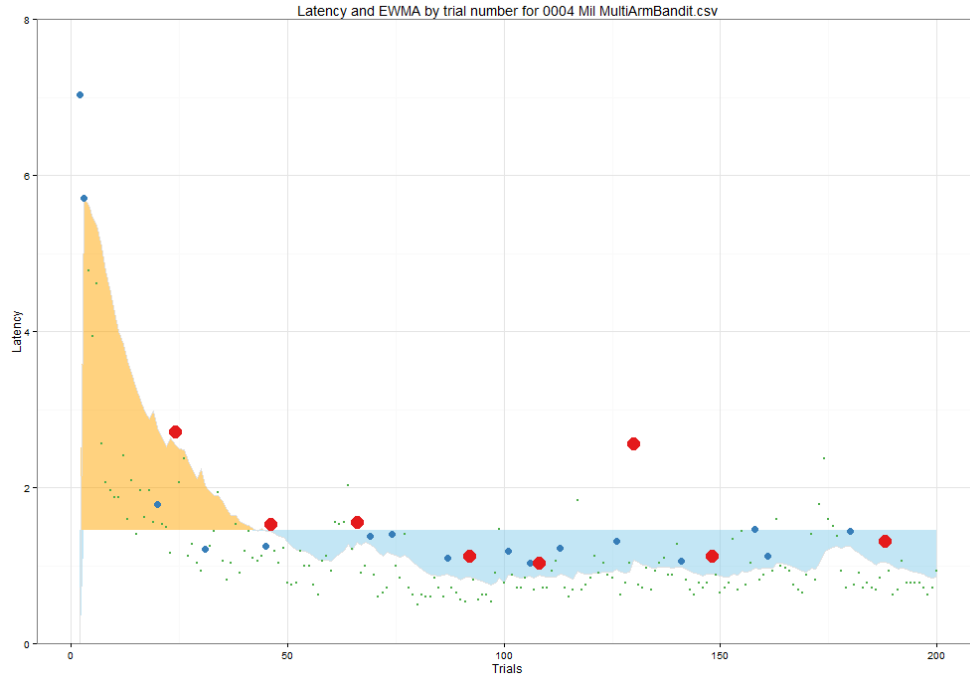


Figure 11. EWMA of latency in decision-making times for Subject 4. The y-axis is latency in decision-making times and the x-axis is the number of trials. The colored dots represent damage incurred and are plotted at the actual latency in decision-making time versus the EWMA. The color and size of the dot is correlated with the level of damage incurred on the preceding trial. Red dots are high damage, blue dots are medium damage, and green dots are low damage. The orange shaded regions are where the EWMA is above the threshold (exploration) and the blue shaded regions are where the EWMA is below the threshold (exploitation).

The following example is of a subject who illustrated optimal exploration of the environment followed by exploitation of the optimal choice. Figure 12 is the EWMA of latency in decision-making times for Subject 14. Subject 14 followed the typical pattern observed for most subjects, by exploring in the beginning (shaded orange region) and then transitioned to exploiting (shaded blue region). Subject 14 transitioned between exploration and exploitation by approximately trial 30. While Subject 14 took some medium damages (medium blue dots) and high damages (large red dots) in the beginning of his/her exploitation phase, he/she eventually converged on the optimal decision and incurred very little

damage throughout the remaining trials. As a result, Subject 14's final damage score was 4700 compared to Subject 4's score of 2050.

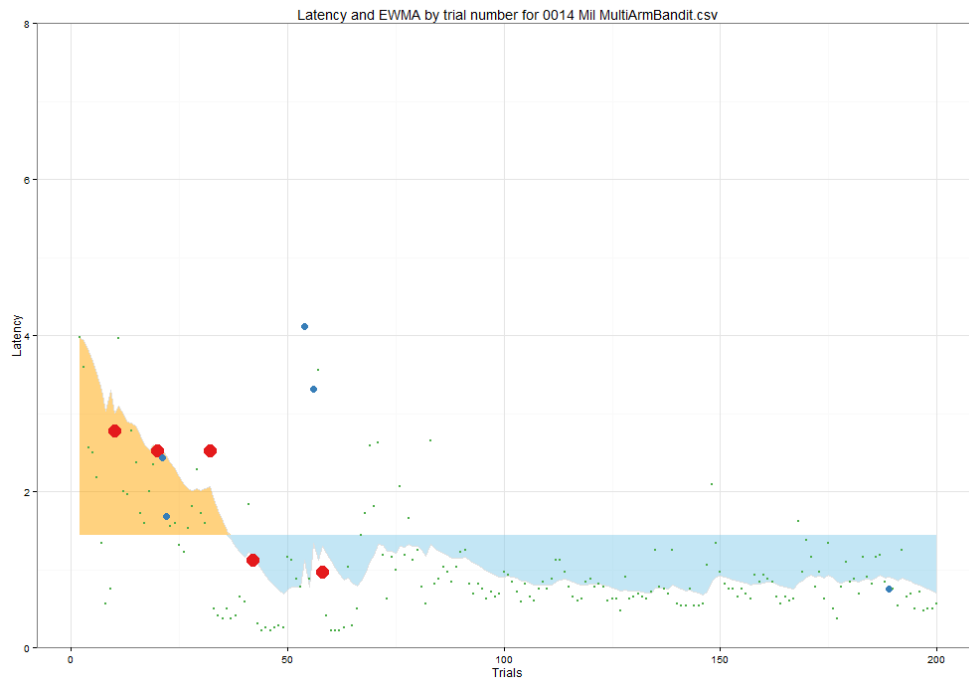


Figure 12. EWMA of latency in decision-making times for Subject 14. The y-axis is latency in decision-making times and the x-axis is the number of trials. The colored dots represent damage incurred and are plotted at the actual latency in decision-making time versus the EWMA. The color and size of the dot is correlated with the level of damage incurred on the previous trial. Red dots are high damage, blue dots are medium damage, and green dots are low damage. The orange shaded regions are where the EWMA is above the threshold (exploration) and the blue shaded regions are where the EWMA is below the threshold (exploitation).

These examples demonstrate that knowing a subject's cognitive state does not provide sufficient insight into their actual decision performance. Subjects 4 and 14 showed similar cognitive state patterns yet had very different decision performance. Thus, the next step was to combine the subject's cognitive states with the categorization of their actual performance (high versus low regret), which was the focus of the research conducted in this thesis.

B. CHANGE POINT ANALYSIS COMBINED WITH COGNITIVE STATE DATA

The cognitive state data from above was then taken and combined with the change point analysis data that delineated between high and low regret. This delineation provided a metric to gauge a subject's actual performance. The combination of actual decision-making performance with cognitive state allowed for the categorization of subjects into CAPTTIM.

1. Delineating High and Low Regret Using Change Point Analysis

Using the change point analysis data, subjects were categorized as having high or low regret on a trial-by-trial basis. The change point analysis returned 15 change points for each of the 34 subjects. These change points represent instances where a subject's process mean changed. The reason that 15 change points were returned was as a result of the method used within R (SegNeigh) to conduct the change point analysis. The number of change points was limited to 15, due to the volatility of the regret data. Regret per trial values vary between 0 and 1250 with intermediate values of 100, 200 and 300. By limiting the number of change points the significant changes were readily identified, while the minor changes were allowed to occur without changing the process mean. If every change point were identified the number of change points would have been too numerous to provide any use for analysis. The change point and its associated process mean were then compared with the median of all 15 process means. This comparison looked at windows of trials on the basis of the process means returned from the change point analysis (see Figure 13). The process mean for that window of trials was then compared with the median of the process means to determine whether it fell above or below the median. If the process mean was above the median, the subject was categorized as having high regret; if the process mean was below the median, the subject was categorized as having low regret. Figure 13 clearly indicates that Subject 4 experienced peaks of high regret throughout his/her 200 trials, which resulted in a much lower final damage score.

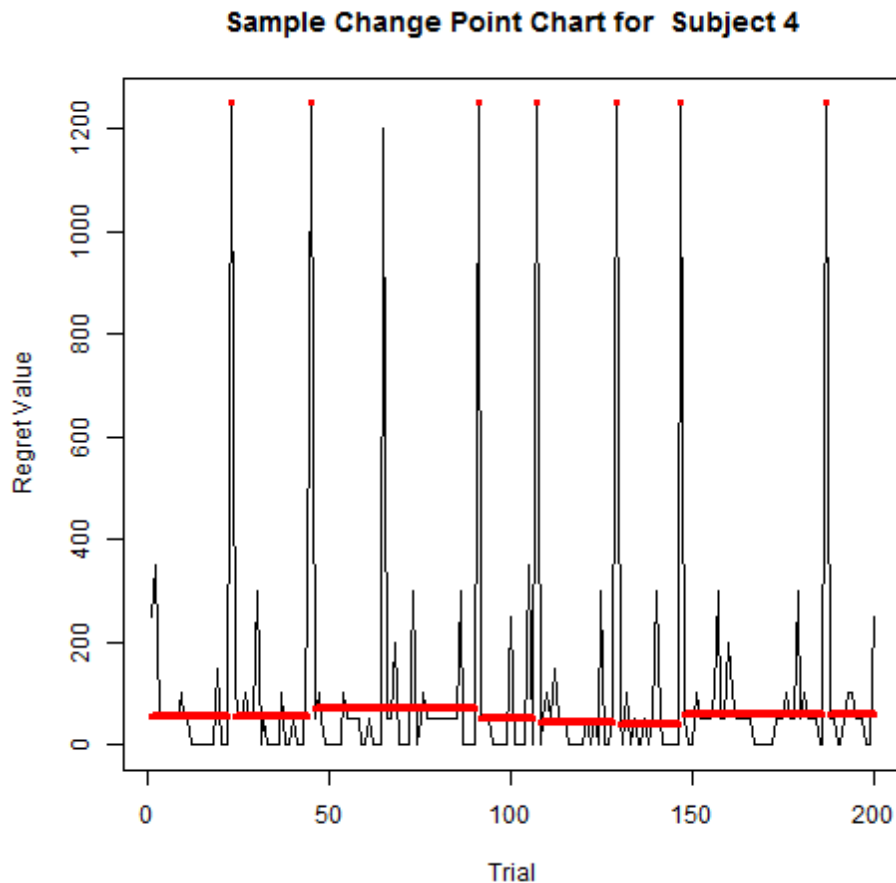


Figure 13. Change point analysis for Subject 4. The y-axis is the regret per trial value, while the x-axis is the trial number. The red lines are the process means returned from the change point analysis. The spikes in the regret value are a result of the subject receiving heavy friendly damage and incurring high regret. These spikes result in a change point that exists over just one trial. The other, longer red lines are where the process mean did not change for that range of trials.

The following information illustrates the change point analysis results for a subject who converged on the optimal choice. Figure 14 is the change point analysis chart for Subject 14. Subject 14 clearly illustrated the ideal exploration phase where heavy damage is expected and encouraged in order for the subject to fully explore the environment and identify the optimal choice. This exploration phase was followed by an ideal exploitation phase, where Subject 14

experienced minimal regret. Because Subject 14 experienced minor regret for the majority of trials, his/her final damage score was much higher than that of Subject 4 (4700 vs. 2050). Another interesting point illustrated by Subject 14, was that he/she experienced numerous change points in the beginning, but after trial 60 (approximately) the process mean remained constant.

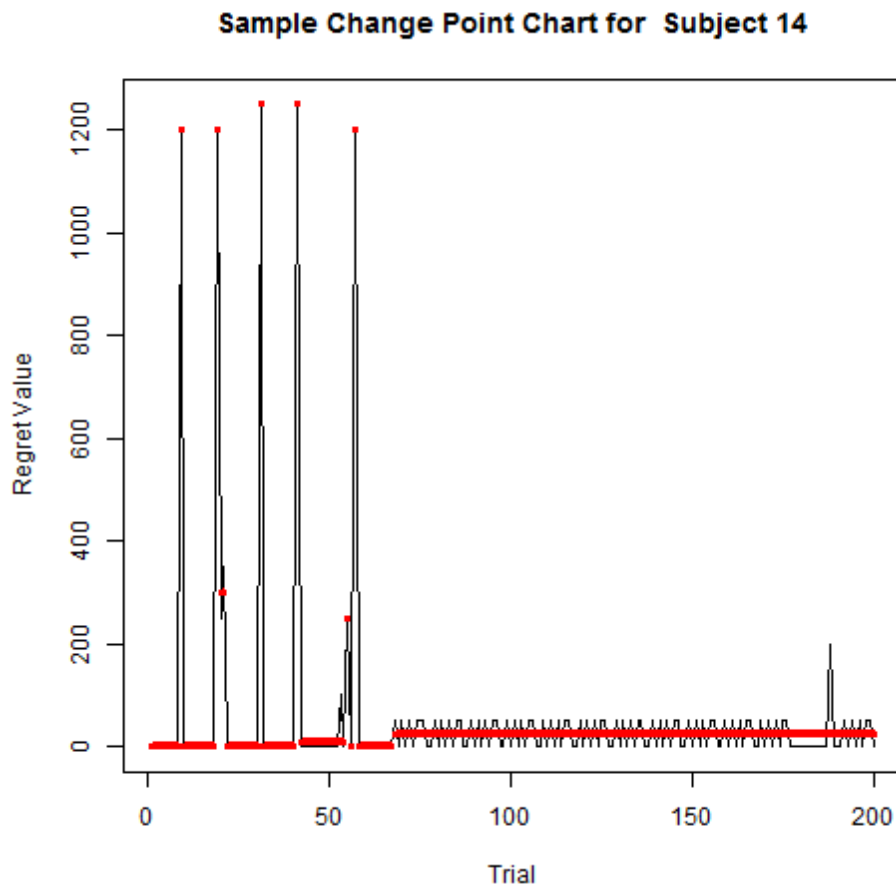


Figure 14. Change point analysis for Subject 14. The y-axis is the regret per trial value, while the x-axis is the trial number. The red lines are the process means returned from the change point analysis. The spikes in the regret value are a result of the subject receiving heavy friendly damage and incurring high regret. These spikes result in a change point that exists over just one trial. The other, longer red lines are where the process mean did not change for that range of trials.

Once a threshold was established that effectively delineated between high and low regret and provided a method for gauging actual decision performance, the research team had all the requisite information required for categorizing subjects within CAPTTIM. This ability to categorize subjects within CAPTTIM fulfilled a primary goal of this thesis.

2. Combining Cognitive State and Decision Performance to Categorize Subjects within CAPTTIM

The combined cognitive state data and decision performance data allowed for the categorization of subjects within CAPTTIM to be accomplished. Figure 15 shows the CAPTTIM categorization algorithm used to properly assign subjects within their appropriate category.

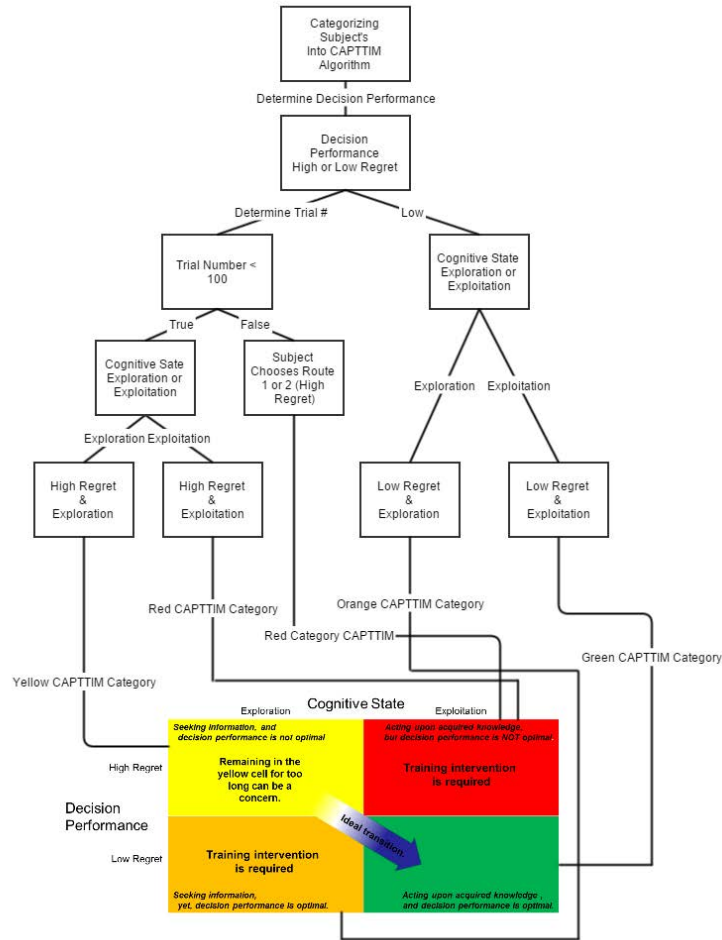


Figure 15. CAPTTIM categorization algorithm. This figure illustrates how each subject is categorized in CAPTTIM based on decision-making performance (measured by regret) and cognitive state (measured by latency in decision-making times).

Because the change point analysis of regret and EWMA of latency in decision-making times delineate between decision performance and cognitive state for a range of trials, a graphical representation was developed that represents what category of CAPTTIM a subject is in on a trial by trial basis. This representation was overlaid on the regret per trial graph in order to illustrate how CAPTTIM could be used to provide instructors information on type and timing of intervention.

Figure 16 is the CAPTTIM categorization chart for Subject 4. Figure 16 clearly shows that Subject 4 experienced high regret at times during his/her exploration phase (yellow block), but never fully explored the entire environment (orange blocks). After a brief exploration phase (approximately 45 trials), Subject 4 transitioned to the exploitation phase. For windows of trials Subject 4 exploited decisions that resulted in low regret (green blocks). However, these windows were often interrupted by exploited decisions that resulted in high regret (red blocks). These repeated exploited decisions with high regret were a clear indicator that Subject 4 did not converge on the optimal choice.

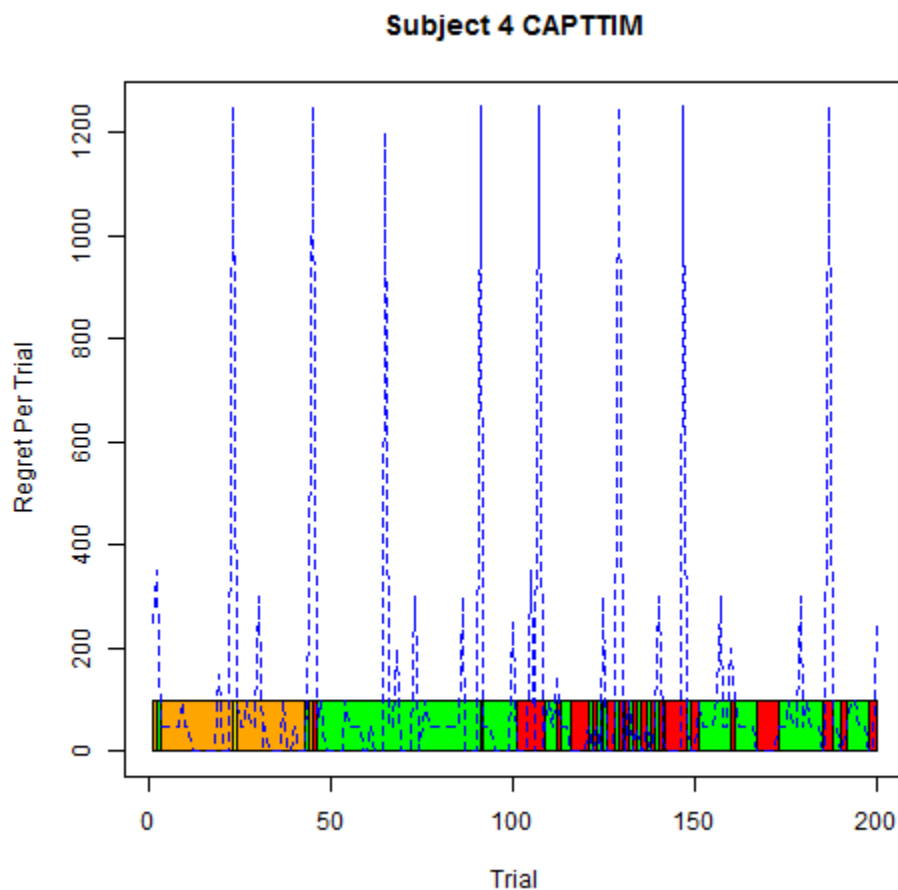


Figure 16. CAPTTIM categorization chart for Subject 4. The color-coded bar at the bottom of the chart correlates to the category color found within the CAPTTIM model. Yellow is high regret and exploration. Orange is low regret and exploration. Red is high regret and exploitation. Green is low regret and exploitation.

Figure 17 is the CAPTTIM categorization chart for Subject 14. This figure accurately portrays that Subject 14 experienced high and low regret during his/her exploration phase (yellow and orange blocks), and even experienced a couple of poor choices during the initial exploitation phase (red blocks). For the vast majority of trials, however, Subject 14 made the ideal transition and converged on the optimal choice (green block) and did not deviate from the optimal choice for the remaining trials.

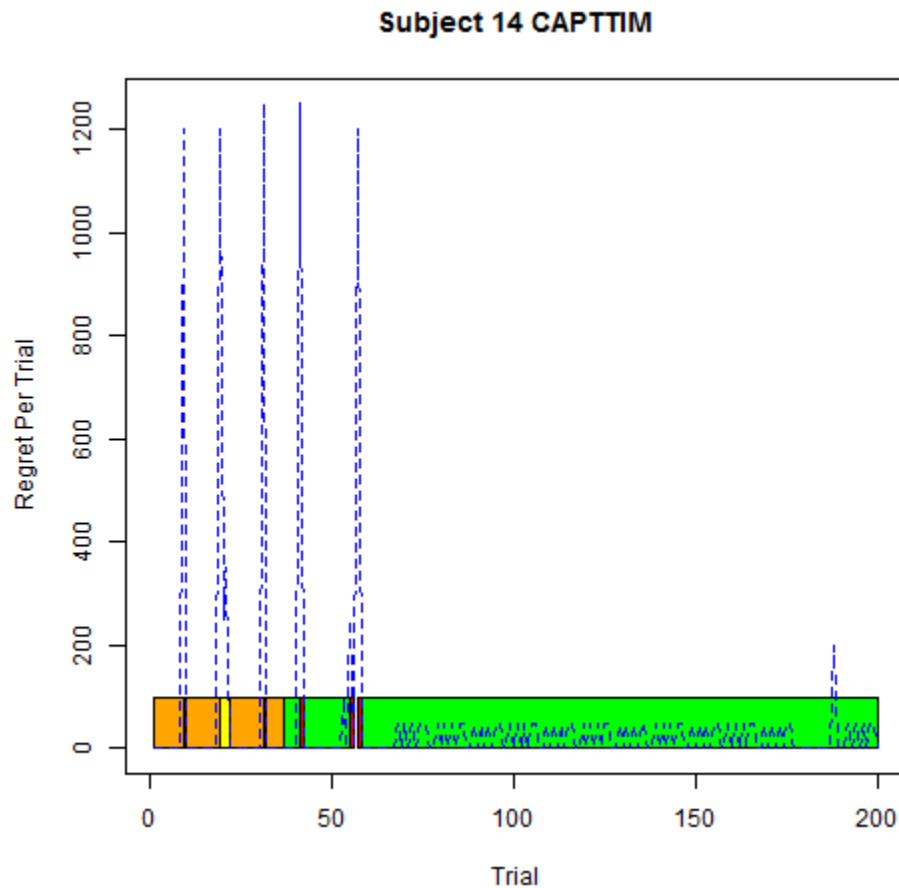


Figure 17. CAPTTIM categorization chart for Subject 14. The color-coded bar at the bottom of the chart correlates to the category color found within the CAPTTIM model. Yellow is high regret and exploration. Orange is low regret and exploration. Red is high regret and exploitation. Green is low regret and exploitation.

The CAPTTIM categorization charts for Subjects 4 and 14 clearly illustrated typical patterns observed across the 34 subjects. Subject 4 illustrated

how the optimal path was never identified and exploited. This decision pattern would have resulted in an instructor intervention based on the CAPTTIM results. Subject 14, however, converged on the optimal choice and exploited. Thus, this decision pattern would have resulted in no instructor intervention being needed. The research team observed that the subjects fell into three typical groups consisting of (1) subjects who explored and eventually identified the optimal choice ($n = 9$), (2) those who explored and exploited non-optimal choices ($n = 21$), and (3) subjects who never transitioned from the exploration cognitive state to the exploitation cognitive state ($n = 4$). This third group would have required instructor intervention, which was accurately identified using the CAPTTIM categorization charts. This third group is illustrated by subject 11 in Figures 18 and 19.

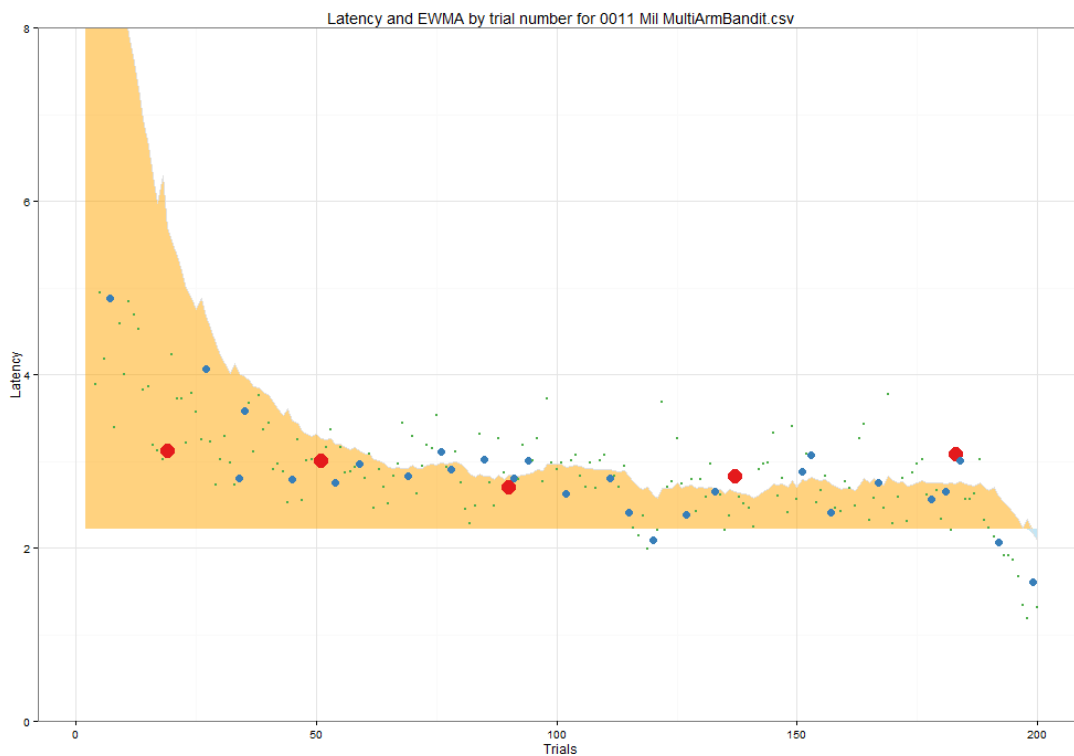


Figure 18. EWMA of latency in decision-making times for Subject 11. The x- and y-axis are the same as the previously described graphs. Note that Subject 11's EWMA of latency in decision-making times never falls below his/her threshold (shaded orange region). This subject spent the entire time exploring the environment and never exploited any decisions.

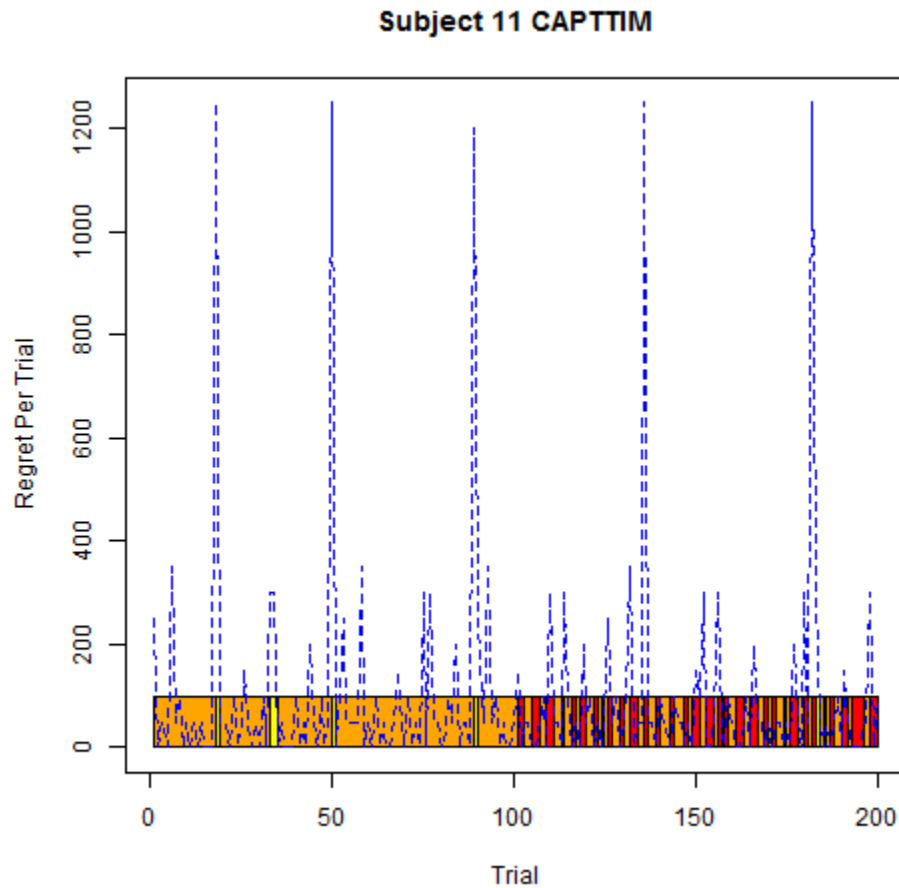


Figure 19. CAPTTIM categorization chart for Subject 11. Note that the values are coded yellow, orange and red. The only reason that Subject 11 was ever categorized as red (high regret and exploitation) within CAPTTIM was due to the fact that subjects are penalized for choosing routes 1 and 2 after trial 100. Subject 11's final damage score was 2200.

Based on the analysis conducted by the research team, the change point analysis of regret provided an accurate delineation between high and low regret. The combination of cognitive state data with the change point analysis in order to generate the CAPTTIM categorization chart is believed to be an effective instructor intervention tool.

C. VALIDATION OF CHANGE POINT ANALYSIS AND COGNITIVE DATA AS CAPTTIM CATEGORIZATION METRICS

All that remained for the research team was to develop a means to validate the effectiveness of using the change point analysis, cognitive state data, and route choice after trial 100. The validation method chosen to validate how well these methods actually categorized subjects within CAPTTIM was a correlation test between number of trials a subject was in the red category and their advantageous selection bias and final damage score. Figures 20 and 21 show the plots for these correlation tests.

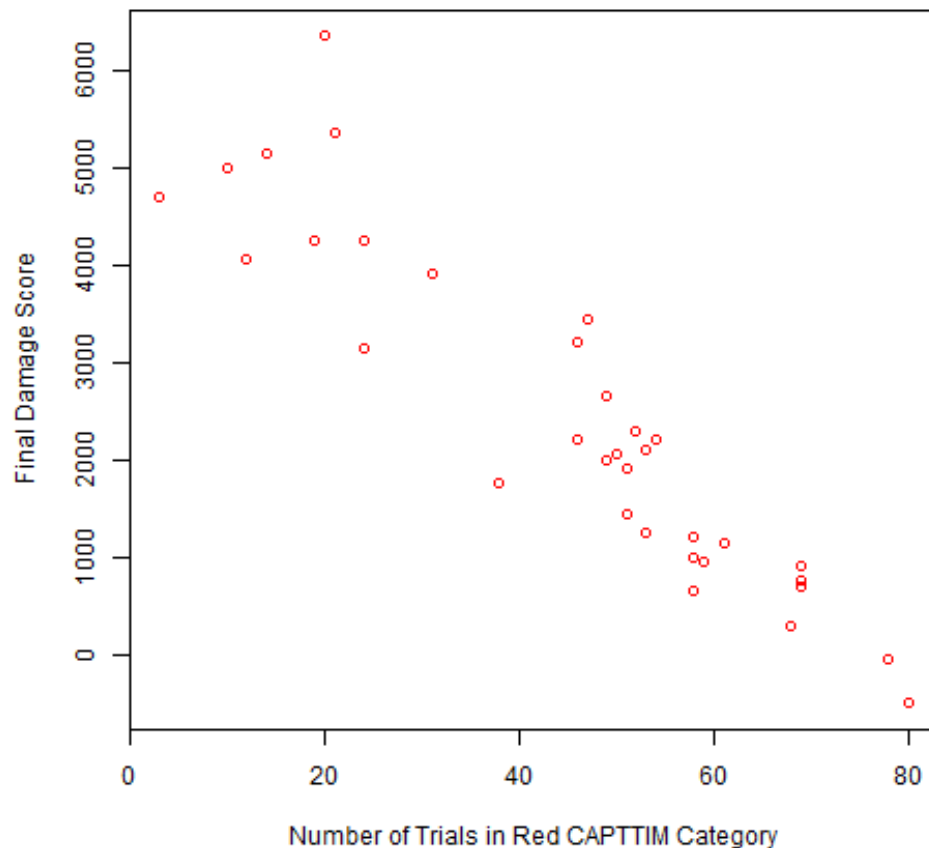


Figure 20. Correlation between final damage score and number of trials spent in the red category of CAPTTIM. The red dots show a strong negative correlation between number of trials spent in the red category and final damage score.

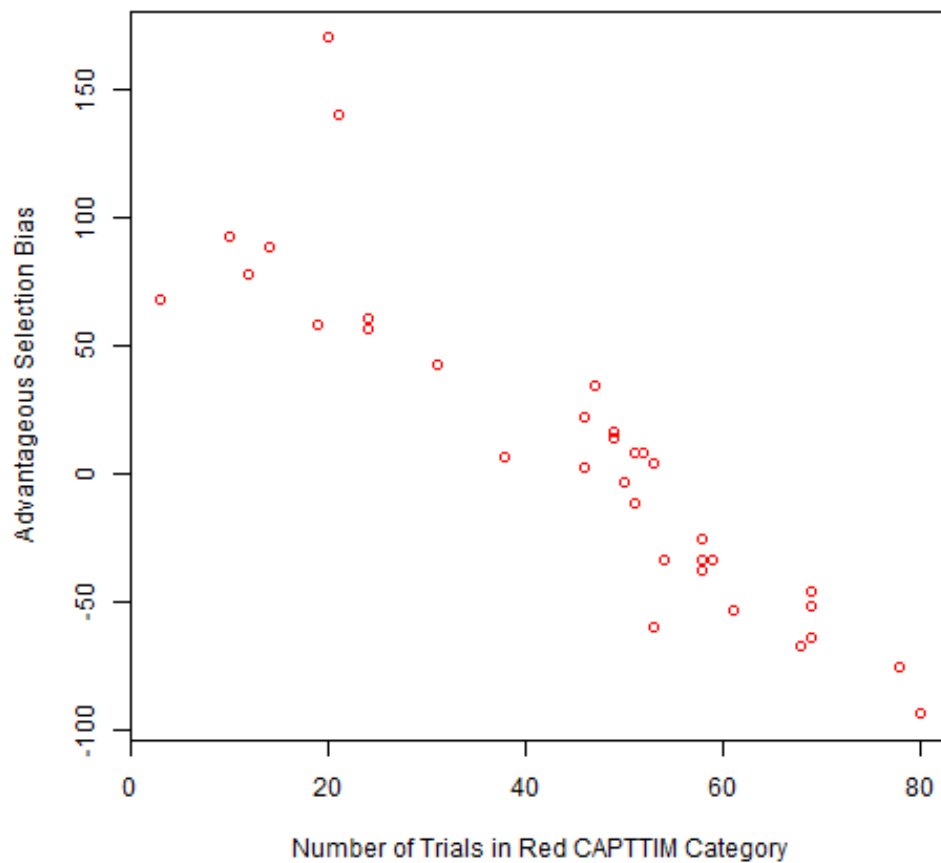


Figure 21. Correlation between advantageous selection bias and number of trials spent in the red category of CAPTTIM. The red dots show a strong negative correlation between number of trials spent in the red category of CAPTTIM and the subject's advantageous selection bias.

The Pearson correlation tests showed a strong negative correlation between the number of trials spent in the red category of CAPTTIM and a subject's final damage score and advantageous selection bias. The correlation test between final damage score and number of trials spent in the red category of CAPTTIM returned a correlation value of -0.92 , $p < .0001$ (95% CI: -0.96 to -0.85), which rejects the null hypothesis that true correlation is equal to 0. The correlation test between advantageous selection bias and number of trials spent

in the red category of CAPTTIM returned a correlation value of -0.90 , $p < .0001$ (95% CI: -0.95 to -0.81), which rejects the null hypothesis that true correlation is equal to 0.

An additional correlation test was suggested by Dr. Kennedy. Because the number of trials spent in the red and green category of CAPTTIM are not necessarily complementary, the same correlation tests described above were conducted looking at the number of trials spent in the green category of CAPTTIM. Figures 22 and 23 show the plots for these correlation tests.

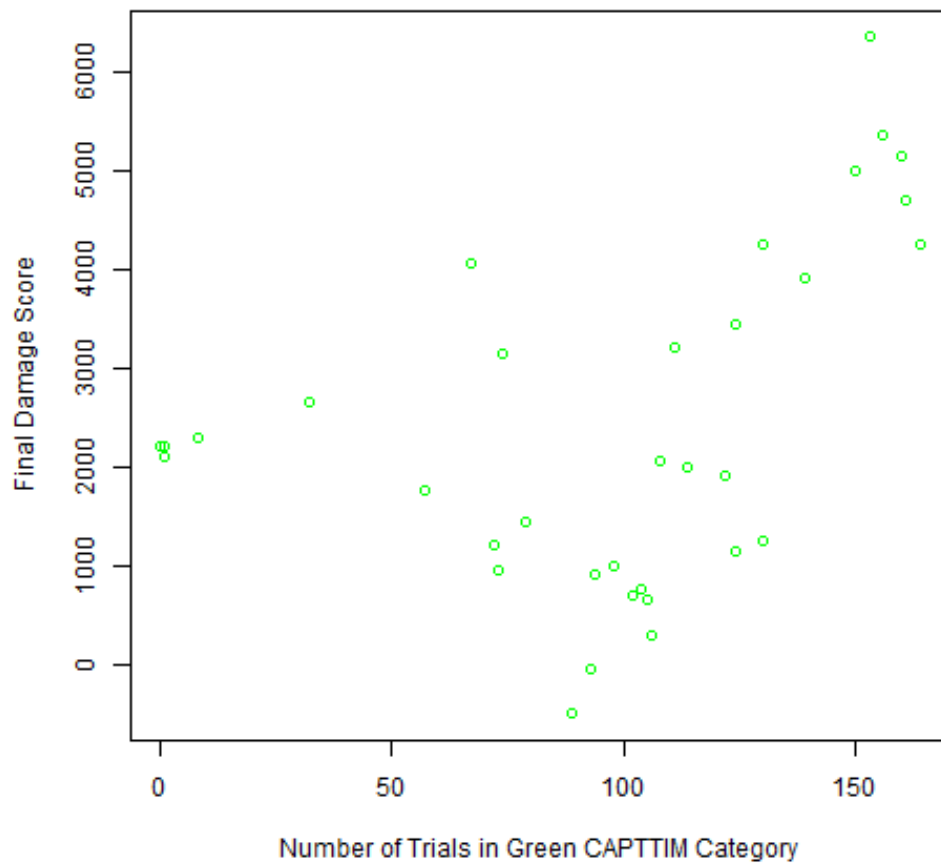


Figure 22. Correlation between final damage score and number of trials spent in the green category of CAPTTIM. The green dots show a moderately strong positive correlation between number of trials spent in the green category and final damage score.

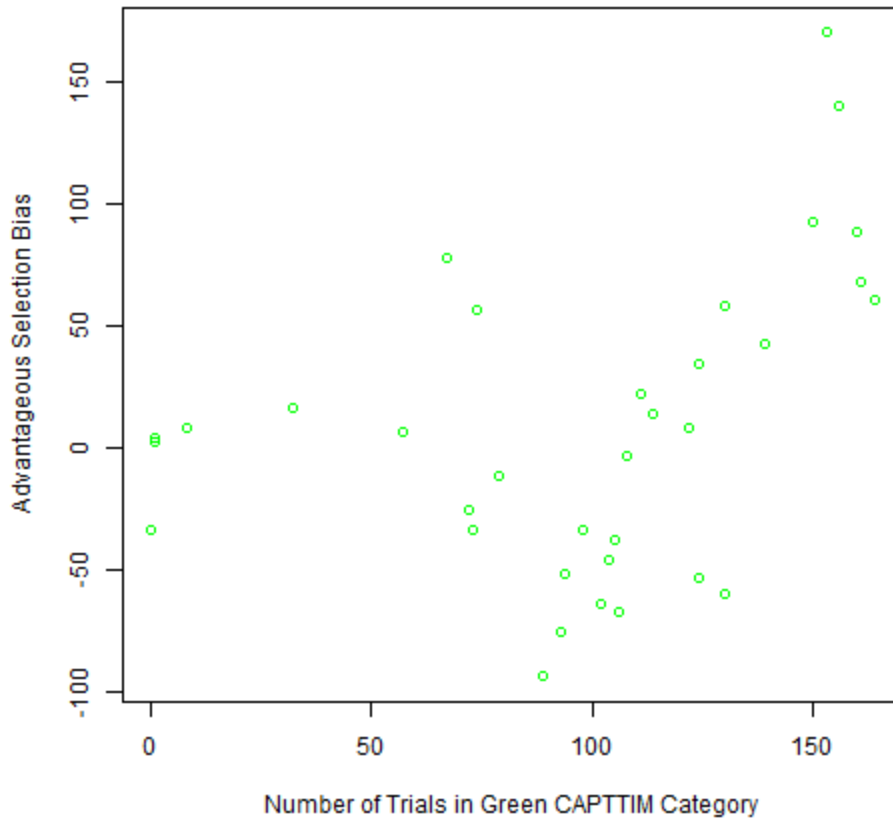


Figure 23. Correlation between advantageous selection bias and number of trials spent in the green category of CAPTTIM. The green dots show a moderately strong positive correlation between number of trials spent in the green category and advantageous selection bias.

Because the plots for these correlations were nonlinear, a Spearman's correlation test was utilized. These tests showed a moderately strong positive correlation between the number of trials spent in the green category of CAPTTIM and a subject's final damage score and advantageous selection bias. The correlation test between final damage score and number of trials spent in the green category of CAPTTIM returned a correlation value of 0.43, $p = .01$, which rejects the null hypothesis that true correlation is equal to 0. The correlation test between advantageous selection bias and number of trials spent in the green

category of CAPTTIM returned a correlation value of 0.38, $p = 0.01$, which rejects the null hypothesis that true correlation is equal to 0.

The weaker correlation between the number of trials spent in the green category of CAPTTIM and final damage score and advantageous selection bias was initially concerning to the research team. However, after further discussion and analysis the weaker correlation made sense. Because the population of high performers (high final damage scores and advantageous selection biases) was smaller within the subject population, the number of trials spent in the green category of CAPTTIM were not as abundant as the number of trials spent in the red category. Additionally, as discussed in the sections above, the third category of subjects were those who never transitioned between the cognitive state of exploration and exploitation. This category of subjects never had the opportunity to experience trials in the green category of CAPTTIM, based on the CAPTTIM categorization algorithm. These observations explained the weaker positive correlation between the numbers of trials spent in the green category compared to the strong negative correlation observed between the numbers of trials spent in the red category.

These results confirmed the use of change point analysis and route choice after trial 100 as an effective method of delineating between high and low regret. When combined with a subject's cognitive state data, these metrics provided an accurate means by which a subject's decision-making pattern could be categorized within the CAPTTIM model.

IV. DISCUSSION

The four primary goals of this thesis were to (1) find a threshold that delineated between high and low regret (decision performance), (2) combine the decision performance data with the cognitive state data, (3) validate these results and CAPTTIM, and (4) develop a visualization method for displaying a subject's CAPTTIM category on a trial by trial basis. All of these primary goals were achieved. This final chapter will summarize the methods used to complete the four primary thesis goals, discuss the implications of the research conducted, discuss future work that could be done to better the CAPTTIM algorithm, and conclude this thesis.

A. SUMMARY OF METHODS USED TO COMPLETE THESIS GOALS

After exploring several analytical approaches, an appropriate method for determining the threshold for regret was found by conducting a change point analysis on the regret per trial that a subject received. The resulting 15 process means returned by the change point analysis were then compared with the median of the subject's 15 process means. The median became the threshold that delineated between high and low regret and categorized the subject's decision performance. An additional metric was introduced based on the number of trials that it took good performers to converge on the ideal decision. On average, the subjects who performed well during the experiment determined that Routes 3 and 4 were the optimal choices by trial 100. Therefore, the additional metric automatically categorized subjects as having high regret if they chose Routes 1 or 2 after trial 100.

This decision performance data was then combined with the cognitive state data that categorized a subject's cognitive state as either exploration or exploitation. The four resulting combinations were (1) high regret and exploration, (2) low regret and exploration, (3) high regret and exploitation, and (4) low regret

and exploitation. As a result of these combinations, a subject's CAPTTIM category could be determined on a trial by trial basis.

The validation of the effectiveness of this CAPTTIM categorization was conducted by performing a Pearson's correlation between the number of trials spent in the red category of CAPTTIM, final damage score, and advantageous selection bias. The Pearson's correlation test was chosen due to the linearity this data exhibited. These correlation results exhibited a very strong negative correlation between these factors. As a result, the number of trials spent in the red category of CAPTTIM proved to be a strong inverse predictor of a subject's final damage score and advantageous selection bias. A Spearman's correlation test was conducted between the number of trials spent in the green category of CAPTTIM, final damage score, and advantageous selection bias. The Spearman's correlation test was chosen due to the nonlinearity this data exhibited. These correlation results showed a moderately strong positive correlation between these factors. As a result, the number of trials spent in the green category of CAPTTIM proved to be a moderate predictor of final damage score and advantageous selection bias.

Finally the visualization of the CAPTTIM category data was designed by creating a bar that exhibited the CAPTTIM category color for each trial. The yellow region of trials is where the subject is experiencing high regret, while their cognitive state is exploration. During a subject's exploration phase, high regret is acceptable and even encouraged. The subject needs to experience high regret in order to gain enough information about the environment to converge and exploit the optimal decision. The orange region of trials is where the subject is experiencing low regret, while their cognitive state is exploration. Long periods of low regret during exploration would require instructor intervention because the subject is ignorantly making the correct decision. Instructor intervention for the orange region could entail letting the subject know that they are making the correct decision or prompting them to sample more of the options to understand

why their decisions are better than the other options. The red region of trials is where a subject is experiencing high regret, while his or her cognitive state is exploitation. Instructor intervention would be required because the subject is exploiting the non-optimal decision believing it to be the optimal decision. The green region of trials is the ideal state in which the subject is experiencing low regret while their cognitive state is exploitation. This yellow, orange, red, and green bar was then overlaid on the regret per trial graph for each subject. This visualization proved to be an effective means of communicating when and where a subject's performance and cognitive state were aligned or misaligned.

B. IMPLICATIONS

The implications of this research are many. CAPTTIM provides feedback on a subject's deviations from the ideal decision path/optimal decision pattern. Based on these deviations, CAPTTIM could provide meaningful feedback to an instructor on the timing and type of intervention that is needed by the trainee. While CAPTTIM is most suited for tasks in which the ideal decision path is known, it could be extrapolated to fit other types of tasks, like rapid response decisions or interactive tactical decision-making games, where understanding optimal decision-making would be beneficial. Another example that CAPTTIM could be extrapolated to fit is wargaming. In wargaming, a commander makes decisions based on the intelligence he/she has received and through trial and error determines the best course of action to execute. The optimal decision path is much more difficult to determine in these examples, but could be determined based on military tactics specific to the wargaming scenario. In these examples inexperienced commanders could conduct wargaming to gain experience that does not involve human lives and receive feedback via CAPTTIM on when and where their performance was aligned or misaligned with their cognitive state.

Another implication of this research is that Army has begun a renewed focus on enhancing the leadership and knowledge of its personnel. The fact that technology has advanced to the degree that countries that used to be inferior in

their military capabilities can now develop quick and innovative solutions that have near peer capabilities, has led the Army to the conclusion that its human resources are its most valuable, adaptable, and flexible assets (Odierno & McHugh, 2015). Based on this conclusion the focus on leadership development tools that train military personnel to be agile, adaptive, and innovative problem solvers in an ambiguous and complex environment has been initiated at the highest level within the Army (Odierno & McHugh, 2015). These leadership development tools range from tasks that aim to improve working memory, comprehending languages, calculating, reasoning, problem solving, and decision-making (Odierno & McHugh, 2015). The ultimate goal of these leadership development tools is to provide technology developed instruction that employs adaptive learning strategies and intelligent tutoring to accelerate learning and education for Soldiers and Army Civilians (Odierno & McHugh, 2015).

The convoy task that was used to collect the data analyzed in this thesis elicits many of the Army's desired leadership development qualities. It requires the user to be adaptive, agile, conduct reasoning, problem solve, and increases working memory and decision-making capabilities. Additionally, the work done in this thesis, specifically the advancement of the model CAPTTIM, has many implications across these leadership development tools. CAPTTIM could be utilized to provide the aspect of intelligent tutoring that could be applied to these technology developed instruction applications that are desired by the Army. Because of CAPTTIM's ability to identify decision performance and cognitive misalignment, it could be used as an intelligent tutor to provide useful feedback to the trainee. Based on these implications the research team believes that CAPTTIM provides a valuable capability to the Army's research on how to develop better leaders and decision makers.

C. FUTURE WORK

As previously stated the delineation between high and low regret and the cognitive states of exploration and exploitation was calculated retrospectively. In order for CAPTTIM to be able to provide “real-time” feedback to an instructor or even a trainee, these delineations must be calculated dynamically. This is the most crucial advancement that must take place in this research in order for CAPTTIM to be a more effective tool for instructors. One way that this can be accomplished is to have a “burn in period” that is a set number of trials where no feedback is provided and a subject is not categorized into any CAPTTIM category. Once this period is complete, a change point analysis of regret per trial can be performed to determine the threshold that delineates between high and low regret. After this threshold is calculated for this period, all future decision performance can be compared to that threshold on a trial by trial basis. The same concept applies to the EWMA of latency in decision-making times in order to provide the delineation between the cognitive states of exploration and exploitation. Once this threshold is calculated for the “burn in period” a subject can be categorized into one of the two cognitive states on subsequent trials. These two delineations can then be combined, as they were in this thesis, to categorize subjects into a CAPTTIM category. An initial analysis of this “burn in period” concept with the research team, suggested that a period of 50–80 trials would be sufficient to calculate a threshold for decision performance and cognitive states.

Other future work would be to (1) test CAPTTIM on a task that differs from the convoy task, and (2) develop the CAPTTIM oriented intervention feedback loop. Testing CAPTTIM on a task like wargaming, rapid response decisions, or tactical decision-making games will help validate CAPTTIM’s adaptability to different tasks. By validating the adaptability of CAPTTIM, the significance of this research to the Army’s leadership development focus will be further solidified. The development of the CAPTTIM oriented intervention feedback loop is necessary to enable the model to be used as an intelligent tutor in computer

based tasks. The ability for a script to be created that utilizes data categorized by CAPTTIM and provides task specific guidance/feedback to a trainee will, again, further illustrate CAPTTIM's implication to the Army's leadership development program.

D. CONCLUSION

Understanding optimal decision-making is a very difficult task, but one that is worth undertaking. The Army and the military as a whole have realized that, due to budget constraints, they are entering into one of the most fiscally austere environments that the military has experienced in decades (Odierno & McHugh, 2015). As a result, they have grasped that the dominance of the United States military will not be accomplished by the unlimited acquisition of newer weapons, vehicles, and technology (Odierno & McHugh, 2015). Thus military dominance will be measured by the ability to develop military professionals that are capable of being effective, agile, adaptive, and innovative decision makers and problem solvers (Odierno & McHugh, 2015). The focus on force development versus the acquisition of material solutions lends gravity to the research conducted in this thesis.

The research team believes that the work done in this thesis has furthered the understanding of decision-making and directly provides a useful tool that could be used to aid leadership development programs. While there is still much to discover when it comes to understanding how humans process information and make decisions, this research has made it more possible to understand and classify decision performance and cognitive state. With this understanding the human mind becomes less of a black box, in which an instructor or intelligent tutor has no insight, and allows a small peek at what is really going on in the subject's decision-making process. This peek is made possible by the ability to understand and identify the alignment or misalignment of cognitive state with decision performance. By looking at a common reinforcement learning task, modified for the military domain, the research team was able to investigate and

better understand a subject's decision-making pattern and how to intelligently influence this pattern if determined to be suboptimal. It will be exciting to see what follow on research discovers, and how CAPTTIM is modified to increase the understanding of optimal decision-making.

THIS PAGE INTENTIONALLY LEFT BLANK

APPENDIX A. PAYOUT SCHEDULE FOR IGT AND CONVOY TASK

IGT Payout Schedule			
Deck A	Deck B	Deck C	Deck D
-150	100	50	50
-250	100	0	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	-1150	0	50
100	100	0	-200
-150	100	50	50
-250	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50

Convoy Task Payout Schedule			
Rout 1	Route 2	Route 3	Route 4
-150	100	50	50
-250	100	0	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	-1150	0	50
100	100	0	-200
-150	100	50	50
-250	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50

100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50

100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50

100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
100	100	0	50
-50	100	50	50
100	100	0	50
-200	100	50	50
100	100	0	50
-100	100	50	50
100	-1150	0	50
-150	100	0	-200
-250	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50

-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50

-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50

-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50

-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200
100	100	50	50
100	100	50	50
-50	100	0	50
100	100	50	50
-200	100	0	50
100	100	50	50
-100	100	0	50
100	100	50	50
-150	-1150	0	50
-250	100	0	-200

APPENDIX B. R SCRIPTS

A. EWMA OF DECISION LATENCY TIMES R SCRIPT

```
print("begin script: ODM multi-arm bandit analysis")
setwd("~/NPS/Thesis/Thesis Data/Data Critz")
require(zoo)
require(ggplot2)
require(fTrading)
require(qcc)
require(RColorBrewer)
require(StatMatch)
IGT <- T # Are we using the published IGT payout schedule?
PlayerInput <- T # Are we analysing a human player?
doRegretA.mb <- T # regret by absolute

Basics <- F # plot basic histograms
BasicsT <- F # plot basic histograms

# Create, test through MC, plot new distributions...

numTrials <- 200 # ignore any more than 200 trials
cog.state <- vector() #Capture cognitive state data
route.select <- vector() #Capture route choice

# Read in payout schedule
IGTresponse <- read.csv("IGTImproved.csv")
numBandits = length(IGTresponse)
numTrials <- 200

# Read in player input
if (PlayerInput){
  files <- list.files(pattern = '*MultiArmBandit*')
  numPlayers <- length(files)
  numBandits <- 4
  subject <- 1
  # Create dataframe for subject specific response
  MA.decision <- data.frame(matrix(0,nrow=200,ncol=numPlayers))
  # Create dataframe for descriptive statistics
  MA.summary <- data.frame(matrix(0,nrow=numPlayers,ncol=35))
  header <- c('Subject','mb.FD.100','mb.numFD.100','mb.numHFD.100',
              'mb.R1.100','mb.R2.100','mb.R3.100','mb.R4.100','mb.adv.sb.100',
```

```

'mb.mean.l.100','mb.med.l.100','mb.sd.100','mb.numFD.SecHalf','mb.numHFD.SecHalf',

'mb.R1.SecHalf','mb.R2.SecHalf','mb.R3.SecHalf','mb.R4.SecHalf','mb.adv.sb.SecHalf',

'mb.mean.l.SecHalf','mb.med.l.SecHalf','mb.sd.SecHalf','mb.FD.200','mb.numFD.200','mb.numHFD.200',
      'mb.R1.200','mb.R2.200','mb.R3.200','mb.R4.200','mb.adv.sb.200',
      'mb.mean.l.200','mb.med.l.200','mb.sd.200','SigLat','perc.regret')
names(MA.summary) <- header

# df used for calculating regret
Regret.mb.df <- data.frame(matrix(0,nrow=0,ncol=5))

#Import Player choices and resulting response by trial
#file <- files[1]
element<-1
for(file in files){

  PlayerID <- file#paste('Subject ',subject,sep='')
  print(PlayerID)
  player <- read.csv(file)
  #print(summary(player))
  LL <- list()
  player<- subset(player, trial<201)
  numTrials <- length(player[,1])

  # add players decision to MA.decision
  colnames(MA.decision)[element]<-as.numeric(noquote(strsplit(PlayerID,"")[[1]])[1])
  MA.decision[element] <- player$routeSel
  decide <- as.numeric(player$routeSel)    # get decision data)
  decide[decide== "1"] <- -1 # recode selections to adv sel scores
  decide[decide== "2"] <- -1
  decide[decide== "3"] <- 1
  decide[decide== "4"] <- 1
  element<-element+1

  # Latency by trial number plot
  numShift <-numTrials-1
  shift <-append(0,head(player$trialLoss,numShift),after=1)
  Damage.before <-factor(player$trialLoss)

```



```

Damage.after <-factor(shift)
size.before <-factor(player$trialLoss)
size.after <-factor(shift)
Damage.color <-factor(player$trialLoss)
damage.cat <-list('none to low (0,50) '=0,'none to low (0,50) '=50,'med
(150,200,250,300,350) '=150,'med (150,200,250,300,350) '=200,
' med (150,200,250,300,350) '=250,'med
(150,200,250,300,350) '=300,'med (150,200,250,300,350) '=350,'high
(1250) '=1250)
damage.size<-
list('10 '=0,'10 '=50,'20 '=150,'20 '=200,'20 '=250,'20 '=300,'20 '=350,'100 '=1250)
damage.color<-
list('3 '=0,'3 '=50,'2 '=150,'2 '=200,'2 '=250,'2 '=300,'2 '=350,'5 '=1250)
levels(Damage.before) <- damage.cat
levels(Damage.after) <- damage.cat
levels(size.before) <- damage.size
levels(size.after) <- damage.size
levels(Damage.color) <- damage.color
myColors <- brewer.pal(5,"Set1")
names(myColors) <- c(100,20,10)
colScale <- scale_colour_manual(name = "damage",values = myColors)

```

```

player<-
cbind(player,Damage.before,Damage.after,size.before,size.after)#,ewmaS)

```

```

####Fill in summary stats for 100 trials
#'Subject'
subject <- as.numeric(noquote(strsplit(PlayerID, " ")[[1]])[1])
MA.summary[subject,1]<- subject
#'Final Damage'
MA.summary[subject,2]<- player$Damage[100]
### trials friendly damage'
MA.summary[subject,3]<- sum(player$trialLoss[1:100]>0)
### trials heavy friendly damage'
MA.summary[subject,4]<- sum(player$trialLoss[1:100]>1000)
#'Route 1'
MA.summary[subject,5]<- sum(player$routeSel[1:100]=='1')/100
#'Route 2'
MA.summary[subject,6]<- sum(player$routeSel[1:100]=='2')/100
#'Route 3'
MA.summary[subject,7]<- sum(player$routeSel[1:100]=='3')/100
#'Route 4'
MA.summary[subject,8]<- sum(player$routeSel[1:100]=='4')/100
#'advantageuos selection bias'

```

```

MA.summary[subject,9]<-
sum(player$routeSel[1:100]=='3')+sum(player$routeSel[1:100]=='4')-
sum(player$routeSel[1:100]=='1')-sum(player$routeSel[1:100]=='2')
#'mean latency time'
MA.summary[subject,10]<- mean(player$latent[2:100])
#'median latency'
MA.summary[subject,11]<- median(player$latent[2:100])
#'standard deviation latency'
MA.summary[subject,12]<- sd(player$latent[2:100])

#Fill in summary stats for second half, 101-200 trials
### trials friendly damage'
MA.summary[subject,13]<- sum(player$trialLoss[101:200]>0)
### trials heavy friendly damage'
MA.summary[subject,14]<- sum(player$trialLoss[101:200]>1000)
#'Route 1'
MA.summary[subject,15]<- sum(player$routeSel[101:200]=='1')/100
#'Route 2'
MA.summary[subject,16]<- sum(player$routeSel[101:200]=='2')/100
#'Route 3'
MA.summary[subject,17]<- sum(player$routeSel[101:200]=='3')/100
#'Route 4'
MA.summary[subject,18]<- sum(player$routeSel[101:200]=='4')/100
#'advantageuos selection bias'
MA.summary[subject,19]<-
sum(player$routeSel[101:200]=='3')+sum(player$routeSel[101:200]=='4')-
sum(player$routeSel[101:200]=='1')-sum(player$routeSel[101:200]=='2')
#'mean latency time'
MA.summary[subject,20]<- mean(player$latent[101:200])
#'median latency'
MA.summary[subject,21]<- median(player$latent[101:200])
#'standard deviation latency'
MA.summary[subject,22]<- sd(player$latent[101:200])

#Fill in summary stats for 200 trials
#'Final Damage'
MA.summary[subject,23]<- player$Damage[numTrials]
### trials friendly damage'
MA.summary[subject,24]<- sum(player$trialLoss>0)
### trials heavy friendly damage'
MA.summary[subject,25]<- sum(player$trialLoss>1000)
#'Route 1'
MA.summary[subject,26]<- sum(player$routeSel=='1')/numTrials
#'Route 2'
MA.summary[subject,27]<- sum(player$routeSel=='2')/numTrials

```

```

#'Route 3'
MA.summary[subject,28]<- sum(player$routeSel=='3')/numTrials
#'Route 4'
MA.summary[subject,29]<- sum(player$routeSel=='4')/numTrials
#'advantageuos selection bias'
MA.summary[subject,30]<-
sum(player$routeSel=='3')+sum(player$routeSel=='4')-
sum(player$routeSel=='1')-sum(player$routeSel=='2')
#'mean latency time'
MA.summary[subject,31]<- mean(player$latent[2:200])
#'median latency'
MA.summary[subject,32]<- median(player$latent[2:200])
#'standard deviation latency'
MA.summary[subject,33]<- sd(player$latent[2:200])
#'Significant latency'
MA.summary[subject,34]<- mean(player$latent[player$size.before==100])

if(doRegretA.mb){
  num.a <- 1 # set the next trial to one for each option
  num.b <- 1
  num.c <- 1
  num.d <- 1
  regret.total <- 0 # initialize total regret
  regret.c <- 0 # initialize regret count
  regret.r <- 0 # initialize regret rate
  for(trial in 1:numTrials){ # for every trial (withing every player loop)
    # The best option value (gain+loss already computed) in the schedule for
each option
    opt.choice.v<-
max(IGTresponse[num.a,1],IGTresponse[num.b,2],IGTresponse[num.c,3],IGTres
ponse[num.d,4])
    # From the records, what they gained and lost
    player.choice.v <- player$trialGain[trial]-player$trialLoss[trial] # positive is
good
    # find the difference
    regret.v <- opt.choice.v - player.choice.v
    if(regret.v>0){regret.c <- regret.c +1}
    regret.r <- regret.c/trial
    # accumulate regret
    regret.total <- regret.total + regret.v
    # normalize by trials
    regret.mean <- regret.total / trial
    # error check

```

```

#
if(regret.v<0){print(paste(num.a,num.b,num.c,num.d,'opt',opt.choice.v,'player',player.choice.v,'regret =',regret.v,' sub ',subject,' trial ',trial))}
# update next available options
if(player$routeSel[trial]==1){num.a<-num.a+1}
if(player$routeSel[trial]==2){num.b<-num.b+1}
if(player$routeSel[trial]==3){num.c<-num.c+1}
if(player$routeSel[trial]==4){num.d<-num.d+1}
# combine into row
trial.regret<-
c(trial,decide[trial],regret.v,regret.total,regret.mean,subject,regret.r)
# add to Regret.df data.frame of all trial/regret measure/player combinations
Regret.mb.df <- rbind(Regret.mb.df,trial.regret)
}
}
#Significant latency'
MA.summary[subject,35]<- regret.r

player <- player[-1,] # Remove first latency observation

#### Sequential Detection Methods for Detecting Exploration-Exploitation Mode Changes

#### Method 1: The Exponentially Weighted Moving Average

# develop single number of standard deviation of all latencies after low damage
threshold <- 2 # threshold multiplier
mb.sd.threshold <- sd(player$latent[player$size.before==10])*threshold

# develop estimate of moving latency from exponential moving  $z_t = \alpha y_t + (1-\alpha) z_{t-1}$ 
EWMAlambda <- .1 # lambda
ewma.latent.lst<-
ewmaSmooth(player$trial[player$size.before==10],player$latent[player$size.before==10],lambda=EWMAlambda) # list of estimate data

# build a dataframe with this data in it
EWMA <- data.frame(matrix('NA',nrow=length(ewma.latent.lst$x),ncol=3))
header <-c('trial','ewma','threshold')
names(EWMA) <- header
EWMA['trial'] <- ewma.latent.lst$x
EWMA['ewma'] <- ewma.latent.lst$y
EWMA['threshold'] <- mb.sd.threshold

# merge it with the other player data

```

```

player <- merge(player,EWMA,by="trial",all.x=T,fill=NA)

# Inpute data from missing high damage +1 trials
# input by 'hot deck', simply continue last value until next observation (estimate in
this case)
ewma.shift<-append(0,head(player$ewma,length(player$ewma)-1),after=1)
#vector from shifting ewma down 1
num.mistakes <-5
for(mistake in 1:num.mistakes){
ewma.shift<-append(0,head(ewma.shift,length(ewma.shift)-1),after=1)#shift
again...
player$ewma[is.na(player$ewma)]<-ewma.shift[is.na(player$ewma)]
}

# build upper and lower bounds for colored ribbons on graph
player['upper.line'] <- apply(cbind(player$threshold,player$ewma),1,max)
player['lower.line'] <- apply(cbind(player$threshold,player$ewma),1,min)
cog.stateTmp <- numeric(200)
cog.stateTmp[1] <- "explore"
cog.stateTmp[2:200] <- ifelse(player$ewma>player$threshold,"explore","exploit")
cog.state <- c(cog.state,cog.stateTmp)
#Due to long latency, we do not count the first route selection.
route.selectTmp <- numeric(200)
route.selectTmp[1] <- 0 #Can be any value for this analysis
route.selectTmp[2:200] <- player$routeSel
route.select <- c(route.select,route.selectTmp)

#### Method 2: Monitoring Sequential Sample Variances

####Create / Save graphs for each subject
#   maxLatent <- 8
#   gtitle <- paste('Latency and EWMA by trial number for',PlayerID)
#   ftitle <- paste0(subject,'TxL.png')
#   LatByTrial<-ggplot(data=player,aes(x=trial,y=latent))+
#
geom_ribbon(aes(ymin=threshold,ymax=upper.line,linetype="NA"),fill="orange",alpha=.5,show_guide=F)+
#
geom_ribbon(aes(ymin=lower.line,ymax=threshold,linetype="NA"),fill="skyblue",alpha=.5,show_guide=F)+
#
labs(title=gtitle)+coord_cartesian(ylim=c(0,maxLatent))+colScale+theme_bw()+xlab("Trials")+ylab("Latency")

```

```

# LatByTrial<-
LatByTrial+geom_line(data=player,aes(x=trial,y=ewma),linetype=1,colour="grey8
8")
# LatByTrial<-
LatByTrial+geom_point(data=player,aes(x=trial,y=latent,color=size.after,size=siz
e.after),show_guide=T)
# #png(file=ftitle,width = 1000, height = 700)
# print(LatByTrial)
# maxLatent <- 8
# gtitle <- paste('Latency and EWMA by trial number for',PlayerID)
# ftitle <- paste0(subject,'TxL.png')
# LatByTrial<-ggplot(data=player,aes(x=trial,y=latent))+
#
geom_ribbon(aes(ymin=threshold,ymax=upper.line,linetype=NA,fill="Explore"),al
pha=.5,show_guide=T)+
#
geom_ribbon(aes(ymin=lower.line,ymax=threshold,linetype=NA,fill="Exploit"),alp
ha=.5,show_guide=F)+
# scale_fill_manual(values=c("Explore"='orange',"Exploit"="skyblue"))+
#
#labs(title=gtitle)+coord_cartesian(ylim=c(0,maxLatent))+theme_bw()+xlab("Trial
s")+ylab("Latency")
#
labs(title=gtitle)+coord_cartesian(ylim=c(0,maxLatent))+colScale+theme_bw()+xl
ab("Trials")+ylab("Latency")
#LatByTrial<-
LatByTrial+geom_line(data=player,aes(x=trial,y=ewma),linetype=1,colour="grey8
8")
#LatByTrial<-
LatByTrial+geom_point(data=player,aes(x=trial,y=latent,color=size.after,size=siz
e.after),show_guide=T)
# #png(file=ftitle,width = 1000, height = 700)
#
# print(LatByTrial)
# dev.off()
#
# gtitle <- paste('Route by trial number for',PlayerID)
# plotBT<- ggplot(player,aes( trial,colour = size.before,factor(routeSel))) +
labs(title = gtitle)+colScale
# plotBT<-plotBT+geom_point(aes(size = size.before),show_guide = F) +
theme_bw()+ xlab("Trials") +ylab("Routes")
# #plotBT<-plotBT+geom_point(aes(colour = Damage.color))#+
scale_fill_continuous(name = "Friendly damage on previous
trial")#+coord_cartesian(ylim=c(0,8))

```

```

#   plotBT<-plotBT + theme(legend.direction = "horizontal", legend.position =
"bottom")#)+annotate("text", x = 0, y = 10, label = "Relationship between x and y")
#   #LatByTrial+ guides(fill = guide_legend(title.theme = element_text(size=15,
face="italic", colour = "red", angle = 45)))
#   ftitle <- paste0(subject,'TxR.png')
#   png(file=ftitle,width = 1000, height = 700)
#   suppressWarnings(print(plotBT))
#   dev.off()

  subject <- subject+1

}

header<-
c('trial','adv.sel.bias','regret.trial','regret.total','regret.mean','subject','regret.rate')
names(Regret.mb.df) <- header
} # end of read in player input (PlayerInput)

survey_data<-
merge(read.csv("survey_data.csv"),read.csv("groups.csv"),by="Subject")
total<-merge(survey_data,MA.summary,by="Subject")
Regret.mb.df$Cog.State <- cog.state
Regret.mb.df$RouteSel <- route.select

save.image("C:/Users/John/Documents/NPS/Thesis/ThesisData/Data
Critz/RegretData.RData")

```

B. CHANGEPOINT ANALYSIS R SCRIPT

```

setwd("~/NPS/Thesis/Thesis Data/Data Critz")
load("C:/Users/John/Documents/NPS/Thesis/ThesisData/Data
Critz/RegretData.RData")

library("changepoint")
subject.vec <- unique(Regret.mb.df$subject) #For all subjects
#subject.vec <- subject.vec[9]
#subject.vec <- c(1,4,8,11,14,15,17,26,28)
regret.vec <- numeric(200)
median.vec <- numeric(200)
med.dev <- numeric(200)
#upperCTLLimit <- numeric(200)
bin <- list()
chngepoint.bin <- list()
bin.vec <- numeric(200)
subject.index <- 1
subject.start <- 1

```

```

subject.difference <- 200

for(index in 1:length(subject.vec)){
  subject.tmp <- which(Regret.mb.df$subject==subject.vec[index])
  test.subj <- Regret.mb.df[subject.tmp[1]:subject.tmp[200],]
  # a <- 1
  # b <- 5
  bin.index <- 1
  tmp.chng <- cpt.mean(test.subj[,3], method="SegNeigh",Q=15)
  chngepoint.bin[[index]] <- tmp.chng
  #Corrected histogram label
  png(paste("RegretHistogramSubject",subject.vec[index],".png",sep=""))
  hist(test.subj[,3],col="blue",xlab="Regret Value",main=paste("Regret Histogram
for Subject ",subject.vec[index],sep=""))
  dev.off()
}

save.image("C:/Users/John/Documents/NPS/Thesis/ThesisData/Data
Critz/RegretData.RData")

```

C. CAPTTIM VISUALIZATION R SCRIPT

```

#Had to create the vector for subject 9 manually
#Source Revised MultiArm
#Source Regret.Mean file
require(data.table) #Required to find unique column elements
#Find the subjects we want
#subject.vec <- unique(Regret.mb.df$subject) #For all subjects
#subject.vec <- c(1,4)
#subject.vec <- c(11)
#index <- 1
#subject.vec <- subject.vec[-c(1:8)]
#subject.vec1 <- subject.vec[-9]

subject.control.vec1 <- vector()
subject.category1 <- vector()
index <- 1
for(index in 1:length(subject.vec)){
  print(paste("Processing Subject ",subject.vec[index]))
  subject.tmp <- which(Regret.mb.df$subject==subject.vec[index])
  test <- Regret.mb.df[subject.tmp[1]:subject.tmp[200],]
  test2 <- chngepoint.bin[[index]]
  chgptmean.vec <- numeric(200) #Creat a vector to collect the changepoints
  i <- 1
  while(i < length(test2@cpts)+1){
    # browser()

```



```

# print(paste("I is ",i))
# print(chgptmean.vec)
if(i==1){
  chgptmean.vec[i] <- test2@param.est$mean[i]
  i <- i + 1
  next
}
if(test2@cpts[i]!=200){
  if(test2@cpts[i]-test2@cpts[i-1]==1){
    chgptmean.vec[test2@cpts[i]] <- test2@param.est$mean[i]
    i <- i + 1
    next
  }
  if(test2@cpts[i+1]-test2@cpts[i]==1){
    chgptmean.vec[(test2@cpts[i-1]+1):(test2@cpts[i])]<-
test2@param.est$mean[i]
    i <- i + 1
    next
  }
}

  if(test2@cpts[i+1]-test2@cpts[i]>1){
    chgptmean.vec[(test2@cpts[i-1]+1):(test2@cpts[i])<-
test2@param.est$mean[i]
    i <- i + 1
    next
  }
}

  if(test2@cpts[i]==200){
    chgptmean.vec[(test2@cpts[i-1]+1):(test2@cpts[i])<-
test2@param.est$mean[i]
    i <- i+1
  }
}

```

test\$Mean.Regret <- chgptmean.vec #Add this to whatever dataframe you would like of the same length

#Now let's add color

#First let's find out which trials were in or out of control

control.vec <- numeric(200)

for(i in 1:200){

```

  if(test$Mean.Regret[i]>median(test2@param.est$mean)){
    control.vec[i] <- "high"
  }

```

```

    if(test$Mean.Regret[i]<=median(test2@param.est$mean)) {
      control.vec[i] <- "low"
    }
  }

test$Control <- control.vec
subject.control.vec1 <- c(subject.control.vec1,control.vec)

#Next, make up a color for each value
color.vec <- numeric(200)
for(i in 1:200){
  if(i <= 100){
    if(test$Cog.State[i]=='explore' & test$Control[i]=="low"){
      color.vec[i] <- "orange"
    }
    if(test$Cog.State[i]=='explore' & test$Control[i]=="high") {
      color.vec[i] <- "yellow"
    }
    if(test$Cog.State[i]=='exploit' & test$Control[i]=="low") {
      color.vec[i] <- "green"
    }
    if(test$Cog.State[i]=='exploit' & test$Control[i]=="high") {
      color.vec[i] <- "red"
    }
  }
  if(i > 100){
    if(test$RouteSel[i]==2) {
      color.vec[i] <- "red"
      next
    }
    if(test$RouteSel[i]==1) {
      color.vec[i] <- "red"
      next
    }
    if(test$Cog.State[i]=='explore' & test$Control[i]=="low"){
      color.vec[i] <- "orange"
    }
    if(test$Cog.State[i]=='explore' & test$Control[i]=="high") {
      color.vec[i] <- "yellow"
    }
    if(test$Cog.State[i]=='exploit' & test$Control[i]=="low") {
      color.vec[i] <- "green"
    }
    if(test$Cog.State[i]=='exploit' & test$Control[i]=="high") {
      color.vec[i] <- "red"
    }
  }
}

```

```

    }
  }

}
#test$Color <- color.vec
subject.category1 <- c(subject.category1,color.vec)
test$Color <- color.vec
png(paste("Subject",subject.vec[index],"CAPTTIMPlot.png",sep=""))
plot(c(1, 200), c(1, 1250), type = "n", main= paste("Subject ",subject.vec[index],"
CAPTTIM",sep=""),
      xlab="Trial",ylab="Regret Per Trial") #Creat a blank plot
color.index <- data.table:::uniqlist(list(test$Color))
i <- 1
while(i < max(color.index)){
  #browser()
  #cat("i is now",i)
  tmp <- which(color.index==i)
  if(length(tmp)==0){
    i <- i+1
    tmp <- which(color.index==i)
  }
  if(length(tmp)==1){
    if(i < max(color.index)){
      if(color.index[tmp+1]-color.index[tmp]==1){ #check for single change points
at a trial
        #cat("i is",i,"\n")
        rect(color.index[tmp],0,color.index[tmp+1],100,col=test$Color[i])
        i <- i+1
        tmp <- which(color.index==i)
      }
    }
    if(length(tmp)!=0 && tmp !=length(color.index)){
      if(color.index[tmp+1]-color.index[tmp]==1){ #check for single change points
at a trial
        #cat("i is",i,"\n")
        rect(color.index[tmp],0,color.index[tmp+1],100,col=test$Color[i])
        i <- i+1
        next
      }
      if(color.index[tmp+1]-color.index[tmp]>1){
        #cat("i is",i,"\n")
        rect(color.index[tmp],0,color.index[tmp+1],100,col=test$Color[i])
        i <- i+1
        tmp <- which(color.index==i)
      }
    }
  }
}

```

```

    }
  }
  if(length(tmp)!=0 && tmp == length(color.index)){
    rect(color.index[tmp],0,200,100,col=test$Color[i])
    break
  }
  else{
    #cat("i is",i,"\n")
    i <- i+1
  }
}
}

lines(test$regret.trial,lty=2,col="blue")
dev.off()

}

Regret.mb.df$Regret.Level <- subject.control.vec1
Regret.mb.df$Capttim.Category <- subject.category1
save.image("C:/Users/John/Documents/NPS/Thesis/ThesisData/Data
Critz/RegretData.RData")
write.csv(Regret.mb.df,file="SubjectData.csv")

```

D. CORRELATION TEST R SCRIPT

```

#Loop through each subject
#Take out row 16 of MA summary
MA.summaryTest <- MA.summary[-16,]
red.count.vec <- vector()
green.count.vec <- vector()
for(i in MA.summaryTest$Subject){
  tmp.df <- Regret.mb.df[Regret.mb.df$subject==i,]
  red.count <- sum(tmp.df$Capttim.Category=='red')
  red.count.vec <- c(red.count.vec,red.count)
  green.count <- sum(tmp.df$Capttim.Category=='green')
  green.count.vec <- c(green.count.vec, green.count)
}

pearsonTest(red.count.vec,MA.summaryTest$mb.FD.200)

pearsonTest(red.count.vec,MA.summaryTest$mb.adv.sb.200)

spearmanTest(green.count.vec,MA.summaryTest$mb.FD.200)

```

```
spearmanTest(green.count.vec,MA.summaryTest$mb.adv.sb.200)
```

```
png(paste("CorrelationTestRedFD.png"))  
plot(xlab = "Number of Trials in Red CAPTTIM Category",  
     ylab = "Final Damage Score",  
     red.count.vec,  
     MA.summaryTest$mb.FD.200, col = "red")  
dev.off()
```

```
png(paste("CorrelationTestRedAdvSelectBias.png"))  
plot(xlab = "Number of Trials in Red CAPTTIM Category",  
     ylab = "Advantageous Selection Bias",  
     red.count.vec,  
     MA.summaryTest$mb.adv.sb.200, col = "red")  
dev.off()
```

```
png(paste("CorrelationTestGreenFD.png"))  
plot(xlab = "Number of Trials in Green CAPTTIM Category",  
     ylab = "Final Damage Score",  
     green.count.vec,  
     MA.summaryTest$mb.FD.200, col = "green")  
dev.off()
```

```
png(paste("CorrelationTestGreenAdvSelectBias.png"))  
plot(xlab = "Number of Trials in Green CAPTTIM Category",  
     ylab = "Advantageous Selection Bias",  
     green.count.vec,  
     MA.summaryTest$mb.adv.sb.200, col = "green")  
dev.off()
```

E. EXECUTE R SCRIPT

```
#Workflow  
rm(list=ls())  
setwd("~/NPS/Thesis/Thesis Data/Data Critz")  
source('RevisedMultiArm_Scrub.v13_Critz.R')  
source('RegretMeanPlots_Critz.R')  
source('RectangleFinalPlot_Critz.R')  
save.image('FinalDataScrub.RData')
```

THIS PAGE INTENTIONALLY LEFT BLANK

LIST OF REFERENCES

- Agrawal, R. (1995). Sample mean based index policies with $O(\log n)$ regret for the multi-armed bandit problem. *Advances in Applied Probability*, 27(4), 1054–1078, 1995.
- Bechara, A., Damasio, A. R., Damasio, H., & Anderson, S. W. (1994). Insensitivity to future consequences following damage to human prefrontal cortex. *Cognition*, 50(1), 7–15.
- Bell, D.E. (1982) Regret in decision-making under uncertainty. *Operations Research* 30(5):961–981. <http://dx.doi.org/10.1287/opre.30.5.961>
- Braimah, O. J., Osanaiye P. A., Omake P. E., Saheed, Y. K., & Eshimokhai, S.A. (2014). On the use of exponentially weighted moving average (Ewma) control chart in monitoring road traffic crashes. *International Journal of Mathematics and Statistics Invention (IJMSI)*, 2(5), 01–09.
- Crowder, S. V. (1989). Design of exponentially weighted moving average schemes. *Technometrics*, 21(3), 155–162.
- Djulgovic, B., Elqayam, S., Reljic, T., Hozo, I., Miladinovic, B., Tsalatsanis, A., Kumar, A., Beckstead, J., Taylor, S., & Cannon-Bowers, J. (2014). How do physicians decide to treat: an empirical evaluation of the threshold model. *BMC Medical Informatics and Decision-making* 2014, 14, 47.
- Fricker, R. (2010). *Introduction to Statistical Methods for Biosurveillance* (1st Ed.). New York, NY: Cambridge University Press.
- Kalgonda, A. A., Koshti, V. V., & Ashokan, K. V. (2011). Exponentially weighted moving average control chart. *Asian Journal of Management Research*, 2(1), 253–263.
- Kennedy, Q., Nesbitt, P., & Alt, J. (2014). Assessment of cognitive components of decision-making with military versions of the IGT and WCST. Human Factors and Ergonomics Society 2014 International Annual Meeting, Oct 27–31, Chicago Illinois.
- Killick, R., & Eckley, I. (2014). Changepoint: an r package for change point analysis. *Journal of Statistical Software*, 58(3). Retrieved April 15, 2015, from <http://www.jstatsoft.org/v58/i03/paper>

- Krain, A. L., Wilson, A. M., Arbuckle, R., Castellanos, F. X., & Milham, M. P. (2006). Distinct neural mechanisms of risk and ambiguity: a meta-analysis of decision-making. *NeuroImage*, 32(1), 477–484. doi:<http://dx.doi.org/10.1016/j.neuroimage.2006.02.047>
- Lopez, T. (2011). Odierno outlines priorities as army chief. *Army News Service*. Retrieved from <http://www.defense.gov/News/NewsArticle.aspx?ID=65292>
- Maidstone, R., Fearnhead, P., & Letchford, A. (2013). Changepoint detection using a segment neighborhood approach. Retrieved May 06, 2015, from <http://www.lancaster.ac.uk/pg/maidston/xmasposter2013.pdf>
- Nesbitt, P., Kennedy, Q., & Alt, J. K. (2015). *Iowa gambling task modified for military domain*. Monterey, CA: Operations Research Department, Naval Postgraduate School.
- Nesbitt, P., Kennedy, Q., Alt, J., Yang, J., Fricker, R., Appleget, J., Huston, J., Patton, S., & Whitaker, L. (2013). *Understanding optimal decision-making in wargaming*. Monterey, CA: TRAC-Monterey Technical Report, TRAC-M-TR-13–063, September 1, 2013.
- Odierno, R., & McHugh, J. (2015). Army human dimension strategy Draft Version 5.3. US Army Publication.
- R Development Core Team (2008). R: A language and environment for statistical computing. Vienna, Austria: R Foundation for Statistical Computing. Retrieved from <http://www.R-project.org>
- Roberts, S. W. (1959). Control chart tests based on geometric moving averages, *Technometrics*, 1, 239–250.
- Sutton, R., & Barto, A. (1998). *Reinforcement learning: An introduction*. Cambridge, MA: MIT Press.
- Taylor, W. (2000). Change-point analysis: a powerful new tool for detecting changes. Retrieved April 15, 2015, from <http://www.variation.com/cpa/tech/changepoint.html>

INITIAL DISTRIBUTION LIST

1. Defense Technical Information Center
Ft. Belvoir, Virginia
2. Dudley Knox Library
Naval Postgraduate School
Monterey, California