



NAVAL POSTGRADUATE SCHOOL

MONTEREY, CALIFORNIA

THESIS

**ACQUISITION OF A STATIC HUMAN TARGET IN
COMPLEX TERRAIN:
STUDY OF PERCEPTUAL LEARNING UTILIZING
VIRTUAL ENVIRONMENTS**

by

Robert L. Kammerzell

September 2008

Thesis Advisor:
Second Reader:

Michael McCauley
Joseph Sullivan

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

THIS PAGE INTENTIONALLY LEFT BLANK

REPORT DOCUMENTATION PAGE			<i>Form Approved OMB No. 0704-0188</i>	
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 September 2008	3. REPORT TYPE AND DATES COVERED Master's Thesis	
4. TITLE AND SUBTITLE Acquisition of a Static Human Target in Complex Terrain: Study of Perceptual Learning Utilizing Virtual Environments			5. FUNDING NUMBERS	
6. AUTHOR(S) Robert Kammerzell				
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.				
12a. DISTRIBUTION / AVAILABILITY STATEMENT Approved for public release; distribution is unlimited			12b. DISTRIBUTION CODE	
13. ABSTRACT (maximum 200 words) Soldiers conducting ground operations must visually detect various dynamic and static threats. While enemy utilization of improvised explosive devices (IEDs) is a constant danger, there is also the requirement to detect the insurgent sniper threat. The U.S. Army has long identified enemy sniper activity as one of great importance to both our individual soldier's survivability and unit operational effectiveness. Specifically, the soldier's visual system and perceptual skills are immediately tasked with categorizing both the environment and any detected threat. This study utilized game engine technology to assess the ability to train subjects in visual target acquisition within a complex virtual environment. The prevalence of computer games within the training realm requires study as to the game engine's ability to support current operations and soldier training. The study's results determined that training improved a subject's target Hit Rate percentage 29% ($p = .0001$), in comparison to the control group, when presented scenes of increased difficulty. Historically, military-themed computer games have succeeded in providing strategic training value. This study indicates that military themed computer games also assist with individual soldier skills training.				
14. SUBJECT TERMS Simulation and Training, Human Systems			15. NUMBER OF PAGES 105	
			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	

NSN 7540-01-280-5500

Standard Form 298 (Rev. 2-89)
Prescribed by ANSI Std. Z39-18

THIS PAGE INTENTIONALLY LEFT BLANK

Approved for public release; distribution is unlimited

**ACQUISITION OF A STATIC HUMAN TARGET IN COMPLEX TERRAIN:
STUDY OF PERCEPTUAL LEARNING UTILIZING VIRTUAL
ENVIRONMENTS**

Robert L. Kammerzell
Major, United States Army
B.S., Colorado State University, 1993
MPA, University of Oklahoma, 2006

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
September 2008**

Author: Robert L. Kammerzell

Approved by: Michael McCauley
Thesis Advisor

Joseph Sullivan
Second Reader

Mathias Kolsch
Chair, MOVES Academic Committee

THIS PAGE INTENTIONALLY LEFT BLANK

ABSTRACT

Soldiers conducting ground operations must visually detect various dynamic and static threats. While enemy utilization of improvised explosive devices (IEDs) is a constant danger, there is also the requirement to detect the insurgent sniper threat. The U.S. Army has long identified enemy sniper activity as one of great importance to both our individual soldier's survivability and unit operational effectiveness. Specifically, the soldier's visual system and perceptual skills are immediately tasked with categorizing both the environment and any detected threat.

This study utilized game engine technology to assess the ability to train subjects in visual target acquisition within a complex virtual environment. The prevalence of computer games within the training realm requires study as to the game engine's ability to support current operations and soldier training.

The study's results determined that training improved a subject's target Hit Rate percentage 29% ($p = .0001$), in comparison to the control group, when presented scenes of increased difficulty. Historically, military-themed computer games have succeeded in providing strategic training value. This study indicates that military themed computer games also assist with individual soldier skills training.

THIS PAGE INTENTIONALLY LEFT BLANK

TABLE OF CONTENTS

I.	INTRODUCTION.....	1
A.	PROBLEM STATEMENT	1
B.	RESEARCH QUESTIONS.....	3
C.	SCOPE	4
D.	ORGANIZATION OF THE THESIS.....	4
II.	LITERATURE REVIEW	7
A.	PERCEPTUAL LEARNING AND VISUAL INTELLIGENCE	7
1.	Visual Construction and How We See	8
2.	Perceptual Learning	15
3.	Expertise and Training.....	20
4.	Perceptual Learning in the Training Environment.....	25
B.	IMPLICATION AND HYPOTHESIS.....	27
III.	METHODOLOGY	29
A.	PARTICIPANTS AND DESIGN	29
1.	Hypotheses.....	29
2.	Design.....	29
3.	Measures	30
4.	Control Procedures.....	30
B.	EQUIPMENT.....	31
C.	VISUAL SCENES.....	31
D.	EXPERIMENT DESIGN AND PROCEDURES.....	34
1.	Participant Screening	35
2.	Experiment Procedures.....	36
IV.	RESULTS	41
A.	DEMOGRAPHICS.....	41
B.	ANALYSIS OF HIT RATES.....	41
1.	Results for Hit Rate by Difficulty and Group	42
2.	Questionnaire User Feedback.....	46
C.	POST-HOC DISCUSSION OF HYPOTHESES	47
V.	DISCUSSION.....	49
A.	PERCEPTUAL LEARNING IN VIRTUAL ENVIRONMENTS.....	49
1.	Perceptual Learning Ramifications.....	50
2.	Perceptual Learning influence on Feedback and Training- Feedback Groups	52
B.	AREAS OF FUTURE RESEARCH.....	53
	APPENDIX A.....	55
A.	LABORATORY PROCEDURES	55
B.	PRETEST QUESTIONNAIRE	60
C.	POST-TEST QUESTIONNAIRE – NO FEEDBACK	61
D.	POST-TEST QUESTIONNAIRE – TRAINING-FEEDBACK.....	62

E.	POST-TEST QUESTIONNAIRE - FEEDBACK	63
APPENDIX B		65
A.	NORMAL COLORED SCENE EXAMPLES.....	65
B.	FALSE (FEEDBACK) COLORED SCENES EXAMPLES.....	66
APPENDIX C		67
A.	GRAPHICS RESEARCH PRIOR TO EXPERIMENTATION	67
1.	Appendix C Abstract	67
2.	The Graphics Pipeline	68
3.	Graphics Cards Design and Marketing	72
4.	Display Technology	74
5.	Conclusion	77
APPENDIX D		79
A.	EXPERIMENT LESSONS LEARNED	79
1.	Selection of Engine	79
2.	Minimap Utilization.....	79
LIST OF REFERENCES		85
INITIAL DISTRIBUTION LIST		89

LIST OF FIGURES

Figure 1.	America’s Army Screen Shot, 2008	2
Figure 2.	Phenomenon of “good continuation” (From Field, Hayes, & Hess, 1993)	10
Figure 3.	Robust Scene (From Field et al., 1993)	11
Figure 4.	Ear Outline Memory Stamp (From Field et al., 1993).....	11
Figure 5.	Target Stimulus Comparison (From Field et al., 1993).....	12
Figure 6.	Fallujah, Iraq.....	14
Figure 7.	Ambiguous pictures and perceptual learning (From McLeod, 2007).....	15
Figure 8.	Leeper’s Gestalt Test (From E. J. Gibson, 1969)	17
Figure 9.	Progression from Novice to Expert (From Wickens, Lee, Liu, & Becker, 2004)	21
Figure 10.	Contributing Roles to Expertise Development (From Wickens et al., 2004) ..	22
Figure 11.	Example Snellen Chart (From Wikipedia, 2008)	35
Figure 12.	Example Ishihara Plate (From Wikipedia, 2008)	36
Figure 13.	Normal Colored Target Scene # 1009	38
Figure 14.	False Colored Feedback Scene #1009	38
Figure 15.	Graph: Hit Rate by Difficulty and Group	43
Figure 16.	ANOVA Hit Rate by Group: Hard	45
Figure 17.	ANOVA Hit Rate by Group: Easy.....	46
Figure 18.	Normal Colored Scene: Easy	65
Figure 19.	Normal Colored Scene: Hard.....	65
Figure 20.	Feedback Colored Scene: Easy	66
Figure 21.	Feedback Colored Scene: Hard.....	66
Figure 22.	Shading and Textures (From Woligroski, 2006)	74
Figure 23.	Spatial Perception (From Woligroski, 2006)	75
Figure 24.	Basic HMD Design(From Weinand, 2005)	76
Figure 25.	3D Display (From Weinand, 2005)	77
Figure 26.	Minimap Centering and Render Error Example	80
Figure 27.	Minimap Scene Render Error Example	81
Figure 28.	Correctly Rendered Scene Comparison to Fig. 27.....	82
Figure 29.	Z-buffer Error Comparison (From Wikipedia, 2008)	83

THIS PAGE INTENTIONALLY LEFT BLANK

LIST OF TABLES

Table 1.	Targets by Scene and Difficulty	34
Table 2.	Demographic Information by Group	41
Table 3.	Hit Rate by Difficulty and Group (Mean and Standard Deviation).....	42
Table 4.	Tukey-Kramer Comparison by Group: Hard.....	44

THIS PAGE INTENTIONALLY LEFT BLANK

ACKNOWLEDGMENTS

Sandra Gail, my Love, my Wife, my Best Friend. Your family support and Faith ground our family and made this NPS experience possible.

Darla and Klara, there are no words for how much I love and appreciate you. Thank you for forgiving me time and again.

Dr. Michael McCauley, my Mentor and Coach. Your professionalism is without peer.

CDR Joseph Sullivan and the MOVES Institute Human Factors instructors, you keep us all grounded in why we are here. Train soldiers and Defend the nation.

Major Derek Sebalj, Canadian Forces. Without your friendship and tutelage, I would not have had the education necessary to undertake this endeavor. I would also be out of shape. Rider On!

LT Patrick Jungkunz, German Navy. My friend and a constant help.

CPT Christopher McClernon, U.S. Air Force. Your ability to relate theory and statistics to a knuckle dragger is most appreciated.

Dr. Christopher Darken. Your support and knowledge are appreciated more than you know.

Dr. Susan Sanchez, your guidance and enthusiasm through the results analysis made the statistics “JMP” off the page.

THIS PAGE INTENTIONALLY LEFT BLANK

I. INTRODUCTION

A. PROBLEM STATEMENT

Soldiers conducting patrols and other ground operations face the challenge of detecting various static threats in a visually complex environment. While the enemy utilization of improvised explosive devices (IEDs) is a constant danger, the visual transition to detecting threat enemy personnel conducting sniper operations is also important to individual soldier survivability and unit operational capability.

As the War on Terror continues in Iraq and Afghanistan, many soldiers, as well as individual augmentee (IA) personnel, have not yet deployed into those theaters of operation. Soldiers and IA personnel face the challenge of internalizing new stimuli from which they will be required to utilize visual cues to detect a static threat target (i.e., a sniper) within an unfamiliar environment.

While there are situational awareness systems under development that may assist the soldier to more quickly determine the target location, these systems provide that information only after the soldier is fired upon. The soldier is then tasked with the visual identification of the targets location within the environment. The soldier's inherent biological systems are tasked with identifying the enemy combatant, given all visual stimuli, within the operational environment based upon an internalized set of perceptual skills, memories, and training.

As soldiers are afforded greater capabilities through weapon effects, communications and personal protection, the human perceptual system is still greatly relied upon to determine possible target locations throughout the visual scene. The human ability to "see" in-depth is challenged when presented with new information or by lack of previous experience.

While game engines have been utilized and studied with regards to cognition, team training and collective tasks, they have not been explored with concern to individual ground operations visual skill training. Specifically, do visual target acquisition skills, given game engine technology, support perceptual learning within a virtual environment?

The U.S. Army has explored utilization of virtual environments with regards to public awareness and recruiting via games and simulations such as America's Army. Launched in 2002, America's Army allows players, bound by Rules of Engagement (ROE), to interact and grow in experience as they work with other players and maneuver through different mission scenarios. These mission scenarios are also described as having training applications for use within the military and government sectors and an "incredible Army Experience" (Figure 1) (America's Army, 2008).



Figure 1. America's Army Screen Shot, 2008

The developers state that they will continue to develop the "game" to replicate the dynamic nature of soldiering and allow players to explore "the Army of today, tomorrow and the future."

America's Army provides civilians with insights on Soldiering...from the battlefields by sending the [designers and developers] to crawl through obstacle courses, fire weapons, observe paratrooper instruction, and participate in a variety of training exercises with elite combat units, all so that you could virtually experience Soldiering in the most realistic way possible (America's Army, 2008).

Previous studies have utilized game engine technology to assess the accuracy to which artificial intelligence algorithms replicate visual target acquisition but not the extent to which game engines might support training military personnel new to ground combat.

The game engine utilized for this research is Delta3D. Delta3D is an Open Source engine that is used for games, simulations, and other graphical applications. Delta3D is fully featured and suitable for a variety of uses including training, education, visualization, and entertainment. Delta3D is unique because it offers features appropriate to the Modeling and Simulation and DoD communities such as High Level Architecture (HLA), After Action Review (AAR), and large scale terrain support. Its modular design integrates other well-known Open Source projects such as Open Scene Graph, Open Dynamics Engine, Character Animation Library, and OpenAL. Delta3D is developed and tested on Windows XP using Microsoft Visual Studio and Linux using GNU Compiler Code (GCC). All the underlying dependencies are cross-platform as well, so just about any platform should be compatible with a few minor modifications to the source. Delta3D is released under the GNU Lesser General Public License (LGPL) (McDowell, 2008).

B. RESEARCH QUESTIONS

Within the operational community, the nagging question remains how to train young soldiers better so they are able to be not just victorious, but that they return home unharmed.

This research will investigate whether game engine technology improves a soldier's ability to readily detect single or multiple static threats. Training target detection, in this case enemy snipers, is influenced by instruction that has a foundation in perceptual learning theory. The questions that specifically motivated this thesis work include:

1. To what extent can the novice soldier be trained within a virtual environment to identify an enemy sniper given that the soldier can only see part of the enemy soldier or some part of the enemy's equipment?

2. Does the individual soldier internalize his visual experiences and training? Is one soldier more likely to visually acquire a target over another?
3. How can we rapidly build an individual soldier's visual mental models and perception in order to improve soldier survivability and effectiveness?

This thesis evolved from an academic curiosity influenced by the Naval Postgraduate School's Human Factors and Human System Integration classes and a desire to explore a "common sense" theoretical concept, that is often ignored, but of great importance to both the soldier and the staff officer. Additionally, as a U.S. Army Simulations Operations officer, I wanted to know if game engines might be utilized to leverage training within a popular media application.

C. SCOPE

The scope of this thesis will focus on addressing question one, above, with a literature review encompassing questions two and three. Consideration was given to how virtual environment technology might be applied to promote perceptual learning for target detection. Specifically, this thesis will utilize a game engine rendered virtual environment to assess the extent to which the individual ability to acquire a visual target is influenced through training and augmented by technology.

This study attempts to determine if there is opportunity to apply perceptual learning theory to enhance individual soldier readiness and training to decrease the soldier's initial risk within the operating environment and improve the soldier's likelihood of visually acquiring a static human target.

D. ORGANIZATION OF THE THESIS

- Chapter I: Introduction. This chapter discusses the motivations for this research, states the research questions, and establishes a framework for pursuing the stated questions.
- Chapter II: Literature Review. This chapter provides a review of the current literature that includes perceptual learning, visual intelligence, and training.

- Chapter III: Method. This chapter describes the design of the experiment and laboratory procedures utilized.
- Chapter IV: Results. This chapter describes the quantitative and qualitative results of the experiment.
- Chapter V: Discussion and Recommendations. This chapter discusses findings, explores the relationship between results and theory, and recommends future research.

THIS PAGE INTENTIONALLY LEFT BLANK

II. LITERATURE REVIEW

This chapter reviews literature providing a theoretical foundation for perceptual learning, visual intelligence and training. Perceptual learning is a concept that gained academic popularity in the 1950s and prominence as a field of study in the latter 1960s. Visual intelligence, a term coined by Donald Hoffman, addresses the overall theoretical concept that there are inherent “rules” that all sighted beings utilize and that these rules incorporate experience and instinct to facilitate the ability to “see” (Hoffman, 1998).

A. PERCEPTUAL LEARNING AND VISUAL INTELLIGENCE

“A particular problem for psychologists is to explain the process by which the physical energy received by sense organs forms the basis of perceptual experience. Sensory inputs are somehow converted into perceptions of desks and computers, flowers and buildings, cars and planes; into sights, sounds, smells, taste and touch experiences” (McLeod, 2007). Perceptual learning refers to a person’s ability to gather new information from the surrounding environment and then integrate the information to an existing body of knowledge stored internally.

Scientists and authors often explore two competing perception theories. One theory espouses that the individual collects all data from the external environment and then forms a perceptual construct of the scene. This “Constructivist Theory” states that there is contextual recognition through pattern recognition. Individual improvement for Constructivists is dependent upon the individual learning to gather the correct data from the stimuli provided. The competing theory, most often called “Direct Theory,” states that perception is data driven and as stimuli is collected at the eye it then proceeds to the visual cortex and is categorized and stored for future comparison and use. Improvement through Direct Theory is dependent upon the development of skills and the ability to correlate information within the memory cells to the stimuli provided (E. J. Gibson, 1969)(J. Gibson, 1969).

1. Visual Construction and How We See

This research requires us to define, in some detail, how soldiers naturally extract critical information visual scenes and target objects. As soldiers progress through their initial training period, the majority of the instruction presented is focused upon individual skills such as marksmanship, drill and ceremony and basic technical skills. Even for ground combat military specialties, the level of field craft and field training exercises are designed so the soldier is capable of following the orders of superiors quite ably. Unfortunately, there remains a lack of target detection skills beyond basic vehicle and the E-type silhouette recognition utilized on the rifle range (G-3, Directorate of Operations & Training, & Current Operations, 2008). As instructors attempt to better equip service members with the capability to operate in a challenging environment, we must first identify the types of perceptual challenges that are faced during everyday ground operations.

Currently, there is much discussion as to how our soldiers can more readily survive the “first 100 days” with greater likelihood of operational success and less likelihood of injury. This first 100 days is commensurate with how the air combat community regards the first few air engagements as a test of how likely a pilot is to survive, and ultimately flourish, within that particular environment. Assignments, tasks, and mission scenarios will challenge the individual’s abilities to extract relevant sensory information from the operational environment.

Ground combat troops must operate in dynamic situations in which the threat changes frequently. While the threat of IEDs is currently the most dangerous and common, there is the additional threat of increasing numbers of enemy snipers. Historically, a sniper encounter has its own deleterious impact upon a unit’s operational capability and morale. It is the sniper threat that is addressed during this research (CALL, 2007) (J. Gibson, 1969).

We are interested in the capacity of new ground forces to identify the threat presented within the operational environment. As the service members’ initial training for their specific military specialty occurs at a military facility and their follow-on home-

station in the United States, there is high likelihood that their internal recognition mechanisms are not yet sufficient to readily identify an enemy threat within the new operational environment.

To first determine what a threat target looks like, humans must have some construct or object imbedded within their memory that facilitates a mental representation. “You...carve your visual world into parts, as a critical step in your construction and recognition of visual objects [and] you can quickly and effortlessly assemble many parts into many objects [by] using color, motion, shape, texture and PRIOR EXPERIENCE” (Hoffman, 1998). These prior experiences are only made possible by two events in any individual’s life. Either through education and training or through the “real life” experience itself. One objective is to determine how to assist the service member to learn how to rapidly acquire static targets within the operational environment. Let us, however, continue to discuss how the service member’s human capacity for mental construction of a scene is facilitated by vision and cognition.

Experiences begin the first moment that the newborn opens its eyes.

[T]hese rules are universal in the sense that all normal kids have the same rules, then although these rules blind them to many possibilities, these rules can also guide them to construct visual worlds about which they have consensus. Two toddlers from different sides of the globe can be shown the same novel image and see, in consequence, the same visual scene. These innate rules, which lead to consensus in the visual constructions of all normal adults despite the infinite ambiguity of images [are] *the rules of universal vision* (Hoffman, 1998).

In many regards, this can be equated to Direct Theory reasoning that would resemble more of an instinctive application of mentally rendering the visual scene. Understanding why one service member can more readily identify a specific target in comparison to another demographically similar service member when presented identical operational environment information is at the core of perceptual learning. For example, while both service members can readily identify major components such as street, building, window and rubble, common sense dictates that both should be capable of identifying a static human target within the environment. Why do individual differences

exist? “One of the peculiarities of cognitive psychology is that it is a science of the unobservable. We can never directly read off the internal symbols or program steps of the mind...[the] purpose of modeling is not just to explain isolated pieces of data, but to discover general principles of the organization of the cognitive processes and, ultimately, a description of the complete processing system” (Matthews, Davies, Westerman, & Stammers, 2000). Both individuals have been exposed to visually interpreting human interaction with their common everyday environment; therefore, both should be capable of identifying the enemy combatant standing inside of the building in the scene presented. Most agree that is not the case and merely equate the difference to “common sense” that people are just different and some are better than others at “seeing” the target. This is undoubtedly true, but how then do we address the issue that due to these differences an understanding that they warrant systematic address of some sort? One key to this issue is understanding how our mind determines what we see. “A key to your success is that you efficiently divide shapes into parts, and describe both the parts and their spatial relationships. You can do this before you know what the objects are. Once you have the parts and their spatial relations you use this information to search the vast list of objects you know, until you find a match” (Hoffman, 1998).

Many studies have considered perception and visual processing. The ability to identify an object that is only partially visible or non-contiguous has been assessed by multiple studies (see Figure 2).

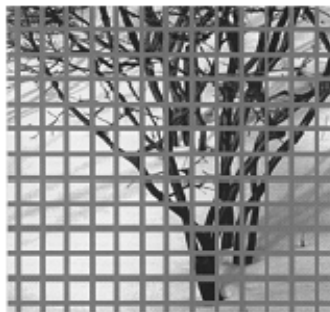


Figure 2. Phenomenon of “good continuation” (From Field, Hayes, & Hess, 1993)

David Field, et al. (1993), conducted a study that assessed an individual's ability to visually acquire a cortical cell path from a background image that incorporated the path into a more robust scene. If a subject, such as you the reader, is tasked with locating the path below that resembles the outline of a human ear, the subject may be successful after conducting a thorough visual search. Based upon the types of memory object files that are located and utilized, the search may or may not be fruitful as the scene is extremely noisy (Figure 3).

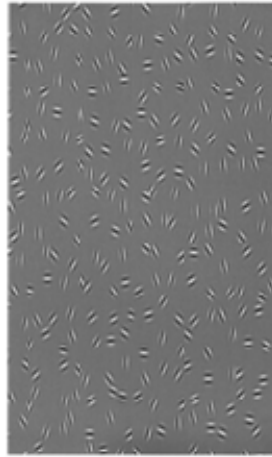


Figure 3. Robust Scene (From Field et al., 1993)

Through memory stamping the object image into the subject's memory folder, Field determined that the visual memory of the image more readily facilitated the acquisition of the target within the robust scene (Figure 4) (Field et al., 1993).

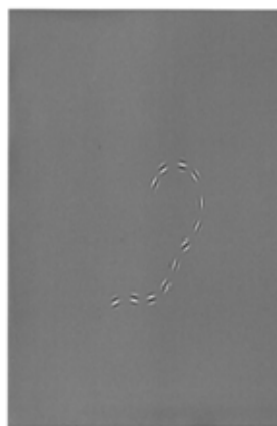


Figure 4. Ear Outline Memory Stamp (From Field et al., 1993)

Now that the subject, you the reader, have the visual representation of the target object presented to you, note the increased ease with which you can locate the desired target given Figure 3 and Figure 4 in a side-by-side comparison (Figure 5).

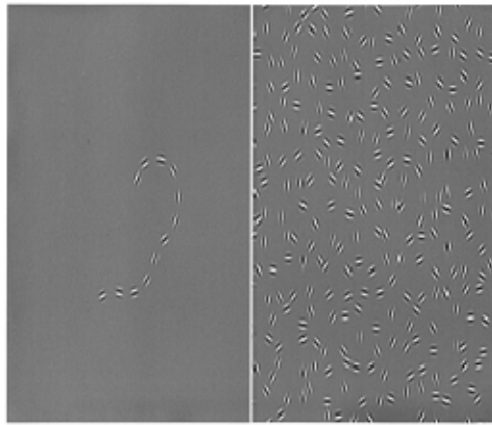


Figure 5. Target Stimulus Comparison (From Field et al., 1993)

Other studies have assessed the ability of the observer to use visual information to match the object image to a specific 3-D classification. “Broadly speaking, visual information has two parts: the image measurements (or features), and the target (or signal) of interest. Image information about the signal, such as 3D description of a giraffe, is typically confounded by uncertainties (or noise) introduced by rendering and projection... image details such as the precise shapes of the spots on the giraffe or the leaves on the bushes are unimportant; what matter are features such as neck length and body size. However, if the task were identification of either the particular giraffe or the kind of bush, then the shapes of the spots and leaves would become important (Olman & Dersten, 2004). This point is very important as it relates directly back to a determination of what the target is and an assessment of the mental models that the service member has available within the “object folder” of their memory. Only then one assess the service member’s capacity to readily access the information from that “file” and apply it to the presented visual scene.

Once this base line assessment is realized, we can begin to understand the necessity to address the perceptual comparison that must be made within the service member’s cognition of the visual scene. As the service member builds a working

memory of objects that fit within the various mental folders that allow for visual scene construction, those objects must begin to conform to rules that will determine their definition and role within the service members operational environment. The ability to assign an individual “value” to each object is extremely important as the service member will be required to rapidly determine the threat location given minimal or austere information.

Kroger et al. (2004), conducted an investigation as to the impact of relational complexity on perceptual comparisons that is at the heart of the issue of visual scene interpretation.

The capacity to recognize that two objects, situations, or events are the same with respect to a certain criterion underlies object recognition, categorization, and analogical reasoning...decisions at the higher levels may be inter-dependent with those at lower levels, in the sense that information from multiple levels contributes to a decision at the nominal level defining the decision. In other words, information about the stimuli at the designated level and all lower levels might be available for processing, and be utilized by participants in determining a response (Kroger, Holyoak, & Hummel, 2004).

The issue then becomes that if we are cognizant of what the target is we are supposed to look for, why are we challenged to locate this specifically defined target in changing environments? The aviator that readily locates a threat aircraft or ground target given a basic description can do so most readily, as indicated by years of research, yet is challenged to notice the partial torso silhouette of a enemy combatant hiding behind a rubble pile (Ciavarelli et al., 2005). The root of this issue is the manner from which our mind locates objects based upon the visual information presented and then searches for the best fit of objects available for labeling the composite parts of the scene presented. If there is no best fit, given the specificity of target type, then target detection becomes an individual construct of the presented environmental stimuli. “You construct visual worlds from ambiguous images in conformance to visual rules” (Hoffman, 1998). The rules that are utilized are those rules that are formed based upon the experiences and training that formed the list of objects from which the service member can draw. Many argue that there is no reason that Navy pilot service member A is unable to perform

duties that include movement through a hostile environment if duly warned and instructed of the type of threat that exists. This, to some extent is true, as their vigilance will undoubtedly increase but, the fact remains that the service member will not be able to translate their expert ability in identifying a ground target from the air to the near-same ability to identify a threat target from ground level. “The answer, in part, is that you construct them according to rules. You can’t do to them what you wish if what you wish violates your rules of construction. Your rules allow you to construct what you see, but they also restrict what you can construct and what you can do with your constructions” (Hoffman, 1998). Until the service members’ visual objects list are amended, there is high likelihood that their original rules of target acquisition will over-ride their initial visual scene construction until the individual memory file is amended via operational experience (Figure 6). Thus, the first time a subject is presented with new visual information, the target detection task may require a significant amount of mental workload if the subject is to perceive and detect a threat within the environment.

As Kroger determined, human visual search will continue until a match is found that the searcher determines is satisfactory. Once a mental match of objects is made, scene construction is discontinued. “Thus, at least a minimal amount of serial, “bottom-to-top” processing seems to be intrinsic to the task. The self-terminating account is more specific in that, in addition to serial “bottom-to-top” processing, it also holds that subjects always stop processing at the lowest level at which they find a match” (Kroger et al., 2004).



Figure 6. Fallujah, Iraq

2. Perceptual Learning

“The criterion of perceptual learning is thus an increase in specificity. What is learned can be described as detection of properties, patterns, and distinctive features” (E. J. Gibson, 1969).

A key aspect of perceptual learning is clearly identifying that which is missing from the construct file of objects the individual has available. Thus far the discussion of visual intelligence has described the foundation upon which this research is laid. There are inherent rules that we all utilize to construct the scene that represents our environment. Some researchers claim that upwards of 90% of the information we assimilate visually is lost by the time it reaches the brain. We can assume, then, that the brain must somehow fill in the gaps of information with stored objects from memory in order for us to identify the environment and other objects included within that environment (McLeod, 2007).

Briefly study the random arrangement of black shapes in Figure 7 and try to determine the subject of the photograph.



Figure 7. Ambiguous pictures and perceptual learning (From McLeod, 2007)

Once the face becomes evident, it is nearly impossible to look away from the photograph and then not see the face upon re-inspection. This constructivist phenomenon

is well documented throughout perceptual literature and is not necessarily surprising, but what is interesting is not just our ability to eventually distinguish the figure, but then to have that ability to match the presented random arrangement to a face more quickly each time we glance at the image.

Once again, our visual intelligence guides us to construct the scene based upon the information at hand and the objects on file within our memory. The challenge is to determine what parts are necessary in order to empower the service member with having access to the requisite parts to complete the building of the scene.

To be useful your parts should satisfy at least four conditions...First, they shouldn't change if you move your view a bit. You need stable parts for stable recognition. Second, they shouldn't change if the object changes its configuration a bit. Again, stable description is the key. Third, you should be able to construct the parts from the retinal images at your eyes. If you can't do it from images you can't do it at all. Fourth, you should be able to construct the parts on a wide variety of objects; the larger the better. If your scheme for constructing parts is not general-purpose, then it might fail you at critical times. If you missed that tiger because you couldn't see its parts, it could ruin your whole day (Hoffman, 1998).

These shapes are very important to how individuals construct their visual scene. As the service member progresses through training, from the first moment of instruction to their final training event prior to transferring to their first unit, they are building an object file specific to the military specialty that they will serve. Hoffman's insight as to how these objects will serve towards perceptual learning are very important as we discuss how visual intelligence and perceptual training are considered within this body of research. If the service member is challenged with identifying a threat based upon identification of only a "part" of the object, the threat can, just as Hoffman stated above, ruin the service member's day.

Historically, studies have proven that once the part image presented has encoded itself to a memory object in the brain, that image is then encoded into the object memory file and can be used to identify the complete object. One such study is that of R.A. Leeper's study of a neglected portion of the field of learning which studied the development of sensory organization.

One of Leeper's experiments employed incomplete figures. They were exposed to groups of subjects under various conditions and with differing amounts of exposure. All the helps given the subjects – repetition, exposure of the complete picture from which the fragmentary one was taken, and knowledge of the class to which the imperfectly portrayed object belonged assisted in achieving the correct reorganization. Once achieved, it was retained (E. J. Gibson, 1969).

As depicted in Figure 8, some of the objects are easily determined, such as the plane, the school bus, the woman holding a dog on her lap and the viola. Interestingly, some other objects must be concentrated upon to be definitively named as a cement truck, a shoe, an old style dual bell alarm clock, a cartoonish alligator, and what appears to be a raccoon rear paw print.



Figure 8. Leeper's Gestalt Test (From E. J. Gibson, 1969)

This ability to identify a visual construct from minimal information is considered a “bottom up” reasoning skill in that there is little reasoning involved by the individual making the assessment. In much the same manner the challenge that the service member faces, when identifying a threat target within the operational environment, is to rapidly identify an object given minimal visual data.

In the example of Figure 8, most people are able to identify many of the objects presented by the pieces as they represent common everyday objects that are encountered not just repetitiously, but in varied environments. The individual then, is able to quickly and without in-depth thought processing piece the individual shapes together to form the visual construct of the object pictured.

By carving an object into parts, and using the visible parts to search your memory for a match, you can recognize it from many different views...[Many] parts are not rigid. If they move, [then their] configuration changes. How shall you recognize a body despite such changes? Again parts come to the rescue...[as] the parts won't change as the configuration does. This gives you a stable description of objects and an **efficient index into your memory of shapes** (Hoffman, 1998).

This research is important in that, much the same as the previously discussed McLeod picture (Figure 7), the subjects were able to identify the greater object based upon the disconnected objects presented that then mapped to a singular larger defined object within the subjects memory. In much the same manner the service member will utilize objects in memory to complete the identification of the threat.

The research and theory of perceptual learning are the theoretical foundations established by Eleanor J. Gibson are presented throughout this thesis. Dr. Gibson's work in perceptual learning has been heavily utilized within the field of education.

Perceptual learning then refers to an increase in the ability to extract information from the environment, as a result of experience and practice with stimulation coming from it. Adaptive modification of perception should result in better correlation with the events and objects that are the sources of stimulation as well as an increase in the capacity to utilize potential stimulation...this definition describes an end result, admittedly, rather than a process (E. J. Gibson, 1969).

Determining the process that will provide the correct information to the service member is of great importance as the objects provided will not only increase the number of object parts within the memory file but, undoubtedly, also determine the response of the service member given the visual stimuli. As stated by Gibson, detection of distinctive features "are differentiated (and thereby identified) by their distinctive features. These features are not constructed by the mind but are discovered by the perceiver. When he is exposed to a new set of objects, what he learns are the distinctive features of each object and of the set" (E. J. Gibson, 1969).

The skills involved to identify a target within a new environment require the service member be afforded opportunity to gain the visual skills requisite with perceptually constructing the scene. “Skill training requires both instruction and the opportunity to practice the task” (Ciavarelli, Asbury, Salinas, Hennesy, & Sharkey, September 2005). The more opportunity that the service member has to practice the task the more likelihood there is that the object parts required to properly identify a threat target are present within memory. Without experience or training, that specifically addresses what the target and target parts look like, visual construction of the scene is highly unlikely.

The visual scene from which the service member must construct his threat identification is dynamic. The role of perceptual learning is to identify that specific target information that the service member does not have currently stored in memory, but instead use the partial information to visualize and construct his mental image of the scene. As stated by Gibson:

Perception is not passive reception. It is active search... Perception is furthermore adaptive and regulatory. It focuses on wanted stimulation and rejects the rest... So perception actively selects and rejects. The search process is adaptive, and I think self-regulatory. From too much available information, it extracts what is salient. From confusion and uncertainty, order, differentiation, and economy are achieved (E. J. Gibson, 1969).

In this manner, the application of visual intelligence can be associated with perceptual learning. As the service member’s internal visual intelligence “rules” are applied to the visual stimuli collected, the service member will naturally begin to construct the mental image of the threat within the scene based upon the partial object information at hand. Through perceptual learning, additional presentations of various operational environments will inherently strengthen the service members’ perceptual ability to identify not just the object part, but the whole object with which the part is associated. There is likelihood that the visual presentation of multiple threat target orientations will facilitate the service member’s ability to identify threat targets given partial presentation through perceptual learning.

I think, therefore, that perceptual learning is taking out from the total stimulus information whatever is invariant about the world, whether it be a distinctive feature or a rule... [the] evidence [indicates] that distinctive features learned in the course of discriminating differences play a role in production as well as in the formation of memory representations (E. J. Gibson, 1969).

Just as Figure 1 presented a tree behind a gridline that was easily distinguished, by most observers, the service member is better equipped to identify a partially presented threat target. The partial presentation to a novice is often not enough to guarantee identification whereas it may be sufficient to allow a practiced observer to identify the threat. The reason is that the practiced observer has adequate internal imagery to complete the visual construction of the scene based upon previous perceptual learning via their experiences and training. Additionally, the practiced observer may have experience that allows specific information search of the environment required to make a target detection merely based upon expectation.

I assume that to visualize a geometrical solid cube is to apprehend its *properties* (e.g. rectangular dihedral angles) and that to visualize a person's head is to know its *features* in both full face and profile. Contrariwise, an "image" of a cube or of a face must be a perspective view of it, either a momentary or a frozen picture, and an "image of memory" can only be the trace of a sensation, i.e., of a retinal image. What is "stored" then, can only be a picture, an engram, an impression. But the kind of memory we refer to as *visualizing* seems to be a knowing of invariants under perspective transformations over time; an awareness of *formless* invariants. To remember is *not* to search through the file of snapshots stored in the brain! There is imageless memory just as there is sensationless perception. To "imagine" is *not* to have an *eidetic* image, an *after* image, or a *pictorial* image, nor is it to *represent* something to oneself. What is it (J. Gibson, 1969)?

3. Expertise and Training

Previous research indicates that there are three generally agreed upon levels of aptitude regarding expertise. Initially the individual has a declarative knowledge of the skill involved and the visual information is not well organized and cannot be employed efficiently (Hoffman, 1998). As the individual gains experience and practice over time,

he progresses to a procedural level of efficiency and employs a more analytical (if – then style of cognitive ability) approach towards the task. If given enough time, practice and experience, the user will then approach a level of automaticity that is employed but what is often described as the expert. As is evident from Figure 9 below, the amount of time required to acquire new skill is not necessarily fluid, nor can it be easily defined with regards to the amount of time each individual needs to progress between levels.

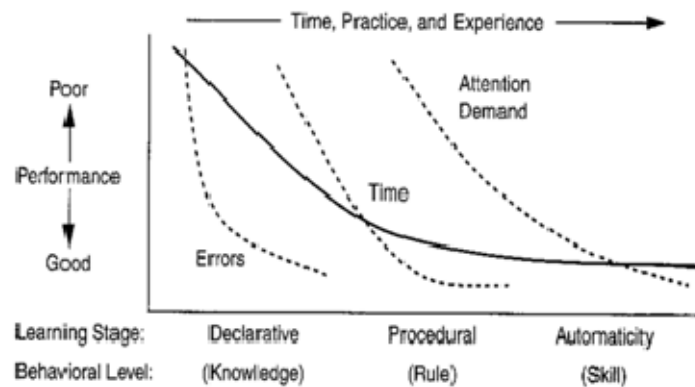


Figure 9. Progression from Novice to Expert (From Wickens, Lee, Liu, & Becker, 2004)

As described by Wickens, et al.,

These three stages generally follow upon each other gradually, continuously and partially overlapping rather than representing sudden jumps. As a consequence, performance in the typical skill improves in a relatively continuous function...the rate of reduction of errors, time, and attention-demand varies from skill to skill...[and] some complex skills may show temporary plateaus in the learning curve (Wickens et al., 2004).

In order to overcome the individual differences the utilization of training aids has been applied to various types of training. Task and function analysis identifies desired skills and allows focused training that improves the novice's skills through exposure to specific information. Training aids, when properly designed and utilized, allow the novice to focus on the task at hand and become more proficient in a shorter period of time (See Figure 10).

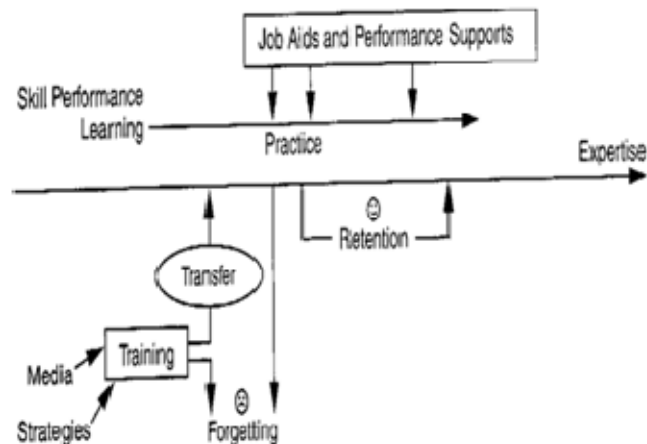


Figure 10. Contributing Roles to Expertise Development (From Wickens et al., 2004)

Perceptual learning facilitates visual target acquisition by making objects in our memory available through their inclusion within our operational environment. As with many other rote style tasks, our visual objects memory is populated in a similar manner as that which allows the identification of symbols for letters and numbers. The challenge with target identification is allowing the service member ample opportunity to determine how the operational environment is inclusive of both common objects and threat objects with respect to their level of expertise. If expertise can only be obtained through operational experience, the service member and those around him are at increased risk. Many studies have concluded that one of the best ways to increase the likelihood of individual success, with regards to visual search, is to provide ample training time that reinforces success.

Overall, the current results show that memory for the visual detail of natural scenes accumulates over time and across separate glances. Memory for the visual details of objects in a scene showed a consistent, linear increase over time. It is important to note that this linear accumulation was the same after viewing a scene once, in a single long trial, or after viewing that scene twice for the same total duration (Melcher, 2006).

Training and expertise are more directly tied to perceptual learning as previous studies have discovered that “observers in the current experiments were actively engaged in learning about the scene, hence their attention was likely drawn to all the items as potential targets for a memory test. In real life, attention may be given primarily to task-relevant items, with ignored objects failing to be retained in memory” (Melcher, 2006).

As individuals are presented relevant new objects, that information is stored in short term memory (STM) until such time as the individual passes it into long term memory (LTM) stores. Research indicates that only 4 objects are stored within the visual short-term memory (VSTM), which will likely influence perceptual learning and the individual’s ability to more rapidly acquire the visual target acquisition skill. Educators and trainers often agree that “practice makes perfect” and research also indicates that, with respect to perceptual learning, this statement holds true for visual skills acquisition as well. As stated by David Melcher:

It is interesting to note that change detection actually showed three distinct levels of performance as the number of items viewed between the target and test period was increased. Best performance was for items that had just been fixated (better than 90%[detection]), followed by performance for items fixation 4 to 10 objects previously (around 83%[detection]), with the lowest level of performance for items tested at the end of the session (around 75% correct[detection]). These three levels of performance were statistically different which argues against the dichotomous STM/LTM model but is consistent with our finding that the information persisting beyond VSTM is not necessarily consolidated into permanent LTM (Melcher, 2006).

Melcher’s discussion ties directly to the internal process our brain assesses the type of objects being visually processed and then mapped to objects stored within memory. The more separated an object is from our STM, the more the mind must conduct a search for the next-best match based upon the visual objects in storage. If no object or next-best match exists, target detection will then likely not occur and the service member will be at risk of injury or mission failure.

Perceptual learning for visual target detection must ensure that there is positive training transfer. Training transfer studies continue to show that similar tasks within virtual environments are often transferable to operational environment settings. As stated by Anthony Ciaverelli:

The similarity between tasks (correspondence between stimulus environments and/or response patterns) and task-environment fidelity from a simulated activity to an operational environment are typically proposed to account for positive and negative transfer effects. But it has often been shown that positive learning transfer can occur between tasks that are quite different in composition or are performed in different environments (Ciavarelli et al., 2005).

While fidelity of training will be addressed within the discussion section of this thesis, two other issues related to of training transfer require short discussion: cognition and attention.

The service member's ability to understand "the structure or nature of the task" and learning "what to pay attention to and look for (salient cue recognition and discrimination)" are directly influenced by perceptual learning (Ciavarelli et al., September 2005).

During the cognitive phase of learning, the individual attempts to "intellectualize the requirements for learning the skill and seek to establish some knowledge about the nature of the tasks to be performed." In this manner, the service member observes the operational environment presented and attempts to "think" their way through the problem. As the individual progresses through the three stages of aptitude, they progress from intellectualizing the environment to a cognitive process that seemingly discriminates and generalizes automatically. A similar example of this type of cognitive visual process is that of scuba divers learning how to identify decorator crabs. As stated by an anonymous graduate school associate:

I'd been diving for years and never could find a decorator crab, so I went to a class. They had a presentation that told you what to look for and it had pictures and stuff like that. At first, I had to concentrate and look for the little signs and at the likely places that they live. Now, I can't dive without going, "Yep, there's one... there's one...there's one." They're everywhere (Anonymous, 2008)!

Very similar results have been reported within the aviation community concerning how pilots identify threat ground targets. “In many situations there are beneficial results (positive transfer) achieved with part-task training of key skill components” (Ciavarelli et al., 2005).

Perceptual learning within the virtual environment must therefore focus upon those aspects of visual target acquisition that facilitate the service members ability to better match the threat target presented with an object within the service member’s mental objects folder.

To facilitate this positive transfer of information and experience, corrective feedback is often used to ensure that the trainee (individual undergoing the training process) is learning effectively. Experience has shown that to be most effective, feedback should be provided to the trainee immediately after the skill is performed in a timely fashion and not while the skill itself is being conducted (Wickens et al., 2004).

4. Perceptual Learning in the Training Environment

This thesis continues research in the quest to answer “many interesting questions regarding perceptual learning (“What is learned?” “How long does learning take and last?” and “How widely does learning transfer?”)” (Goldstone, 1998). These questions, while approached from various academic and experimental angles over the past 40-plus years, continue to challenge us to answer how learning occurs and how to train individuals faster and with better training transfer.

A key aspect to this particular research is to determine the validity of perceptual learning through stimulus imprinting. As stated by Robert Goldstone:

The term imprinting captures the ideas that the form of the detector is shaped by the impinging stimulus. Internalized detectors develop...and increase the speed, accuracy, and general fluency with which the stimuli are processed (Goldstone, 1998).

Perceptual research indicates that individuals more readily adapt to an environment by directly imprinting to it. Imprinting allows detectors, sometimes referred to as receptors, to develop that are especially attuned to a stimuli or parts of stimuli.

Historically, perceptual learning and imprinting have allowed doctors to more readily diagnose disorders. When confronted by a new case similar to a previously presented case, even when the new patient is presenting symptoms irrelevant to the previous diagnosis, the doctor's perceptual skills facilitate the diagnosis (Goldstone, 1998).

Common sense and theory continue to find that people become better through education, training and maturation. The age old idiom of "If I show you, you will do it better and if you do it a lot, you will get faster" is often implicitly stated and agreed upon. The challenge is to determine just what it is within that expression that facilitates "better" and "faster." The human brain's power to perceptually identify ill-defined visual stimuli, when the subject has been afforded previous like-exposure, is both a quandary and extraordinary! The implication that a minimal amount of exposure training improves a person's ability to detect a target demonstrates the benefit of utilizing perceptual learning prior to inserting the inexperienced soldier into an operational environment.

As stated by Goldstone:

Although this effect is traditionally discussed in terms of implicit memory for exposed items, it also provides a robust example of perceptual learning. The identification advantage for familiarized instances...requires as few as one previous presentation of an item, and is often tied to the specific physical properties of the initial exposure of the item. In brief, instance memories that are strong and quickly developed facilitate subsequent perceptual tasks involving highly similar items (Goldstone, 1998).

The issue at hand, regardless of one's stance on perceptual learning, training, and human expertise, is "how does one capture what one sees and then put it to use in future situations?" Specifically, as trainers of soldiers, we must ask ourselves: Is there more to getting better at visual target acquisition than just practice, luck and survival?

B. IMPLICATION AND HYPOTHESIS

This attempt at investigating perceptual learning must extend previous studies to the particular population of interest. Based on the literature discussed above, the following hypothesis is proposed: A subject's ability to acquire a target, as measured by a Hit Rate, is affected by the difficulty of the scene and by the visual feedback the subject receives.

THIS PAGE INTENTIONALLY LEFT BLANK

III. METHODOLOGY

A. PARTICIPANTS AND DESIGN

Participants of this research will be students at NPS comprising a cross section of service and varied level of operational experience and military specialty. The participants were screened through a voluntary completion of a questionnaire (Appendix A) that assessed their experience level with regards age, years of service, and experience utilizing video gaming technology. No volunteers were screened out as the demographic information was gathered for predictive and qualitative analysis.

1. Hypotheses

Null hypothesis one, H_{o1} , is that the Hit Rate of the Control, Feedback, and Training-Feedback groups is the same.

Alternative hypothesis one, H_{a1} , is that the Hit Rate of the groups is not equal and differences exist.

Null hypothesis two, H_{o2} , is that the Hit Rate of difficulty levels (Easy Scene vs. Hard Scene) is the same.

Alternative hypothesis two, H_{a2} , is that the Hit Rate between the difficulty levels is not equal and differences exist.

2. Design

A 2 by 3 between subjects design was used to explore the target acquisition skill of participants as a function of the scene's difficulty and the feedback provided. There were two levels of scene difficulty, Easy and Hard. There were three group assignments possible to which each subject was randomly assigned: Control, Feedback or Training-Feedback. To minimize possible confounds, a pilot test was conducted to rehearse the laboratory procedures and ensure researcher-subject interaction as practiced and standardized.

The Control Group was required to point-and-click on targets detected in each scene and then presented no feedback.

The Feedback Group was required to point-and-click on targets detected in each scene and then provided a side-by-side visual feedback feature.

The Training-Feedback Group was provided a five-minute instruction class on the U.S. Army's ground search technique. Following instruction, the Training-Feedback Group was required to point-and-click on targets detected in each scene and then provided a side-by-side visual feedback feature.

The feedback feature utilized the Fujitsu laptop to display the targets in a side-by-side comparison of each test scene. This allowed the Feedback and Training-Feedback Groups to identify missed targets, through comparison of vivid red (False Colored) targets to the same Normal Colored targets, within each presented scene.

3. Measures

Dependant measures. The game engine application, using data collected from each scene sequence, captured the subject's interaction via the computer mouse point-and-click target selection. The Hit Rate was then calculated by assessing the number of targets correctly selected by the user, given the number of targets presented. In this manner, the Hit Rate was captured for all 24 scenes for each subject. Time of first point-and-click target selection was also collected, but variation between subjects was not statistically significant and is not included within this report.

4. Control Procedures

The control procedure utilized to eliminate extraneous variables included scripted directions to all participants, randomized assignment to experiment group, and counter balanced scene presentation order. These procedures ensured a randomized block design was used to assign equal numbers of participants to each of the six conditions.

B. EQUIPMENT

Equipment utilized for this experiment included a Dell XPS XPS720 computer with an Intel Core2 Quad CPU operating at 2.40 GHz and 3.0 GB RAM. The graphics were displayed utilizing a NVIDIA GeForce 8800GT performance video card.

The scenes were displayed utilizing a Dell flat screen computer monitor, measuring 24 inches diagonally, and set to 1600 x 1200 screen resolution and 60 Hertz refresh rate. The monitor was adjusted ensuring the top of the monitor is roughly at eye level with the center of the screen being at approximately 20 degrees declination (top leaned away from the subject) and one arms length from the subject. The screen was tilted back from the subject at approximately a 10-degree angle.

A Fujitsu LifeBook T Series was utilized to provide the feedback scenes to the experimental/treatment groups. The Fujitsu uses a Intel Core2 Duo CPU and operates at 2.4 GHz and a Mobile Intel 965 Express Chipset display adapter. The Fujitsu flat screen display measures 12 inches diagonally and was set to 1024 x 768 screen resolution and 60 Hertz refresh rate. The feedback scenes were copied as bitmap format to Microsoft Office PowerPoint 2003 and presented as a Slide Show which matched the subject's assigned scene sequence. Normal colored and false colored scenes can be found in Appendix B.

The test area was a windowless eight foot by twelve-foot room at the Naval Postgraduate School's Human System Integration Laboratory. The room was darkened to provide the optimal viewing experience for each subject and external light sources were minimized to reduce glare on the display screen.

C. VISUAL SCENES

As the present research will initiate a proof-of-concept, the training will focus on the target objects and utilize the virtual environment of the Delta3D game engine. The subject task was to detect a static target within a virtual environment and Hit Rate data was analyzed to determine the extent that perceptual learning influences target detection.

The scenes utilized for this experiment were the result of ongoing collaboration and research efforts. As stated by the Ryan Wainwright:

The stimulus package used in this target detection research was the result of a collaborative effort. Mr. Michael Dunhour developed the basic structure of the scene-generating software with the use of Delta-3D, an open-source computer game development engine. Mr. Daniel McCue developed the software that enabled the placement of targets, the adjustment of scene conditions, and the recording of mouse-click data. Since previous studies tend to highlight the importance of having a large number of target detection opportunities, the experimenter used these products to develop 24 unique scenes (Wainwright, 2008).

In addition, I created nine additional scenes, of which seven were used during subject orientation and practice sessions. Four of these scenes were utilized during the “training” phase of each subject’s orientation to the software and experimental design and three scenes were utilized during the Training-Feedback group’s pre-test training session.

As with previous experiments, this research utilized multiple target presentation opportunities within each scene. All scenes are designed with not less than one target per scene to eliminate subject bias towards associating the inability of visually identifying a target with the possibility that no target exists. As all scenes within this experiment had at least one target, the subjects all understood that there was never less than one enemy combatant represented in each scene.

These factors include the amount of target within view (how much the target was masked by the surrounding environment – often described as “cover and concealment”), target brightness, and target contrast. Additionally, the number of “easy” scenes and “hard” scenes was balanced with the number of targets displayed in each scene. Scene difficulty was assessed by myself as either Easy or Hard.

Easy difficulty required that the target outline be noticeable to most subjects as defined by non-matching contrast with surroundings, over 33% target exposed, and poor edge matching (such as a silhouette in a window).

Hard difficulty required that the target outline be near-matching in contrast (good camouflage), less than 33% of target exposed (using cover and concealment) and good edge matching (sniper behind a rubble wall, etc.).

Therefore, the scene package developed previously, and utilized for this research, is “24 scenes consisting of one target in each of two easy and two hard scenes, two targets in each of two easy and two hard scenes, and so on up to six targets in each of two easy and two hard scenes. Twelve easy scenes contained 42 targets, as did twelve hard scenes and by the conclusion of the [scene] package each [subject] had been presented 84 targets” (Wainwright, 2008). While this definition of “Easy” and “Hard” is somewhat subjective, as applied by the scene’s various artists and the researcher, it is an honest appraisal of each scene’s expected difficulty with regards to the subjects ability to visually identify the targets.

Table 1 displays the experiments target scene number, target(s) presented, and assigned difficulty.

Table 1. Targets by Scene and Difficulty

Scene Number	Targets	Difficulty
1	4	Easy
13	1	Easy
23	2	Easy
25	6	Easy
38	6	Easy
40	3	Easy
41	1	Easy
42	5	Hard
44	3	Easy
45	5	Easy
46	1	Hard
47	6	Hard
48	4	Hard
49	2	Hard
50	5	Hard
51	3	Hard
52	1	Hard
53	2	Hard
55	4	Hard
56	5	Easy
57	4	Easy
58	6	Hard
59	3	Hard
61	2	Easy
1009	4	Easy
Total Targets	88	
Average Targets / scene	3.5	
Variance per scene	2.9	
Std deviation	1.7	
Most number trgts	6	
Least number trgts	1	

D. EXPERIMENT DESIGN AND PROCEDURES

All subjects were systematically assigned to one of the three experiment groups previously described. The first participant was placed in the Control group, the second subject in the Feedback group, the third subject in the Training-Feedback group and then the assignment process repeated itself until the experiment concluded. As the assignment

was not made until the participants “walked through the door,” the systematic sequential assignment resulted in an unbiased, pseudo-random assignment to groups.

1. Participant Screening

Prior to experimentation, the researcher tested each individual’s visual acuity and color vision. The visual acuity of each individual was at least 20/40 and ensured that the test subject is capable of focusing upon the virtual scene depicted. A Snellen Chart was placed upon a wall and the subject stood 10 feet away and was asked to read from line 5 while covering first the right eye and then the left. Once the subject’s vision was determined to be at least 20/40, for each eye, the subject proceeded to the color vision test. While this test for acuity is not a true diagnosis of the subject’s precise visual acuity, the test provided an acceptable level of confidence for the researcher that the subject would have the visual capacity to identify the features of the presented targets. Figure 11 is an example of a common Snellen Chart similar to the one utilized during the subject’s visual acuity assessment (Wainwright, 2008).

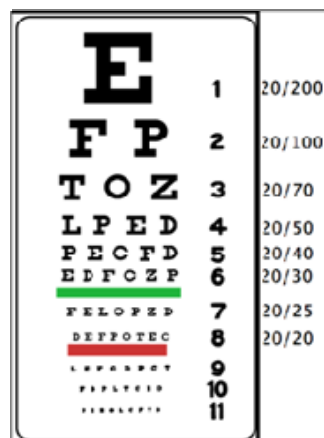


Figure 11. Example Snellen Chart (From Wikipedia, 2008)

Regarding target acquisition, previous research described both the advantage and disadvantage of color vision and color vision’s relationship to visual search tasks. Prior to experimentation, the subject’s color vision was tested as it could influence the subject’s visual identification the target object in the complex environment. Each participant, therefore, was administered a basic color vision test to determine any level of

color deficiency. Although no subject was denied experimental participation, a standard Ishihara test was administered as an expedient method to classify the color vision of the subjects utilizing a Ishihara's Tests For Colour Deficiency 24 Plates Edition (2006). For purposes of this experiment, the subjects were asked to identify the number visible in the first five of the six presented plates (see Figure 12). This information was utilized to determine the subjects color vision as the final visual evaluation step prior to testing.

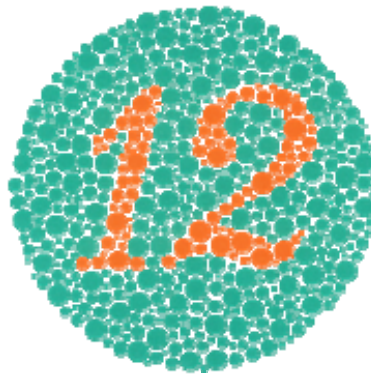


Figure 12. Example Ishihara Plate (From Wikipedia, 2008)

Once these two tests are complete the subject shall be administered a questionnaire that will assess their age, gender, military experience, and an individual basic health assessment. This information, while not directly studied within this particular body of research, will be used to determine test subject demographics. A post-test questionnaire will be utilized to qualitatively assess each subject's opinion as to the training received during the experimentation. The questionnaires may be viewed in Appendix A.

2. Experiment Procedures

The experiment initiates with the participant orientation to the HSIL laboratory and the subject's informed consent to participate in the experiment. The subject then completed the demographic questionnaire and visual acuity and color vision exam.

Once assigned to one of the three test groups, the subject was provided a short situational brief that described a tactical situation in which they are a soldier conducting a foot patrol in an urban contemporary operating environment (COE). All subjects were briefed that only enemy combatants were present in all scenes and that current rules of engagement (ROE) allowed firing immediately upon threat target identification. The intent of this scenario briefing is two-fold: to place the subject in the proper mindset similar to that of a soldier on patrol, and to focus the subject on the visual search and target acquisition task.

Instructions regarding subject interaction with the stimulus package were specific to each group and no group was privy to the other group's instructions. Once assigned to one of the three experiment groups, each subject was exposed to all 24 test scenes.

Individuals assigned to the Training-Feedback group were given a 5 minute class on the basic fundamentals of unaided human visual search and detection of static targets in the virtual environment. The training utilized the U.S. Army's ground search technique as described in the laboratory procedures section of this paper. The instruction was given on three similar scenes as those utilized for the experiment and was provided by the researcher, a SME, with 15 years of infantry experience.

The Training-Feedback (experimental) group, Feedback (second experimental) and Control (no-feedback) group received instruction that directed their available number of shots was not restricted and that, while there is never more than six targets per scene, they could fire as many times as they felt there were targets.

Each subject had a practice session with the three scenes, comprised of four targets within each scene, to ensure they understand the interface and how interaction with the system occurs. The three scenes were the same for all participants and not counterbalanced in any way to protect the subjects from confounds.

Subjects were reminded that the instructions and procedures remained the same and no changes were required between the practice set and the experimental set.

Subjects completed the visual target acquisition task for all 24 scenes without interruption and were allowed 15 seconds to scan each normal colored target scene

(Figure 13). The Researcher then presented the same 24 target detection scenes to all participants with an equal number of easy and hard scenes and the scenes were randomly ordered for each participant.

Subjects assigned to the control group completed the visual target acquisition task for all 24 scenes without interruption. Subjects were allowed 15 seconds to scan each scene (Figure 13), proceeding from one scene to the next, until the randomized set was complete. Total testing time for the Control group was approximately 30 minutes.



Figure 13. Normal Colored Target Scene # 1009

Subjects assigned to the feedback and training-feedback groups were provided side-by-side comparison of enemy target information for the scene they just viewed by being provided the “false colored” screen shot of the tested scene they just viewed. This false coloring of the target highlights that portion of the target visible to the viewer in red (Figure 14).



Figure 14. False Colored Feedback Scene #1009

The Feedback and Training-Feedback subjects were allowed to view and compare the scenes for up to 15 seconds and then asked to proceed to the next scene until all 24 scenes had been viewed and the visual target acquisition task was complete. The subjects were, therefore, able to compare the false colored feedback scene to the normal colored (test) scene in a side-by-side comparison to facilitate identification of missed target location and properties (head exposed behind wall, etc.). Total testing time was approximately 45 minutes for the Feedback group and 50 minutes for the Training-Feedback group.

Following the completion of the experiment, each subject was asked to answer a questionnaire used to qualitatively assess their perceived training as it related to their ability both pre and post experiment.

At the completion of the experiment, all participants stated that they understood the importance of not compromising future data collection efforts by discussing the experimental session with others. This was very important, as this is a body of research that has opportunity for expansion and reutilization.

A complete copy of the laboratory procedure, instructions to participants and the pre and post-test questionnaires is attached in Appendix A.

THIS PAGE INTENTIONALLY LEFT BLANK

IV. RESULTS

A. DEMOGRAPHICS

This study utilized 23 male and 4 female participants with an average age of 33.2 years (4.7) and 11 years (5.4) military service. Demographic information for each group is depicted below in Table 2.

Table 2. Demographic Information by Group

	Group		
	Control	Feedback	Training-Feedback
Male	8	7	8
Female	1	2	1
Average Years of Service	10	11	12
Average Age (Years)	32	35	33

B. ANALYSIS OF HIT RATES

The data recorded from each participant for each scene presented included the number of hit targets, number of targets presented, number of shots (mouse clicks), and the time each target was shot. A total of 2,272 targets were presented, 2,654 shots taken, and 1,863 hits were observed. This data was entered into the JMP software statistics package for analysis. Overall results for Hit Rate by Group are depicted in Table 3.

Table 3. Hit Rate by Difficulty and Group (Mean and Standard Deviation)

		Group			All Subjects
		Control Group	Feedback Group	Training-Feedback Group	Overall Hit Rate
Difficulty	Hard	0.57 (0.36)	0.76 (0.29)	0.70 (0.35)	0.67 (0.01)
	Easy	0.93 (0.16)	0.97 (0.11)	0.95 (0.11)	0.94 (0.13)

The Hit Rate of the 27 participants for all scenes was 0.819 (standard deviation 0.29). The overall shots per hit target was 1.06 for easy scenes, 1.43 for hard scenes and 1.22 overall.

The analysis allowed a between subjects comparison of experimental groups and a within-subjects comparison of scene difficulty with regard to Hit Rate.

1. Results for Hit Rate by Difficulty and Group

The data were analyzed using a 2 (Difficulty: Easy or Hard) by 3 (Group: Control, Feedback, or Training-Feedback) ANOVA, where Difficulty was a within-subjects variable and Group was a between-subjects variable. Description of findings and representative figures are below.

The ANOVA for Group was significant, $F(2, 646) = 8.16$, Group $p = .0003$.

The ANOVA for Difficulty was significant, $F(1, 647) = 179.04$, Difficulty $p = <.0001$

The ANOVA for the interaction between Group and Difficulty was also significant, $F(5,643) = 43.37$, and $\text{Difficulty} \times \text{Group } p = 0.0001$.

The ANOVA results suggest that both the Difficulty and Group assignment and their interaction influenced target detection performance.

Further analysis determined that the Hit Rate for the Control group decreased more rapidly as the scene changed between Easy and Hard difficulty. The Hit Rate for the Feedback and Training-Feedback group also decreased with Hard scenes, but at a lesser rate (See Figure 15).

Group assignment indicates that all groups had similar hit rates for Easy scenes, but for the Hard scenes, there is a noticeable difference between the treatment groups and the Control group with the Feedback group demonstrating the greatest Hit Rate advantage over the Control group. Graphic representation of this interaction is evident in Figure 15.

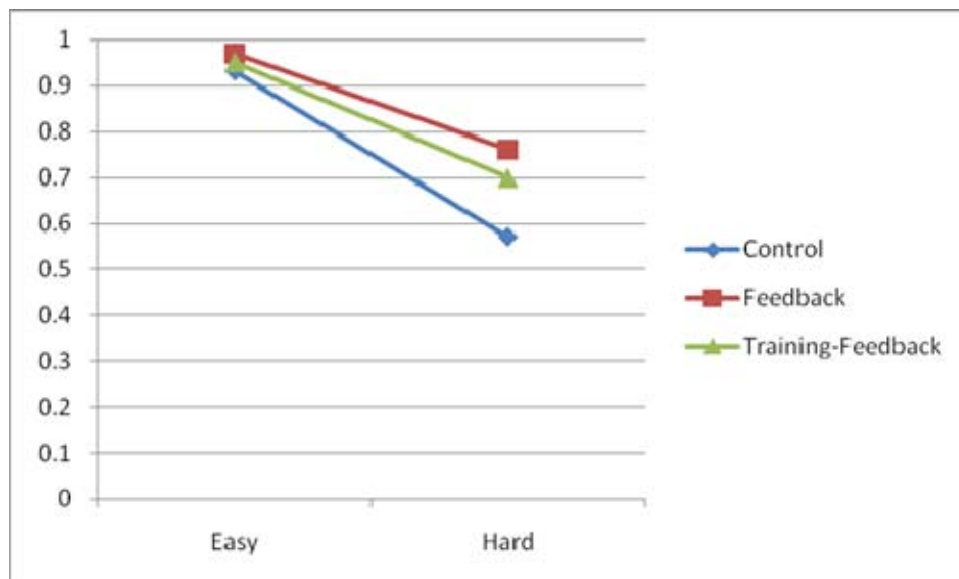


Figure 15. Graph: Hit Rate by Difficulty and Group

For the Hit Rate, Group (G; $t(9)=8.16$, $p=0.0003$) and Difficulty (D; $t(27)=179.04$, $p<0.0001$) explained a significant proportion of the variance in mean Hit Rate (HR) performance (HR; $R^2=0.252$, $F(5,643)=43.38$, $p<0.0001$):

$$HR = 67.26 - 0.06G - 13.66D - 0.77G*D + \epsilon$$

Indication of significant differences between the control group and the treatment group were found through analysis and comparison of means utilizing the Tukey-Kramer Honestly Significant Difference test. This Test uses the distribution of the maximum range among all data collected. The comparison of means indicates that there was a statistical significance between the Feedback and Training-Feedback group when compared to the Control group. There is no significant difference, however, between the Feedback and Training-Feedback Group (See Table 4).

Table 4. Tukey-Kramer Comparison by Group: Hard

Group	Tukey-Kramer Letter Assignment	Mean
Feedback	A	0.76
Training-Feedback	A	0.70
Control	B	0.57
Groups not connected by same letter are significantly different.		

The Hit Rate analysis for Hard Difficulty by Subject the overall Minimum Mean Hit Rate = 0.36 (0.097) and overall Maximum Hit Rate = 0.90 (0.097) with the remaining 25 subjects normally distributed. Analysis indicates that for the 12 Hard Difficulty scenes, scene difficulty had measurable effect upon Hit Rate (Figure 16).

Each of the three experimental groups had 9 subjects assigned and all subjects observed 12 Hard difficulty scenes. This ensured a total of 108 Hard scene presentations per group. Control group scenes $N = 108$, producing a Mean Hit Rate = 0.57 (0.36). The Feedback group scenes $N = 108$, producing a Mean Hit Rate = 0.76 (0.29). The Training-Feedback group scenes $N = 108$, producing Mean Hit Rate = 0.70 (0.35). Analysis indicates that for the 12 Hard difficulty scenes, Group assignment had significant effect on Hit Rate (Figure 16). Given the Hard scene difficulty, ANOVA for Group was significant, $F(2, 322) = 8.35$, Group $p = .0003$.

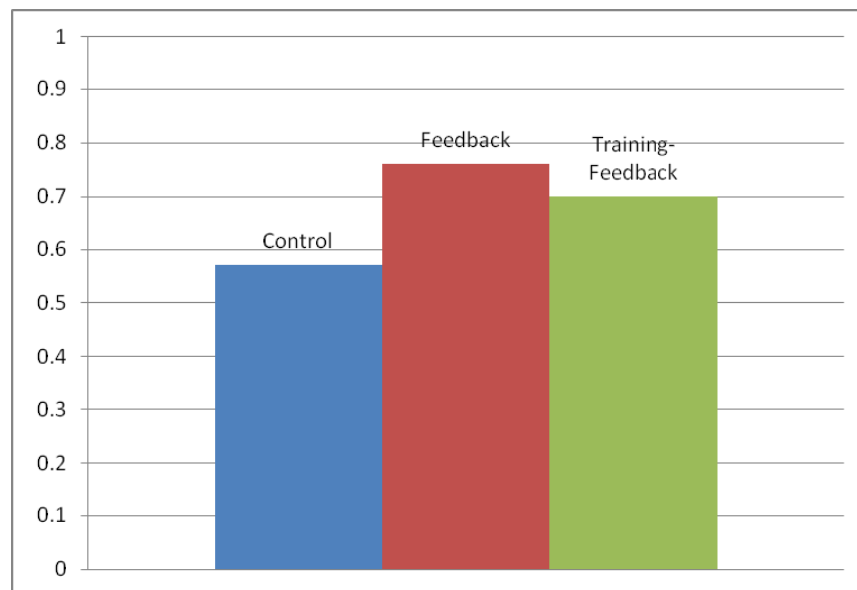


Figure 16. ANOVA Hit Rate by Group: Hard

Each of the three experimental groups had 9 subjects assigned and each subject was presented 12 Easy difficulty scenes. This ensured a total of 108 Easy scene presentations per group. Control group scenes $N = 108$, producing a Mean Hit Rate = 0.93 (0.16). The Feedback group scenes $N = 108$, producing a Mean Hit Rate = 0.97 (0.12). The Training-Feedback group scenes $N = 108$, producing a Mean Hit Rate = 0.95 (0.11). Analysis indicates that for Easy scenes Group assignment had no significant

effect upon Hit Rate. The ANOVA for Group was not significant, $F(2, 322) = 2.67$, Group $p = .07$ (Figure 17). This result is consistent with the interpretation that a “ceiling effect” was obtained with the Easy scenes.

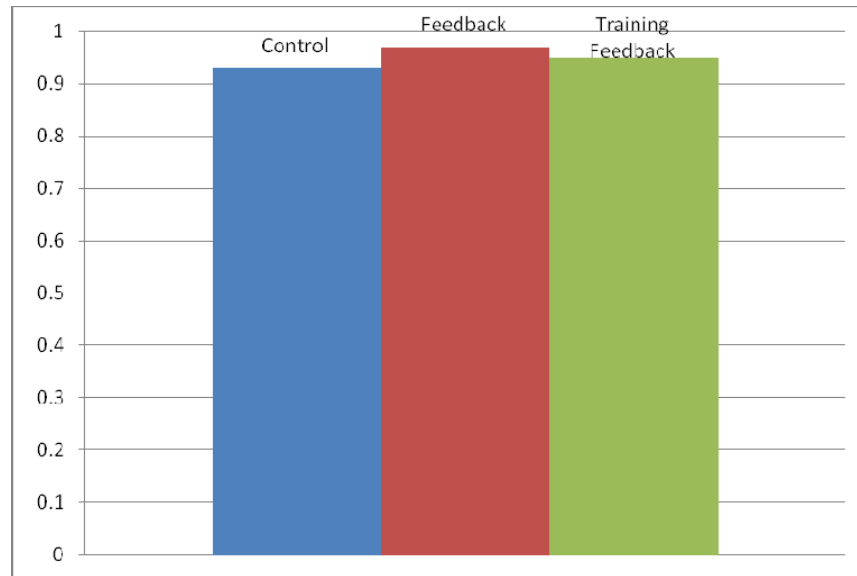


Figure 17. ANOVA Hit Rate by Group: Easy

2. Questionnaire User Feedback

All subjects were presented a post experiment questionnaire to obtain attitudes and opinions about virtual environment training. The responses were gathered utilizing a five point gradient scale ranging from Strongly Disagree to Strongly Agree.

When questioned if they believed that the training would “improve their target acquisition skills” in an operational environment, 100% of the Training-Feedback and 89% of the Control and Feedback groups responded within the Agree/Strongly Agree category.

The degree to which the subject perceived targets as “easy to acquire in the scenes” displayed greater variance. 44% of the Feedback and 67% of the Control group Agreed/Strongly Agreed while 78% of the Training-Feedback group replied in same. This finding is interesting in that the subjects in the group with the lowest Hit Rate did

not perceive the targets as difficult to acquire. Subjects often stated, “Lighted scenes are much easier to find targets and the mid range targets are easier to spot than close or far range targets.”

Virtually all subjects responded that they “improved as they were exposed to more scenes,” which indicates that the comfort level of the subject increases directly in proportion to exposure over time. A common theme found within verbal response of subjects was “I became familiar with scenes and was able to not look at urban features and instead focus on targets.”

With regards to the Feedback and Training-Feedback groups and the perception that the feedback feature made “target acquisition faster,” 89% of the former and 100% of the latter Agreed/Strongly Agreed. This lends credence to perceptual learning theory as presented in the Discussion section. As demonstrated in the above response, one subject’s unsolicited comment regarding the feedback feature may allude to perceptual learning even in the short term. As stated, “The side-by-side feedback feature was a great advantage as it taught me what to specifically look for.”

Although the Training-Feedback group’s Hit Rate was not as high as the Feedback group, 89% Agreed/Strongly Agreed that the training “improved visual target acquisition.” One U.S. Army officer’s comment lends itself to the debate as to game based training systems and utilization, “The apparatus used was excellent...I could see how this application could be used to train soldiers how to scan and what features to look for in an urban combat environment.”

The various subjects’ unsolicited verbal feedback, both constructive and critical, was useful and appreciated by the research and is addressed within the body of the Discussion section.

C. POST-HOC DISCUSSION OF HYPOTHESES

The results confirm that differences in Hit Rate exist with regards to the Group assignment and the Difficulty of the scene.

Null hypothesis one, H_{01} , that the Hit Rate between the Control, Feedback, and Training-Feedback groups is the same is rejected as there is a statistically significant effect of Group.

Null hypothesis two, H_{02} , that the Hit Rate between difficulty levels (Easy Scene vs. Hard Scene) is the same is rejected, as there is statistical significance between scene difficulty and Hit Rate.

When presented Hard Scene difficulty, the Feedback treatment group demonstrated superior target detection skill in comparison to both the Control and the Training-Feedback group. The Training-Feedback group demonstrated superior target detection performance in comparison to the Control group, but was slightly less effective than the Feedback group.

When presented scenes of increased difficulty, subjects had a much lower Hit Rate for Hard Scenes than for Easy Scenes. Targets in the Hard Scenes, proved much more difficult for subjects to acquire when compared to the Easy scenes with the same number of targets.

V. DISCUSSION

A. PERCEPTUAL LEARNING IN VIRTUAL ENVIRONMENTS

This study indicates that utilization of popular game engine technology does promote perceptual learning within virtual environments. There is indication that, given further testing and evaluation, this type of application may improve a soldier's operational perception skills. This perceptual skill improvement within operational environments would undoubtedly lead to an increase in soldier survivability, improved morale, and unit capability.

One of the challenges with determining the relationship between Hit Rate and perceptual learning is the lack of understanding of how individuals internalize and process the presented environmental stimuli. As stated by Robert Goldstone:

[A] more general problem with treating object recognition as a separate process from learning is that no account is given for how object descriptions are initially internalized. Even under the assumption that we recognize objects by decomposing them into elements, we still need processes that learn object descriptions. One might assume that the first time an object is viewed, a description is formed and a trace is laid down for it. After this initial registration, standard object recognition routines are applied. This approach would preserve the separation of object recognition and learning, but given the strong influence of object familiarity on recognition, it is too gross a simplification. Object learning occurs simultaneously to, and interacts with, object recognition (Goldstone, Schyns, & Medin, 1997).

As discussed within this thesis, in order to encourage positive transfer of information and experience, feedback was used to ensure that the trainee (individual undergoing the training process) is learning effectively. Experience has shown that to be most effective, feedback should be provided to the trainee immediately after the skill is performed in a timely fashion and not while the skill itself is being conducted (Ciavarelli et al., September 2005) (Wickens et al., 2004) .

The flickering red of the hit target likely provided learning and feedback to each group, to include the Control group, and was addressed by all subjects during the post-test questionnaire. This flickering effect created by the False Positives application facilitated both mental models within the STM and allowed all groups to develop strategy of where to search for targets (likely locations).

Another interesting phenomenon that occurred was the number of shots for Feedback and Training-Feedback were greater as they were made aware that they were missing targets. This likely led to higher incidence of false positives (TYPE 1 error) as the subjects engaged those objects in the scene that resembled previous targets given the selected object's contrast, shape and location.

1. Perceptual Learning Ramifications

Perceptual learning involves the long-term changes utilized by an organism's perceptual system that improve its ability to interact with its environment. The extent to which a short duration training event influences perceptual learning is challenging and requires further experimentation and evaluation.

Given the type of training conducted within the bounds of this experiment there are several perceptual learning theories that may explain the results.

All subjects may have utilized attentional weighting. Attentional weighting allows the subject to increase the attention paid to those perceptual dimensions and features that they deem important to the task at hand. This experiment likely leveraged this process as the subjects were focused in the area of category learning and internalized the stimulus aspects that tied directly to their task success: the human in the scene. The focus of the subject became the learning and categorizing of the shape, contrast, and texture of those portions of the targets that were evident within each scene. The curve of the human form, contrast of textures given the light source (light objects and shapes existing within shadow), or geometric pattern interruption (target in a doorway). These categorical patterns allowed the subject to internalize both the shape of the exposed target and experience the targets interaction with the environment (Goldstone, 1998). An

advantage given to the treatment groups was the side-by-side comparison of the rendered screen and a false colored screen. This treatment likely encouraged internalized stimulus imprinting by the subject. Stimulus imprinting adapts one's perception to the environment by directly imprinting to it. Detectors are developed, over time, that facilitate the ability to perceive both the whole stimuli and partial stimuli much in the fashion that a glance at an object allows a discernment of the whole even with only partial information. Internalized detectors are "shaped" from repeated exposure and subjects likely performed better as they became attuned to the targets presented. Historically, this effect has proved a robust example of perceptual learning and some studies have shown that "instance memories are strong and quickly developed facilitate subsequent perceptual tasks involving highly similar items" (Goldstone, 1998).

Feature imprinting, a short term memory (STM) skill utilized to quickly recognize stimuli, may have proved to be one of the strongest perceptual learning tools used by the subjects in all groups.

Rather than imprinting on entire stimuli, there is also evidence that people imprint on parts or features of a stimulus. If a stimulus part is important, varies independently of other parts, or occurs frequently, people may develop a specialized detector for that part. This is a valuable process because it leads to the development of new building blocks for describing stimuli (Schyns et al 1998, Schyns & Murphy 1994). Parts that are developed in one context can be used to efficiently describe subsequent objects. Efficient representations are promoted because the parts have been extracted because of their prevalence in an environment, and thus are tailored to the environment (Goldstone, 1998).

This imprinting is most evident when comparing the Hit Rate of the Feedback group to the Control group. The ability to determine not only where the missed targets were within the scene but also to view how the targets were "hidden" within the scene encouraged perceptual learning. A understanding of where targets could position themselves in the scene, how much of the target was exposed and what that portion of exposure looked like, and how the targets physical interaction with the environment likely positively influenced both experimental groups Hit Rate.

2. Perceptual Learning influence on Feedback and Training-Feedback Groups

An interesting phenomenon that occurred within the Training-Feedback group was that the training had negative influence upon the Training-Feedback's Hit Rate. A possible explanation of this negative relationship is that the training affected the cognitive capacity of the Training-Feedback group. Possibly, instead of interacting with the scene and locating targets based upon previous experience building upon newly gained knowledge, the Training-Feedback subjects were challenged with learning new visual scanning procedures. By integrating this short training event into the experimental design, the researcher expected to see the Training-Feedback group have the highest Hit Rate. The Training-Feedback group trailed the Feedback group by approximately 6% (70% to 76%, respectively) even though they were presented the same side-by-side comparison.

As discussed by Robert Goldstone:

In many cases, perceptual learning involves acquiring new procedures for actively probing one's environment (Gibson 1969), such as learning procedures for efficiently scanning the edges of an object (Hochberg 1997, Salapatek & Kessen 1973). Perhaps the only reason to selectively highlight perceptual learning is to stress that flexible and consistent responses often involve adjusting initial representations of stimuli. Perceptual learning exerts a profound influence on behavior precisely because it occurs early during information processing and thus shifts the foundation for all subsequent processes (Goldstone, 1998).

Time to engage the first target was nearly identical for all three groups. This may stem from the environment, as there was similarity between scenes concerning rendered and perceived angle, distance, and likely target locations. In essence, all subjects came to expect the target in a certain location within the rendered screen.

A possible confound within this study was determined from a manual study of the computer collected data set. This evaluation determined that each control group subject was allowed approx 10 seconds time to view the screen while manual data collection was taking place. In essence, this may have impacted the subject's perception of the targets within the scene and provided some minimal training feedback to the control group.

Ultimately, differences did exist between the control and treatment groups and indicate that utilizing different technologies may positively influence soldier target-detection skills. The use of game engine technology for improving basic soldier skills may prove advantageous over time as the application may be easily nested within a larger and already popular gaming application. As game engine technology continues to improve and fidelity and resolution become more “near real” there is an opportunity to familiarize soldiers with situations and environments not yet experienced.

B. AREAS OF FUTURE RESEARCH

The purpose of this experiment was to explore the utilization of perceptual learning techniques for target acquisition. The goal of the experiment was to determine whether game engine technology could support perceptual learning of visual target acquisition skills. While the results of this experiment suggest that perceptual learning does take place, there is still much to be learned regarding training transfer and individual differences.

Assessing training transfer between the virtual environment and the real world operational environment is still required. Operational testing would assist the U.S. Army Simulations Operations community, and the Department of Defense, with further understanding how commercial off-the-shelf (COTS) applications may be utilized not just as scenario trainers and collective training apparatus, but how the “games” may also assist with improving the individual training level within Live, Virtual and Constructive simulations.

The length of time that perceptual learning and training transfer remains with the individual must be evaluated in order to understand the length of time the training remains valid and when refresher training would be beneficial. A longitudinal study that tests all of the subjects’ ability to detect targets, given the same scenes, would determine if perceptual learning occurred.

As graphics and game engine applications continue to improve, those conducting future research may wish to utilize an experimental design that also allows navigation

through the virtual scene. This would allow the subject to move through the virtual environment and gain perceptual learning in both target acquisition and environment interaction (cover and concealment, sniper lanes, etc.).

APPENDIX A

A. LABORATORY PROCEDURES

Before Subjects Arrival:

1. Ensure recording of previous subjects data

Excel Spreadsheet

Data Results File renamed as subject's number

Data Output File ran with Python app and renamed to subject number

2. Assign next participant number and no-feedback (odd #) or feedback

Prepare laptop for feedback if required

3. Prepare Forms

Consent: Subject number at top

Questionnaire: Subject number at top

Annotate Email Address on email address sheet

4. Prepare Equipment/Experiment

Vision Test:

Ishihara color machine

Stopwatch

Apply training set of scene sequences to the output file in the FalsePos application

Ensure that the random set sequence for that subject number is readily available (file name and subject name shall match)

5. Have training application prepared for subject 3, 6, 9, 12

Subject's Arrival:

Record Arrival time and participant number on all forms (verify)

Orient Subject to the Laboratory and read/sign the consent form

Confirm responses on the screening questionnaire.

1. Administer vision tests

2. Color: "We are using the Ishihara's Tests For Color Deficiency 24 Plates Edition to assess your color vision. Note that you need only read the number incorporated into the the plates that have been tabulated for you and that there are six plates total. Please either state the number that you can read or "no number visible." Proceed."

3. 20/40 Acuity: "Step behind the tape on the floor with your toes touching the edge closest to you. Cover your left eye and read this line (point to the 20/40

line), from left to right. Now, cover your right eye and read this line (point to 20/40) you from right to left. Now, with both eyes uncovered read this line (point to 20/20), from left to right.”

Read tactical scenario to participant (instructions to participants)

1. Participant 1, 4, 7, 10: No feedback scenario
2. Participant 2, 5, 8, 11: Training- Feedback scenario
3. Participant 3, 6, 9, 12: Feedback scenario

Briefback Questions asked after reading the scenario to each participant (with answers in parenthesis)

1. How many targets are present in each scene? (From 1 to 6)
2. How long can you search any given scene? (15 seconds)
3. How can you tell when you have properly “engaged” a target? (It turns red/shaded)
4. Are you ever going to be presented a scene with no targets? (No, there will always be at least 1 and not more than 6 targets in each scene. Between 1 and 6 targets, always).
5. Should you indiscriminantly fire (mouse click) all over the scene in the hope of finding a target? (No. And if you do you may “lock up” the application!)

Target Detection Sequence:

1. Verify the subject’s random scene sequence is in the output file in FPA
2. “I will demonstrate how the FPA application works on the first scene. Upon completion of the demonstration you will be presented 3 scenes with enemy combatants in the open. Each scene will have 4 combatants and you will have 15 seconds to identify the combatant, place the cursor over the targets with the mouse, and left click to engage the targets. You must “point and click” each target for the target to be highlighted in Red, in other words “shot.” I will notify you when 15 seconds is up and ask you to press the enter key on the keyboard. The enter key, once depressed will cycle the application to the next scene at which time the 15 seconds will immediately begin and you will repeat the target ID and point-and-click cycle. Do you have any questions? Begin.”

3. Click the .exe file for the subject to begin the FPA sequence and demonstrate on scene 1.
4. Upon completion of the three scenes:
 - a. Ask if there are any problems or questions
 - b. Ask if the participant understands that the 15 second time requirement will remain the same.
5. Ask subject: “Are you prepared to conduct the experiment at this time?”
 - a. No-Feedback: “This experiment will take approximately 10 minutes to conduct from the time you begin the sequence. Are you ready? Begin.”
 - b. Feedback: “This experiment will take approximately 20 minutes to conduct from the time you begin. Are you ready? Begin.”
6. No Feedback:
 - a. Observer: Time each scene every 15 seconds.
 - b. Observer: Manually enter the number of targets the subject Hit in Column C
 - c. At end of 15 seconds: “Stop. (Ensure Hits are annotated) Please depress the enter key and proceed to the next scene”
 - d. Repeat steps A through C until complete.
 - e. After all scenes: “Thank You. Please let me verify that the output data has correctly updated.” Verify results text file is in place
 - f. Record completion time
 - g. Remind participant to not discuss the project or any of the project details until the data collection is complete and we notify you by email.
 - h. Verify email account is annotated on the IRB release form.
 - i. Thank You!

7. Feedback:

- a. Observer: Time each scene every 15 seconds.
- b. Observer: Manually enter the number of targets the subject Hit in Column C
- c. Observer: Verify the feedback scene is the same as the scene the subject is engaging.
- d. At end of 15 seconds: “Stop. (Ensure Hits are annotated) Please turn your attention to the feedback scene and note the location, and any visual information as it appears to you, of the red targets. The red targets are the same enemy targets presented to you in the experiment scene. You may study the feedback scene for 15 seconds.”
- e. After 15 seconds: “Stop. Please turn your attention back to the FPA monitor and depress the enter key to proceed to the next scene”
- f. Repeat Steps A through E until complete.
- g. After all scenes: “Thank You. Please let me verify that the output data has correctly updated.” Verify results text file is in place
- h. Record completion time
- i. Verify all sheets have subject’s number annotated except for the IRB sheet.
- j. Remind participant to not discuss the project or any of the project details until the data collection is complete and we notify you by email.
- k. Verify email account is annotated on the IRB release form.
- l. Thank You!

Visual Target Detection in Complex Scenes – No Feedback.

Administrative instructions: “You will be presented a series of scenes representing a modern urban combat environment. Each scene will contain between one and six targets. When you detect a target, place the mouse pointer over the target and press the left mouse button once. Do not double click as that will be recorded as two shots by the system and may skew the data. If, before the scenes allotted time is up (15 seconds), you are satisfied that you have visually identified and engaged all targets within

a scene, you may state “Done.” I will ensure that I have manually collected all required data and then ask you to proceed to the next scene. After I state “proceed to next scene,” please depress the enter key once for the next scene to be automatically displayed. Time begins as soon as the scene is displayed. Treat each mouse click as an aimed rifle shot and remember that in “real life” operations, those that see first and shoot first survive. Scan, identify and shoot!

You cannot undo any click and you cannot return to any previous scene. Do you have any questions pertaining to the administration of the test?

You must detect and respond to all targets in each scene within 15 seconds.

Scenario Instructions: Your unit is currently conducting combat operations in the middle eastern rogue nation of Ishtar. You are newly assigned to a rifle platoon operating in al-Basri, and have been in contact with an enemy force that is armed with modern weapons and equipped with personal protective gear to include ballistic helmets and body armor. The enemy units are utilizing sniper and ambush tactics and will fire upon you as soon as you are identified. Your Rules of Engagement allow you to fire upon any target as soon as you visually identify it as an enemy combatant. There are no firing restrictions and ammunition resupply is not an issue. You currently have 100% ammunition level and resupply is already enroute. The villagers have all fled, so the only personnel left in the village are the combatants. All enemy are to your front, all friendly are to the rear. You are “on point.”

Visual Target Detection in Complex Scenes – Feedback.

Administrative instructions: “You will be presented a series of scenes representing a modern urban combat environment. Each scene will contain between one and six targets. When you detect a target, place the mouse pointer over the target and press the left mouse button once. Do not double click as that will be recorded as two shots by the system and may skew the data. If, before the scenes allotted time is up (15 seconds), you are satisfied that you have visually identified and engaged all targets within a scene, you may state “Done.” I will ensure that I have manually collected all required data and then ask you to proceed to the next scene. After I state “proceed to next scene,” please depress the enter key once for the next scene to be automatically displayed. Time begins as soon as the scene is displayed. Treat each mouse click as an aimed rifle shot and remember that in “real life” operations, those that see first and shoot first survive. Scan, identify and shoot! You cannot undo any click and you cannot return to any previous scene. Do you have any questions pertaining to the administration of the test? You must detect and respond to all targets in each scene within 15 seconds. *At the end of the 15 seconds, we will have a 15 second pause during which time I will show you the locations of the enemy combatants on my “master” scene display. After the 15 seconds are up, you will be directed when to proceed to the next scene.

Scenario Instructions: Your unit is currently conducting combat operations in the middle eastern rogue nation of Ishtar. You are newly assigned to a rifle platoon operating in al-Basri, and have been in contact with an enemy force that is armed with modern weapons and equipped with personal protective gear to include ballistic helmets and body

armor. The enemy units are utilizing sniper and ambush tactics and will fire upon you as soon as you are identified. Your Rules of Engagement allow you to fire upon any target as soon as you visually identify it as an enemy combatant. There are no firing restrictions and ammunition resupply is not an issue. You currently have 100% ammunition level and resupply is already enroute. The villagers have all fled, so the only personnel left in the village are the combatants. All enemy are to your front, all friendly are to the rear. You are “on point.”

B. PRETEST QUESTIONNAIRE

Participant Number: _____

Ishihara: 12 8 5 6 5 No Number/Control
Snellen: Left: Right

1. Age: _____

2. Gender: Male Female

3. Manual Dexterity: Right Left

4. Military Rank: _____

5. Years of Military Service: _____

6. How many hours of sleep did you get last night (nearest ½ hour)? _____

a. How many hours of sleep do you average per night?

7. Do you use play video games? Yes No

a. If Yes, Type and Quantity played: _____

8. Do you have simulation experience? Yes No

a. If Yes, Type and Quantity: _____

C. POST-TEST QUESTIONNAIRE – NO FEEDBACK

Participant Number: _____

Please answer the following short questions based upon your experience with this particular experiment. Rating scale is 1 (Strongly Disagree) to 5 (Strongly Agree)

- | | | | | | |
|---|---|---|---|---|---|
| 1. This experience improved my target acquisition skills. | 1 | 2 | 3 | 4 | 5 |
| 2. I improved as I was exposed to more scenes (I improved over time). | 1 | 2 | 3 | 4 | 5 |
| 3. The targets I acquired in the scenes made each follow on target acquisition easier . | 1 | 2 | 3 | 4 | 5 |
| 4. The targets were easy to find in most scenes. | 1 | 2 | 3 | 4 | 5 |
| 5. Feedback as to how I was doing would have helped improve my acquisition skills. | 1 | 2 | 3 | 4 | 5 |
| 6. The ROE information provided helped me concentrate and acquire targets. | 1 | 2 | 3 | 4 | 5 |
| 7. I relied on the hit target flickering red to teach me how the enemy hid in the scenes. | 1 | 2 | 3 | 4 | 5 |

Please do not discuss any portion of this experiment with other students.

-----RETURN THIS SHEET TO EXPERIMENTER----

D. POST-TEST QUESTIONNAIRE – TRAINING-FEEDBACK

Participant Number: _____

Please answer the following short questions based upon your experience with this particular experiment. Rating scale is 1 (Strongly Disagree) to 5 (Strongly Agree)

- | | | | | | |
|--|---|---|---|---|---|
| 1. This experience improved my target acquisition skills. | 1 | 2 | 3 | 4 | 5 |
| 2. I found the feedback feature useful. | 1 | 2 | 3 | 4 | 5 |
| 3. The feedback feature made target acquisition easier. | 1 | 2 | 3 | 4 | 5 |
| 4. The feedback feature made target acquisition faster. | 1 | 2 | 3 | 4 | 5 |
| 5. My ability to acquire the target improved through use of the feedback feature. | 1 | 2 | 3 | 4 | 5 |
| 6. The training gave me the information I needed to make a target location decision. | 1 | 2 | 3 | 4 | 5 |
| 7. I relied on the minimap more than my own visual abilities to acquire the target. | 1 | 2 | 3 | 4 | 5 |

Please do not discuss any portion of this experiment with other students.

E. POST-TEST QUESTIONNAIRE - FEEDBACK

Participant Number: _____

Please answer the following short questions based upon your experience with this particular experiment. Rating scale is 1 (Strongly Disagree) to 5 (Strongly Agree)

- | | | | | | |
|---|---|---|---|---|---|
| 1. This experience improved my target acquisition skills. | 1 | 2 | 3 | 4 | 5 |
| 2. I found the feedback feature useful. | 1 | 2 | 3 | 4 | 5 |
| 3. The feedback feature made target acquisition easier. | 1 | 2 | 3 | 4 | 5 |
| 4. The feedback feature made target acquisition faster. | 1 | 2 | 3 | 4 | 5 |
| 5. My ability to acquire the target improved through use of the feedback feature. | 1 | 2 | 3 | 4 | 5 |
| 6. The minimap gave me the information I needed to make a target location decision. | 1 | 2 | 3 | 4 | 5 |
| 7. I relied on the minimap more than my own visual abilities to acquire the target. | 1 | 2 | 3 | 4 | 5 |

Please do not discuss any portion of this experiment with other students.

THIS PAGE INTENTIONALLY LEFT BLANK

APPENDIX B

A. NORMAL COLORED SCENE EXAMPLES



Figure 18. Normal Colored Scene: Easy



Figure 19. Normal Colored Scene: Hard

B. FALSE (FEEDBACK) COLORED SCENES EXAMPLES

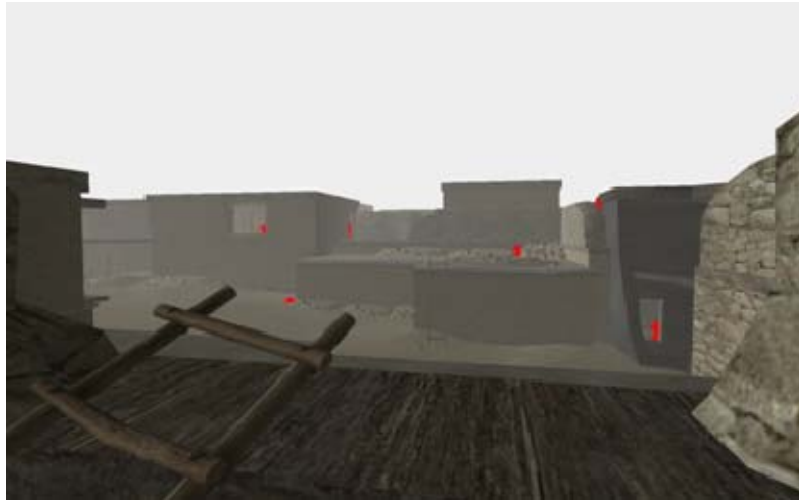


Figure 20. Feedback Colored Scene: Easy



Figure 21. Feedback Colored Scene: Hard

APPENDIX C

A. GRAPHICS RESEARCH PRIOR TO EXPERIMENTATION

1. Appendix C Abstract

This appendix addresses the fundamentals of the graphics pipeline, how it relates to graphics card design, and finally the utilization of common display technology. My intent as to the inclusion of this appendix within the thesis is to assist the reader with comprehension of how Human Factors and technology interact to facilitate improved perceptual learning and training. As virtual environments are becoming more and more capable in replicating the visual environment, the training and perceptual learning of the human subject is improved. Of key importance to me, as the researcher, is the general understanding of how the subject's ability to visually acquire a target in a virtual environment may one day assist with rapidly improving the novice soldier towards a more practiced practitioner and have greater likelihood of a victorious "first contact" in a hostile environment. This experiment's initial design was to utilize commercial off the shelf software (COTS) to facilitate the virtual environmental design and this appendix will briefly address how hardware companies account for software design and user/gamer desires when designing or upgrading their graphics processing units (GPU) (After Polkowski, 2007). Additionally, while this appendix discusses graphics with regards to real-time rendering, the fundamentals for graphically displaying Delta3D's static scenes remains much the same.

Issues relating with the graphics pipeline were encountered when authoring the FPA scenes in create mode that were likely caused by "pipeline" issues and required a robust gaming-type PC. Those issues were not a factor once the scenes were created and saved and had no impact upon experimentation.

As this experiment transitioned to utilizing open source software available through Delta3D, the reader should be aware that many of the same graphics issues still

apply, concerning the graphics pipeline, but more research is necessary to fully understand the issues of licensing and use. That particular topic, however, is not covered within the bounds of this appendix.

2. The Graphics Pipeline

The main function of the graphics pipeline is, given a virtual camera, three dimensional object, some type of light source (ambient, directional, etc.), textures and shading, to render the two dimensional image within the image scene (Akenine-Moller & Haines, 2002).

We also know that the locations, shape, and state of objects in any image is determined by the defined geometry of the object, the placement of the camera and the material, lighting and texture models associated with the image scene. While some would argue that the term “pipeline” is technically incorrect, a good example as to how the graphics pipeline relates to the real world is that of an oil pipeline. Oil cannot move from the first stage of a pipeline until the oil in the second stage has moved to the third stage and so-on. The speed that all the oil is ultimately capable of moving through the pipeline is determined by the oil’s ability to move through each stage’s process so the slowest stage will control the rate of flow regardless of the speed of any other stage (Akenine-Moller & Haines, 2002).

This pipeline is a great analogy as to how “real-time” computer graphics are influenced by the three basic stages of the rendering pipeline: Application, Geometry, and Rasterizer stages. “It is the slowest of the pipeline stages that determines the rendering speed, the update speed of the images...[since it is a pipeline] it does not suffice to add up the time it takes for all the data we want to render to pass through the entire pipeline. If we [find] the bottleneck [we] could compute the rendering speed...or throughput (Akenine-Moller & Haines, 2002).” As stated by the authors:

As the name implies, the application stage is driven by the application and is therefore implemented in software. This stage may, for example contain collision detection, acceleration algorithms, animations, force feedback, etc. The next step, implemented either in software or in hardware, depending on the architecture, is the geometry stage, which

deals with transforms, projections, lighting, etc. that is, this stage computes what is to be drawn, how it should be drawn, and where it should be drawn. Finally, the rasterizer stage draws (renders) an image with use of the data that the previous stage generated (Akenine-Moller & Haines, 2002).

The Application Stage is dependent upon the application and implementation of the software giving the developer the majority of control as to how this stage executes. The utilization of geometry within the design of the software can affect not just the amount of time needed to pass data to the next stage, but also impact upon the eventual quality of the rendered object (Woligroski, 2006).

A common process that occurs in this stage is collision detection as it relates to multiple object proximity. Feedback may be processed as a priority and is assigned the the objects which are given a status or priority as to how and where they are rendered in relation to each other. All processes that are initialized during this stage include texture animation, transforms, geometry morphing, and any initial calculations that are not performed during follow-on stages (Akenine-Moller & Haines, 2002).

From a training standpoint, this stage is important to understand as input-output devices, such as keyboard-mouse or head-mounted display (HMD), are also first introduced to the application during the Application Stage. If the application is not aware of the type of input device being utilized, incorrect geometry and image rendering may affect the ability of a user to interact with a given scene.

The second stage in this 3-stage model is the Geometry Stage, which is responsible “for the majority of the per-polygon operations or per-vertex operations (Akenine-Moller & Haines, 2002).”

This stage can be further divided into multiple sub-stages that are each responsible for a particular function within the geometry process.

The first sub stage will likely define the model and view transform (Akenine-Moller & Haines, 2002). Initially, a model will exist in its own model space so it is not relative to the global space that it will eventually be rendered in. In short, it has not been transformed at all and is relative on to “itself.” The vertices and normals of a model are

transformed by the model transform and only the models that the camera sees are rendered. The camera has its own location and direction within the world space so that it is necessary, to facilitate projection and clipping, that the camera and all the models are transformed with a view transform. The view transform places the camera at the origin and directs its angle of view and the camera's "eye" is looking at the appropriate object from the right perspective (Akenine-Moller & Haines, 2002).

The next sub stage is the lighting and shading process that gives the model a more realistic appearance and can give a scene additional lighting. The models may also have textures applied to them allowing more of a three-dimensional effect and enjoy a more realistic appearance. This texturing of an object can improve the appearance even with minimal lighting changes. Should an object require affectation by light sources, a light equation will likely be applied during this stage enabling the mathematical equations applied by the software to have a realistic interaction with the objects surface.

Additionally, the application of shading (such as Gouraud shading) during this stage will simulate more elaborate lighting effects prior to more elaborate lighting effects that take place later in the process (Akenine-Moller & Haines, 2002).

All these initial processes seem to utilize a similar approach to the "painter's algorithm," discussed in Naval Postgraduate School's MV3202, during the initial calculations the software is beginning to apply effects to the object that will build depth and realism from beginning to end (Peitso, 2007).

Following the lighting stage, the next process that will take place is the Projection Stage, which transforms the view space into a unit cube that will utilize the two projection methods of *Orthographic* projection and *Perspective* projection.

The orthographic view is normally a rectangular box that is then transformed into a unit cube through the orthographic projection transform. The main characteristic of an orthographic projection is that parallel lines will continue to remain parallel following the transform because of the combination of translation and scaling (Akenine-Moller & Haines, 2002).

The perspective projection is more complex in that the farther away an object lies from the camera, the smaller it will appear after projection and parallel lines may actually converge at the horizon. It is in this fashion that the perspective projection mimics how the human eye perceives an objects size (Akenine-Moller & Haines, 2002).

One sub-stage that saves time in throughput is the application of clipping only primitives wholly or partially within the viewing plane are actually passed to the final stage. Primitives that are totally out of the view are no longer required for the scene and are not passed forward along the pipeline process.

The final sub stage within the Geometry Stage is the Screen Mapping process that takes the x-y coordinates of the primitives, following the clipping stage, and applies them to the screen coordinates for correct positioning. Next the z coordinate, corresponding to the x-y coordinate of each primitive, is applied so that there is now a window coordinate that allow the object to be passed on to the next main stage; the Rasterizer Stage (Akenine-Moller & Haines, 2002).

The Rasterizer Stage takes all the data that has traveled through the Application Stage and the Geometry Stage and assigns the correct colors to the pixels to render the image correctly. In this fashion the Rasterizer Stage converts the two-dimensional vertices with their corresponding depth value (the x-y-z coordinates), an assigned color, possibly a texture, and applies them to a pixel on the screen or display device. In this fashion, the Rasterizer Stage is responsible for the pixel-by-pixel operation associated with rendering the image to the screen. Due to the amount of data required to define each pixels most high-performance graphics require that the rasterization process take place in hardware. Through proper application of hardware and buffering, the human eye is able to “view” a fully rendered object and not merely the primitives that actually constitute the object. Utilization of methods such as double buffering or back buffering is primarily utilized so that the scenes are already rendered prior to reaching the display device (Akenine-Moller & Haines, 2002).

The depth of the scene is controlled by the utilization of the z-buffer that holds the value associated with each x-y pair pixel coordinate. When a primitive is being rendered

to a certain pixel, the z-value of the primitive at that x-y coordinate is being computed and compared to the z-buffer at the same pixel. Through mathematical comparison, the renderer can determine that a smaller z-value equates to a primitive that is closer to the camera at that particular pixel. This allows the z-value and color of each primitive being drawn to be updated accordingly. Once the primitives have been rendered application of a texture can further enhance the realism of the object scene. “When the primitives have reached and passed the rasterizer stage, those that are visible...are displayed...[and]are rendered using an appropriate shading model, and they appear textured if textures were applied to them (Akenine-Moller & Haines, 2002).

One issue, worthy of consideration, is that as technology continues to improve “[the] place in this pipeline where software leaves off and accelerations hardware takes over is constantly shifting” is evident when researching the next two topics of this paper: graphics card design and marketing and display technology.

3. Graphics Cards Design and Marketing

Throughout the research for this thesis, there was an ever-present theme of change, “upgrade,” and redesign. Of note were studies that not only discussed the technical aspects of graphics cards, but also how hardware companies determined graphics card design and how the marketplace impacts the capabilities that those cards possess. “Not only do hardware companies work with...developers to help get software products to market...they also work on the back end to make sure the games look and play their best (Polkowski, 2007).” There are four driving forces, or stakeholders, behind the development of the graphics cards that are primarily manufactured: The gamer (users), the developer (aka “coders”), the publisher (EA, etc.), and the hardware maker (AMD, etc.). Due to this interaction between the four stakeholders, there is a certain financial symbiosis that exists and in turn drives the technology of GPUs in specific or certain direction. As stated by the one author:

If it weren’t for the money we spend for entertainment, none of these companies would be in business. Game developers need investment capital to create games we want to play...Publishers are the guys that help seed money to studios to make, mass produce, market, and sell

games...Hardware makers are the folks we need to make these games show up in a way we can experience them (Polkowski, 2007).

One certain advantage that hardware makers have with GPU development is the ability to educate the other players that are involved within the graphics industry. As an example, AMD held seminars focused on DX10 and future changes at the recent gaming convention in Germany where they were able to tell the developers how to best adjust their coding to meet the hardware requirements that would exist six months later. In essence, the hardware company assisted the developers with meeting consumer future demands for hardware that did not yet “exist” (Polkowski, 2007). One might argue that this is leading the developers solely down the path that the hardware companies, either for financial gain or for other reasons, desire to go. This may be true to some extent, but the fact remains that there is also likelihood that through educating the “game” developers, the hardware company is ensuring that the end product (i.e., the scene that the user looks at) is always going to be discernibly better than the scene either rendered through old hardware or by non-optimized software.

Furthermore, the hardware manufacturers create tools to “help improve what is being developed.” These tools will definitely impact graphics pipeline throughput as they assist with building shaders, procedural textures and “any other components of a scene.” The hardware companies even post resources to download from their respective websites that enable developers to constantly upgrade their tools and stay inline with what the manufacture’s GPUs capabilities. For example, AMD’s website has “tools” ranging from performance evaluation to interpreting meshes...the most well known...is probably Render Monkey, a tool that helps compose shaders (Polkowski, 2007).

Further examples of how the hardware manufacturers and developers are continuing to build their relationship are the laboratories that both AMD and Nvidia have developed that are purely devoted to testing current game builds and providing analysis back to developers. Nvidia has over 50 employees testing games on 200 platforms at its Moscow lab and states that it has certified 350 titles, each given a full analysis with recommendations. AMD states that its Boston lab devotes 20 to 50 testing hours per build (not including engineer-to-developer prep time, meetings, and publisher

interaction). While these analyzed products don't necessarily mean immediate improvements to either the gaming experience, or more importantly, to graphics hardware it is readily evident that the interaction will continue to improve both the capability of GPUs and the software design that utilizes them (Polkowski, 2007).

The close ties that the hardware companies are maintaining with the other stakeholders are vitally important. While much of the publicity may try to sway the consumer from buying graphics card A over graphics card B, the fact remains that this symbiotic relationship in the end equates to better capabilities and improved visual interaction. This has allowed hardware manufacturers and developers to take into consideration the abilities of older software while at the same time, as is the case with DirectX, leveraging the new capabilities of graphics programming(see Figure 24) (Woligroski, 2006).



Figure 22. Shading and Textures (From Woligroski, 2006)

4. Display Technology

Historically, most visual interaction with computers remains tied to the two-dimensional (2D) flat screen panels that come standard with our laptop or desktop systems. By design, these display systems render the majority of scenes sent to them through the graphics pipeline in a perceived flat visual “landscape.” Regardless of how fast our three dimensional (3D) GPU passes primitives to our screen, they are doomed to

be displayed in a 2D environment. 2D displays are, largely, incapable of allowing the user to utilize naturally the stereoscopic vision that our physiology grants us (see Figure 25).

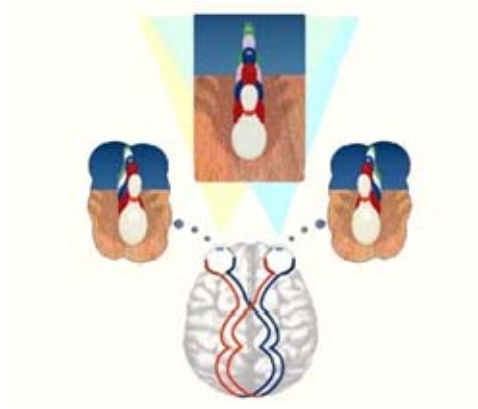


Figure 23. Spatial Perception (From Woligroski, 2006)

Since computer screens are basically flat projection areas, creating a true 3D scene without technological assistance is impossible as the screen will always display a “flat” image that is seen by both eyes simultaneously (Weinand, 2005).

One system that is gaining in popularity is the head-mounted display (HMD) that consists of a “visor” and small screens mounted directly in front of each eye (see Figure 26). One of the main drawbacks to utilizing HMDs is cost as systems range in price from \$800-900 USD (U.S. dollars) for units that have monitoring capability only to about \$24,000 for a HMD such as the nVisor SX that boasts a resolution of 1280X1024 pixels (Weinand, 2005).

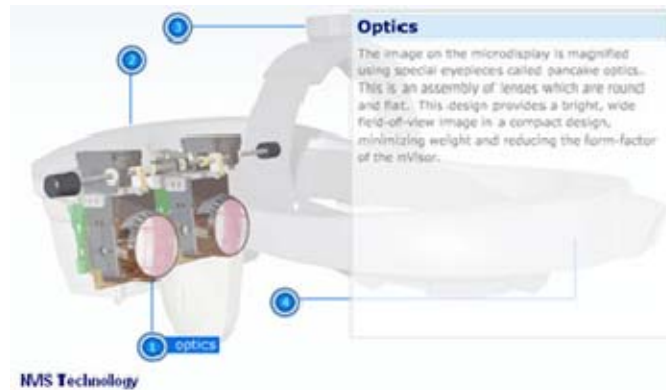


Figure 24. Basic HMD Design(From Weinand, 2005)

While the visual resolution and application of HMDs are ideal for complete immersion in a 3D environment, they are also limiting in their ability for the user to take care of other tasks at the same time. Ultimately, HMDs appear to be advantageous for applications that require excellent spatial representation of an environment, but not for tasks that include both visual and motor skills implementation (Weinand, 2005). This is the same issue that the U.S. Army is facing as it tries to identify what type of system is best suited for soldier training given basic soldier tasks.

Another viable, though not necessarily more economical, option are 3D displays that create a spatial representation of objects. When combined with the proper graphics card, the displays are capable of supporting “a large number of 3D applications” though many will not support the majority of games currently available “with the exception of a few OpenGL titles (Weinand, 2005).”

Much like the HMDs, these displays are not aimed at the whole consumer market and are, consequently, much higher priced than their standard TFT flat-panel counterparts are. Currently, an entry-level system with no tracking system costs about \$4,000, IR tracking models cost \$12,000, and the flagship eye-tracking model costs \$29,500! The pictures are, however, quite interesting even when viewed on a laptop screen (see Figure 27) (Weinand, 2005).



Figure 25. 3D Display (From Weinand, 2005)

The development of 3D technology continues to advance and X3D (not to be confused with our NPS application) has recently announced release of a stereoscopic 17" TFT display for PC gamers that should cost somewhere around \$1000. Hardware support is in place and software is available that will allow 3D software to be displayed in 3D stereo (Weinand, 2005).

The fiscal challenges will continue to be the Achilles heel of both the PC gamer and the military trainer as the primary issue with properly displaying the capabilities of the current generation of GPUs and software continues to be just that: our displays.

5. Conclusion

The real time graphics pipeline continues to increase the throughput of objects and images as more and more of its stages and processes become near simultaneous through buffering and application of multi-card solutions. As more and more algorithmic software capabilities are placed upon hardware, the GPUs will facilitate even more detailed scene images as throughput speed is increased.

Graphics card capabilities will continue to increase over time. Historically, it would appear that a new card (or an updated version) reaches the market place every 9-12 months. Due to the interaction between the hardware manufacturers and the developers there is high likelihood that while these new GPUs may not hit the market any faster, the

developers' ability to utilize the GPUs capabilities to the fullest will undoubtedly be fulfilled (much to the delight of the gamer in us all).

Finally, understanding how the graphics pipeline facilitates the construction of the virtual scene allows us to understand that training human subjects on a specific task must take into consideration the type of visual information displayed. A clear understanding of the task to be trained, the level to which the task will be evaluated, the desired endstate of training (transfer issues), and a estimation of the requirements to fulfill the training objectives must be ascertained prior to creating the first line of code or choosing the engine that will produce the training environment.

APPENDIX D

A. EXPERIMENT LESSONS LEARNED

1. Selection of Engine

Selection of the engine that would provide the rendered scene to the subject was influenced by the availability of both COTS and OSS game engines.

Initially, this experiment was to utilize the Unreal Engine to both design the virtual environment scene (also referred to as a “level”) and evaluate the subject’s ability to visually identify the target.

While the user interface was extremely easy to master, with regards to allowing the author/artist to design a level, and there is a vast library of textures and lighting that assist with rendering a visually realistic virtual scene, issues arose with regards to the AI agent interface.

Specifically, there was no way (found by the author or the course instructor) to make the enemy combatant actually remain still — one might expect a “sniper” to act. The agent would essentially spin in place as its AI attempted to locate the next waypoint to which the agent expected to navigate. This application would likely be advantageous if one were attempting to create an environment to assess the subject’s ability to conduct combat operations (such as a movement to contact that requires meeting and engaging an opponent). My assessment was that it would not allow statistical analysis of the subject’s ability to visually identify a static target of varying degrees of difficulty (i.e., only head and part of shoulder of sniper is visible within a rubble pile).

2. Minimap Utilization

Another facet of perceptual learning that this study was, in the end, unable to ascertain was the augmentation of the subject’s situational awareness. As the Delta3D developers continue to augment and improve their product it was thought that a minimap

application would be readily applied to the False Positive application enabling both friendly (Blue) and enemy (Red) positions location identification by the subject. As most of the military's situation awareness tools now depict known friendly and suspected enemy positions, and technology development seeks to provide networked information that detects enemy locations. We desired to incorporate that situational awareness common operating picture (COP) into the experimental design through application of a minimap function utilized in other Delta3D projects. While many promising advances were made, multiple issues denied the ability to incorporate this particular aspect of the initial experimental design into the experiment conducted for this body of research.

Specific issues arose with the minimap's ability to center on the subject's current position and when the scene was rendered in full-screen mode the minimap would offset. This issue was very difficult to overcome as the user would be required to "fly the camera" over the entire scene until the Subjects Blue position was within the minimap "island" screen (Figure 26).

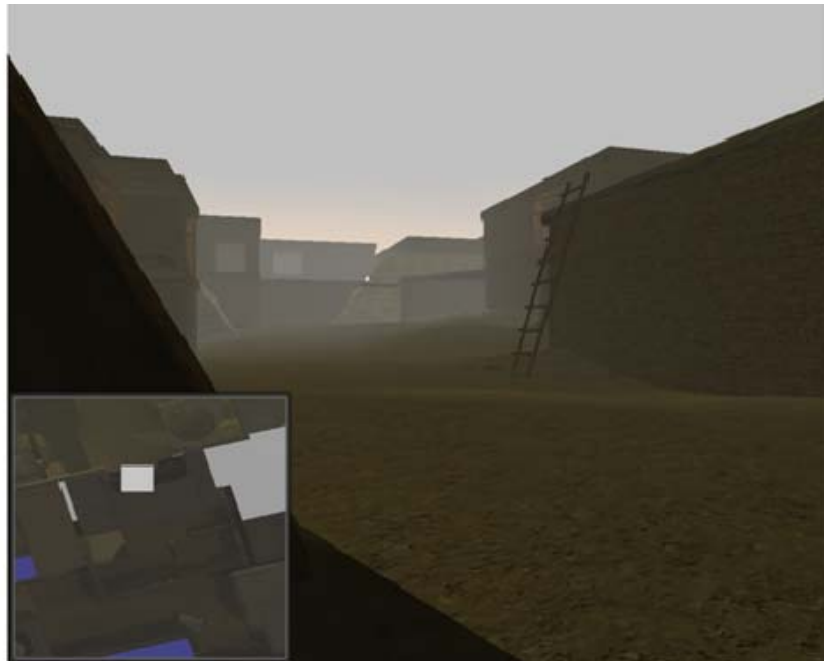


Figure 26. Minimap Centering and Render Error Example

Additionally, all of the scenes designed by myself and previous authors were not usable for the experiment because the minimap application loaded them with contrasts either too low (targets black) or too light (targets white). Additional work to rebuild the scenes in the minimap application proved unsuccessful because the targets inserted into the scene were subject to the same contrasting and rendering issues as the “previous-built” scenes (Figures 27 and 28) .

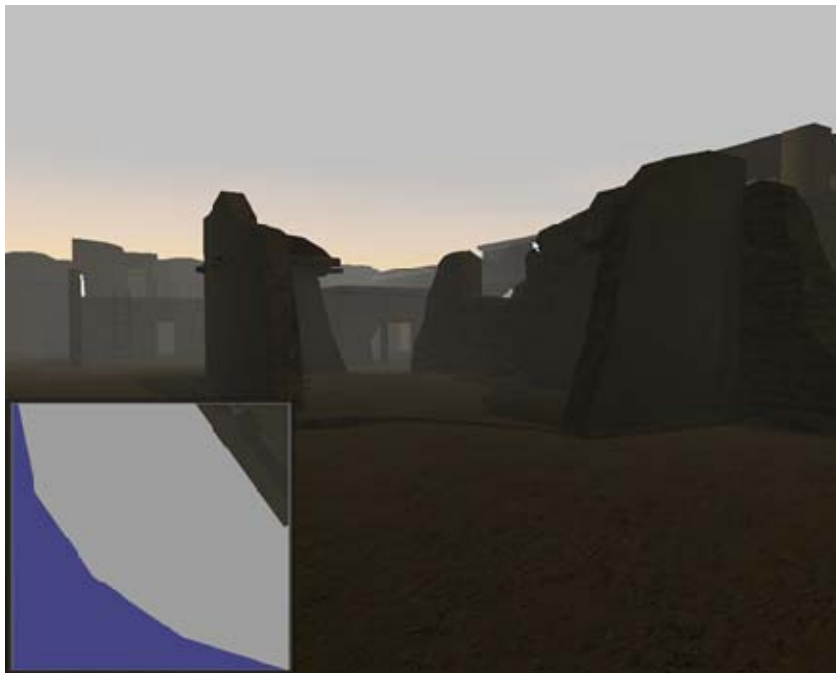


Figure 27. Minimap Scene Render Error Example



Figure 28. Correctly Rendered Scene Comparison to Fig. 27

This may have to do with the shading utilized by Delta3D and an issue with the how the minimap application passes the rgb (red, green, blue) value to the z-buffer for rasterizing. If those values for some reason are becoming more associated with light or dark colors, then the entire target will always be incorrect with regards to shading. As the scenes loaded into the minimap rendered correctly (much like the top picture in Figure 29 below) but the targets rendered similar to the bottom picture, I would gather that incorrect data is being passed between the scene function and the target function within the minimap Delta3D code. At this juncture, however, I must yield that particular finding to those more suited to overcoming that coding issue.

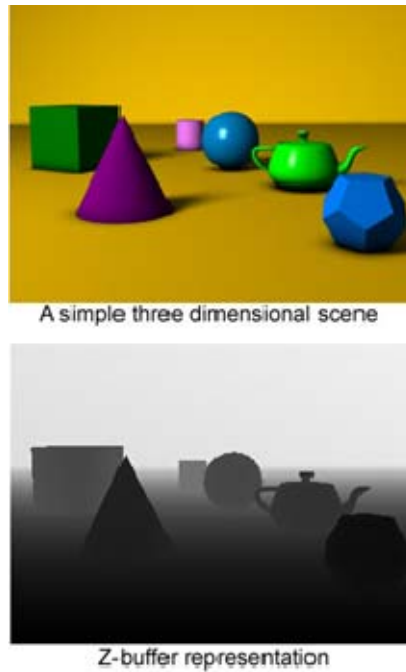


Figure 29. Z-buffer Error Comparison (From Wikipedia, 2008)

My only other thought is that it might be related to the pixel shader, but I will leave further analysis to those interested.

THIS PAGE INTENTIONALLY LEFT BLANK

LIST OF REFERENCES

- America's Army. Retrieved September 9, 2008, from <http://www.americasarmy.com/about/>.
- Ariga, A., & Kawahara, J. (2004). The Perceptual and Cognitive Distractor-previewing Effect. *Journal of Vision*, (4), 891-903.
- Akenine-Moller, T., & Haines, E. (2002). *Real-time Rendering* (1st ed.). Natick, MA: A K Peters, Ltd.
- CALL. (2007). *Soldiers' Handbook (For Official Use Only)* (07-15 ed.). Fort Leavenworth, KS: Center for Army Lessons Learned.
- Ciavarelli, A. P., Asbury, C. N., Salinas, A. Y., Hennesy, R. T., & Sharkey, T. J. (September 2005). *Skill Training Using Adaptive Technology: A Better Way to Hover* (Unclassified No. 2005-10). Arlington, VA: United States Army Research Institute for the Behavioral and Social Sciences.
- Coren, S., Ward, L. M., & Enns, J. T. (2004). *Sensation and Perception* (6th ed.). Hoboken, NJ 07030: John Wiley & Sons, Inc.
- Dynamics, G. (2008). *Land Warrior*. Retrieved March 2, 2008, from <http://www.gdc4s.com/content/detail.cfm?item=aa0d1b86-ac8d-47ed-b59d-f8c2157beb7e>.
- Endsly, M. R. (1996). Automation and Situation Awareness. In R. Parasuraman, & M. Mouloua (Eds.), *Automation and Human Performance: Theory and Applications* (pp. 163-182). Mahwah, New Jersey 07430: Lawrence Erlbaum Associates, Inc.
- Estes, W. K. (Ed.). (1978). *Handbook of Learning and Cognitive Processes Volume 5 Human Information Processing*. Hillsdale, New Jersey 07642: Lawrence Erlbaum Associates, Inc.
- Fahle, M. (2004). Perceptual learning: A case for Early Selection. *Journal of Vision*, (4), 879-890.
- Field, D. J., Hayes, A., & Hess, R. F. (1993). Contour Integration by the Human Visual System: Evidence for a Local Association Field. *172.20.80.102 Vision Research*, 33(2), 173-193.
- G-3, Directorate of Operations & Training & Current Operations. *Infantry One Station Unit Training Cycle Template*. Retrieved June 6, 2008, from <https://www.benning.army.mil/198th/CycleTemplate/Osut%20Training%20Template.htm>.

- Gibson, E. J. (1969). *Principles of Perceptual Learning and Development*. New York, New York 10016: Meredith Corporation.
- Gibson, J. (1969). *The Purple Perils of James Gibson: Does the Ability to Visualize Depend on Visual Images? Cornell University; October 1969*. Unpublished manuscript. Retrieved March 2, 2008, from <http://huwi.org/gibson/visualize.php>; <http://huwi.org/gibson/toc.php>; <http://huwi.org/index.php>.
- Goldstone, R. L. (1998). Perceptual learning. *Annual Review of Psychology*, 49, 585.
- Goldstone, R. L., Schyns, P. G., & Medin, D. L. (1997). Learning to Bridge between Perception and Cognition. *Psychology of Learning and Motivation*.
- Gruenbaum, D. G. (2008). Exploiting Commercial War Games for Teaching, Training and Mission Rehearsal Exercises. *Modeling and Simulation Newsletter*, (Summer), 20-24. Retrieved from <http://www.ms.army.mil/>.
- Hoffman, D. D. (1998). *Visual Intelligence: How We Create What We See*. New York, New York 10110: W.W. Norton & Company.
- Kroger, J. K., Holyoak, K. J., & Hummel, J. E. (2004). Varieties of Sameness: The Impact of Relational Complexity on Perceptual Comparisons. 2004, 28, 335-358.
- Matthews, G., Davies, D. R., Westerman, S. J., & Stammers, R. B. (2000). *Human Performance: Cognition, Stress and Individual Differences*. 27 Church Road, Hove, East Sussex, BN3 2FA: Psychology Press.
- McDowell, P. (2008). *Delta3D Open Source Gaming and Simulation Engine*. Retrieved June 6, 2008, from <http://www.delta3d.org/index.php>.
- McLeod, S. A. (2007, July 22, 2008). Simply Psychology. Message posted to <http://www.simplypsychology.pwp.blueyonder.co.uk/perception-theories.html>.
- Moulaoua, M. (2006). Accumulation and Persistence of Memory for Natural Scenes. *Journal of Vision*, (6), 8-17.
- Olman, C., & Dersten, D. (2004). Classification Objects, Ideal Observers & Generative Models. *Cognitive Science*, (28), 227-239.
- Peters, G. A., & Peters, B. J. (2006). *Human Error: Causes and Control*. Boca Raton, FL: CRC Press.
- Polkowski, D. (2007). *Hardware Companies Outside the Box*. Retrieved September 11, 2007, from http://www.tomshardware.com/2007/09/11/hardware_companies_outside_the-box/.

- Program, L. W. (2005). *Land Warrior & Stryker Warrior Programs*. Retrieved February 14, 2008, from <http://www.defense-update.com/features/du-4-04/land-warrior-1.htm>.
- Reason, J. (1990). *Human Error*. 40 West 20th Street, New York, NY 10011: Cambridge University Press.
- Sadr, J., & Pawan, S. (2004). Object Recognition and Random Image Structure Evolution. *2004*, (28), 259-287.
- Shattuck, L., & Miller, N. (2006). Extending Naturalistic Decision Making to Complex Organizations: A Dynamic Model of Situated Cognition. *Organization Studies*, 27(7).
- Shirley, P., Ashikhmin, M., Gleicher, M., Marschner, S. R., Reinhar, E., Sung, K., et al. (2005). *Fundamentals of Computer Graphics* (Second ed.). Wellesley, MA: A K Peters, Ltd.
- Smallman, H. S., & St. John, M. (2005). Naive Realism: Limits of Realism as a Display Principle. Paper presented at the *Proceedings of the Human Factors and Ergonomics Society 49th Annual Meeting*, Santa Monica, CA.
- Thomas, L. C., & Wickens, C. D. (2006). Effects of Battlefield Display Frames of Reference on Navigation Tasks, Spatial Judgments, and Change Detection. *Ergonomics*, 49(12-13), 1154-1173.
- Vaughan, B. D. (July 2006). *Soldier-in-the-Loop Target Acquisition Performance Prediction Through 2001: Integration of Perceptual and Cognitive Models* (Unclassified No. ARL-TR-3833). Aberdeen Proving Ground, MD 21005-5425: Army Research Laboratory.
- Wainwright, R. K. (2008). Look again: An Investigation of False Positive Detections in Combat Models. (Master of Science in Operations Research, Naval Postgraduate School). (Unclassified).
- Weinand, L. (2005). *3D Stereo Technology: Is it Ready for Prime Time?* Retrieved September 12, 2007, from http://www.tomshardware.com/2005/05/3d_stereo_technology.
- Wickens, C. D., Lee, J., Liu, Y., & Becker, S. G. (2004). *An Introduction to Human Factors Engineering* (2nd ed.). Upper Saddle River, New Jersey, 07458: Pearson Education, Inc.
- Wilken, P., & Ma, W. J. (2004). A Detection Theory Account of Change Detection. *Journal of Vision*, (4), 1120-1135.

- Woligroski, D. (2006). *Graphics Beginners' Guide, Part 2: Graphics Technology*. Retrieved September 12, 2007, from http://www.tomshardware.com/2006/07/31/graphics_beginners_2/.
- Yu, C., Klein, S. A., & Levi, D. M. (2004). Perceptual Learning in Contrast Discrimination and the (Minimal) Role of Context. *Journal of Vision*, (4), 169-182.

INITIAL DISTRIBUTION LIST

1. Defense Technical Information Center
Ft. Belvoir, Virginia
2. Dudley Knox Library
Naval Postgraduate School
Monterey, California
3. Dr. Michael E. McCauley
Naval Postgraduate School
Monterey, California
4. CDR Joseph A. Sullivan
Naval Postgraduate School
Monterey, California
5. MAJ Robert L. Kammerzell
I Corps Headquarters
Fort Lewis, Washington