

why

The Philosophy and Science of Multivariate Reasoning

Hal Campbell, Ph.D.

Copyright © 2010

All rights reserved.

CONTENT AND LIABILITY DISCLAIMER

The author and publisher shall not be responsible for any errors or omissions contained within this book and further reserve the right to make changes without notice. Accordingly, all original and third party information is provided **"AS IS"**.

In no event shall the author or the publisher be held liable for any damages whatsoever, and in particular the author and publisher shall not be liable for special, indirect, consequential, or incidental damages, or damages for lost profits, loss of revenue, or loss of use, arising out of or related to this book, whether such damages arise in contract, negligence, tort, under statute, in equity, at law or otherwise.

Library of Congress Cataloging in Publication Data:

Campbell, Harold, 2010

ISBN 978-0-557-35648-5

Printed in the United States of America

Note to Faculty

This book has been specifically written and edited to support classes dealing with the subjects of critical thinking and reasoning, introduction to scientific methods, or basic research methods. A variety of examples and references have been incorporated specifically to support students pursuing degrees in the sciences, social sciences, and liberal arts. The writing style and grammar used in the book have been purposely tailored to support the needs of undergraduate education, but can also be valuable in support of a graduate level review course. The concepts presented, as well as the presentation style and explanations should be well within the capability of most freshman and sophomore level undergraduate students. The Flesch-Kincaid Grade Level for this textbook computes to 16.7.

The book can be used as a primary text in support of critical thinking and basic research classes, or it can be used as a supplementary text to help elevate the level of student understanding pertinent to the principles and concepts contained herein. The book has been structured to support semester based courses. Faculty teaching in quarter based programs will likely discover that using the first ten to twelve chapters (in sequence) will garner the best results in communication of critical level information and student comprehension.

Note to Students

As you commence your review of this book, please know that it was written specifically with you in mind to help you understand a very difficult, but incredibly important concept of [multivariate] reasoning. I strongly encourage that you endeavor to not only familiarize yourself with the principles contained herein, but that you take the time to master these concepts. That will require practice. Multivariate reasoning is probably the most important thing that you will learn during your college career. Everything else depends upon it. Without a solid understanding of how to formulate arguments, how to evaluate information for truthfulness, as well as the ability to derive conclusions that make sense, all of the other things that you will learn in college become ineffectual. They become just facts with no foundation to pin them on to facilitate your understanding. College is an exciting time, but we should never lose focus of the process involved in higher education. The pedagogy involves teaching you to reason and then providing you with an education regarding all of the facts associated with your major area of study. This Socratic Method empowers you to apply all of the facts you've learned during your undergraduate studies within a context and provides a familiarity with the logical process that allow you to formulate your own conclusions about the world and support your efforts to make new discoveries.

I sincerely hope that the information contained within this book helps you in your journey toward understanding the world through multivariate reasoning.

About the Author

Dr. Hal Campbell received his Ph.D. from Claremont Graduate University in 1983. Professor Campbell selected Claremont based on the university's comprehensive doctoral studies requirements for graduate students to specialize in multiple academic disciplines. During his time at the university, Professor Campbell completed rigorous studies and academic specializations in the areas of statistics and empirical analysis, systems analysis and operations research, public law, and philosophy. He spent over ten years in planning and analysis with the County of Los Angeles prior to joining the ranks of academe. From 1989 to 2009 he served as a Professor for the California State University, College of Natural Resources and Sciences, Department of Mathematics and Computing Science. In addition to his affiliation with the California State University, Professor Campbell has also been retained by several public and private universities to teach in the area of statistics, law, and justice studies. Throughout his career, Dr. Campbell has authored a significant number of scholarly publications that address a wide variety of issues and disciplines. These include scientific papers that provide insights and perspectives on alternative energy development, global warming and climate change, terrorism and national security, the application of technology to health care, computer generated camouflage and defensive systems, distance learning and technology mediated instruction, educational administration and planning, and a variety of writings dealing with the subject of multivariate reasoning and logic. He remains active in higher education through his involvement in distance education.

To all of those people who have crossed my path throughout a lifetime of discovery and shared your insights about those things that really mattered.

Lyle you are at the top of the list. Thank you.

Contents

<i>Chapter 1</i>	<i>Introduction</i>
<i>Chapter 2</i>	<i>Argument Structures</i>
<i>Chapter 3</i>	<i>Introduction to Multivariate Reasoning</i>
<i>Chapter 4</i>	<i>Vertical and Perpendicular Logic</i>
<i>Chapter 5</i>	<i>Scientific Evaluation Process</i>
<i>Chapter 6</i>	<i>Proving a Premise with Z ratio</i>
<i>Chapter 7</i>	<i>Proving a Premise with Student's T</i>
<i>Chapter 8</i>	<i>Proving a Premise with Chi Square</i>
<i>Chapter 9</i>	<i>Proving a Premise with Correlations</i>
<i>Chapter 10</i>	<i>Multivariate Correlations and Discriminant Function Analysis</i>
<i>Chapter 11</i>	<i>Applied Spatial Modeling</i>
<i>Chapter 12</i>	<i>Conclusion and Final Thoughts</i>

Preface

As the title of this book implies, the answer to the question “why” is often elusive, and is discovered only after the thoughtful deliberation of the many factors and influences related to the topic under deliberation. From the extraordinarily complex to the perceptively simplistic, the search for truth requires a painstaking process of contemplation, research, theorization, hypothesis formulation, data collection, evaluation, and assessment of the results. It is only after all steps have been taken into consideration and empirical evidence is analyzed, that one is safe in rendering a final judgment.

Truth is not at all relative, but instead is absolute, most of the time. What is relative is the amount of time and energy that most people are willing to devote to the quest of being correct in their judgments. Many times, people simply grab hold of the first reasonable “univariate” explanation that occurs to them, rather than suspending judgment until they have thoroughly examined all possibilities and potential influences that could be attributed to the phenomena. Needless to say, this approach rarely results in the attainment of truth.

This book is devoted to examining the processes and methodologies used in complex analysis. The intent of this text is to support the efforts of students, as they endeavor to learn about critical thinking and to help them achieve a solid understanding of the scientific methodologies related

to reasoning. The concepts and techniques examined within this book relate directly to the search for truth, using logic and critical reasoning, as opposed to speculation and supposition. Unlike other texts, the processes and methodologies articulated here will focus exclusively on the more scientific forms of reasoning used to differentiate mere speculation from proven understanding, using scientific methods.

One's ability to master critical thinking and reasoning is an evolutionary process. It is only after years of practice at employing the techniques outlined within this book that a person can garner any degree of confidence in their ability to employ the scientific principles associated with isolating the truth. This is not to suggest that they will, with any degree of predictability, conjure up the inherently correct answer, rather that they are prodigious in employing scientific reasoning to minimize the probability of their being incorrect. Sometimes truth is elusive no matter how scientific or methodical you try to be. Circumstance does not always permit the thoughtful evaluation of all the factors and variables exerting influence in an equation, or it may be that knowledge has not evolved to a point where it allows you to reach the correct answer, but even a person's cursory judgments can be greatly enhanced through the adoption of scientific reasoning methods.

Hopefully, this text will help to refine the reasoning skills of those familiar with its contents and lessen the likelihood that they will fall victim to illogic.

Chapter 1

The Importance of Perspective

Throughout human history (perspective) has played a significant role in shaping ideology, influencing social beliefs, guiding scientific discovery, and determining our value systems. Perspective is a complex issue because each of us develops our individual perspectives based on a wide variety of factors, which are predicated upon a lifetime of experiences, and perceived knowledge. Perspective is shaped by aggregating all of the facts and discoveries that we possess about a particular issue and then forming conclusions based on the information and understanding that we have amassed since our birth. We (hopefully) then integrate this knowledge within a decision making process to derive conclusions and beliefs. The result is our perspective on the issue.

What is important to recognize about perspective is not that it is an ending point, but rather, that our perspective also serves as a point of departure for all of our subsequent decisions and judgments about future issues. We use perspective as an anchor point for our beliefs about the world and gauge each new fact or discovery relative to its pertinence and impact upon our existing core beliefs (or perspective). Unfortunately, perspective (or the lack thereof) can be a contaminating influence in the search for truth. People who firmly believed that the

earth was flat made subsequent judgments about our importance within the universe, religion, and political beliefs based on this fallacy. Values were formed that placed the earth in a special position within the universe. Religious dogma was shaped to reinforce the belief that since the stars revolved around the earth then our planet must be the center of all things, and extrapolated this to mean that we humans were exceptional creatures in comparison to all others and as such, merited a position of supreme importance. Obviously, later discoveries by Copernicus and Galileo altered this view of the universe, resulting in a need to reassess everything we knew to be true and absolute one minute prior to their discoveries. Many however simply could not accept these revelations and rejected their implications because the significance of this new information stood in direct conflict with their long held perspective. For many, delusion was preferable to enlightenment.

Needless to say, perspective matters in the search for truth and not just by those that seek scientific discovery, but also in all forms of deliberation by those who must integrate new discoveries within their own belief systems. Social values and belief systems of all kinds are shaped by perspective. In fact, many of the atrocities perpetuated by one civilization upon another have been directly influenced by perspective, or the lack of perspective. Demonization of an entire culture based on misunderstanding and intolerance for differing social values is not an uncommon event and has resulted in a

significant number of conflicts and atrocities throughout human history. Judgments of all kinds are influenced by perspective. These include determinations about right from wrong, good versus evil, acceptable and unacceptable behavior, and even the inherent value of scientific discoveries versus a perception of insignificant relevance based upon the failed recognition of the importance of new information. Every judgment we make is reached (at least in part) based on our perspective about the issue and the value that we perceive that the new information has relative to the core values that we presently embrace. Not until a thorough examination of these influences is conducted and assimilated, can we reassess our conclusion about a specific phenomenon, which in turn becomes our new perspective for the next evolution in the never-ending process of discovering the truth.

Determining beliefs about what is truth and what is fallacy is a much easier process for those who have not yet formed an opinion or perspective about an issue, provided they are open minded. Students who are learning information for the first time about a specific topic are probably more empowered than the rest of us, because they lack the impediments to learning possessed by those of us who have been studying the subject for years. The reason for this is based, in part, on the previously mentioned statement about perspective serving as a point of departure in the assessment process. It makes sense that if you have not yet formed an opinion about a topic,

then you are less likely to bring a preconceived notion or perspective about the issue. Subsequently, your judgments about the issue are less likely to be contaminated by bias or fallacious insights. You can see evidence of this notion all around where students of a subject are less likely to present obstacles to learning and more readily accept the relevance of new discoveries, as opposed to scientists who have been studying the subject for years. Imagine the difference in the process of accepting the implications of a new found discovery between a new student and a professor of that subject. If the discovery is so significant that it serves to mitigate many of the conclusions previously accepted as gospel by the scholarly community, then resistance is automatic by many of those seasoned people who must contend with re-evaluating the implications of new information as it relates to the beliefs that were previously held, based on that misunderstanding. In contrast, a new student of the subject can more easily embrace the discovery and its significance without the requirement for a total reconsideration of everything they knew to be true one minute before the revelation. It is no wonder that some of the greatest scientific minds chronicled in the history books were either put to death or banished from society because of the impact their discoveries had on unraveling the fabric of social beliefs and values. This same process happens today as new discoveries are made that force us to reconsider just how misinformed and incorrect we have been about things for most of our lives. The good news is that we are less prone to be put to death for the merits

and impact of our scientific discoveries. The bad news is that there are a plethora of suppositions and presumptions that are untrue flowing through our civilization at the speed of electrons and it is difficult, if not impossible, for most people to differentiate truth from fallacy.

An excellent demonstration of the importance of perspective can be found in the question, “how many directions are you moving, this very minute”. I have used this question for decades to profile the importance of perspective in critical thinking and invariably, when I ask this question of a classroom full of students, I get a plethora of responses. Mostly the students simply guess and shout out a number that they hope is correct, but when I challenge their assertion by saying, “you’re not moving at all are you, because you’re sitting in this classroom”, they unanimously agree that it cannot be possible to be moving while sitting still. This is the same group mindset and group dynamic that has caused countless cultures to form inaccurate conclusions and values about the world, since the dawn of humanity. As I explain that their perspective of sitting still and motionless is wrong, and in fact, they are moving in four distinctly different directions at once, they begin to notice a fracture in their collective perspective. I start by reminding them that they are moving in an arc around the planet as it orbits its axis at about nine hundred miles per hour, as well as moving around the Sun as the earth travels in its orbital path at nineteen miles per second. To make

matters more complicated, the Sun is traveling around the spiral arm of the Milky Way galaxy and travels about one million miles a day, while the galaxy itself travels away from the center of the big bang. Therefore, essentially, each of us is moving in four different directions at the same time, while sitting perfectly still. Our perspective however is limited by our failure to sense any of these motions and it's not hard to imagine how it is that people conclude (falsely) that they are anchored firmly to their seat and sitting perfectly still.

The importance of this exercise can be found in pointing out that false conclusions are an easy thing to fall victim to and all too often, many of our perspectives are found to be based on previously inaccurate premises. As scientists however, it is important for us to remember that perspective matters in determining absolute truth and we need to be careful to avoid falling victim to haphazard acceptance of prior beliefs, as we endeavor to extend the reach of human knowledge. Nothing should be taken for granted and everything that we think we know for certain should be reassessed within the confines of our experiments. Simply having our eyes open does not assure the attainment of truth, especially if our minds are closed because of limited perspective. If we fail to fully understand the point of departure we occupy before we search for new discoveries, and if that point of departure is predicated on a fallacy that we have accepted as true, then our perspective is inaccurate and everything we build on top of it is contaminated and incorrect. This is precisely

why professors are bound by the covenant of never telling a student something that they are not certain of themselves as factual. The consequence will likely be the acceptance of a fallacy into the knowledge base of that student, which will contaminate all future judgments because of the erroneous nature of the information.

Emotional Intelligence

Emotional Intelligence (EI) is also a critical component of the logic process and is as important to the endeavor as dedication, commitment, and vision. Basically, it is the ability to process emotional information, particularly as it involves the perception, assimilation, understanding, and management of emotion. The very nature of logic presupposes that scientists will approach the endeavor in an objective manner and will, at all times, exercise open mindedness, inquisitiveness, and unbiased perspective as they search for knowledge.

Emotional Intelligence has a significant impact in any setting, but we are specifically interested in its applications in the search for truth and scientific discovery. Not only does each individual have an Emotional Intelligence Quotient (EIQ), but whenever a team of researchers is created, that team possesses a collective EIQ as well. Finding the proper balance of excitement and energy, while assuring an optimal level of emotional intelligence is maintained, is often a significant challenge.

Pure logical ability, as far as I can tell, does not exist here on earth. Perhaps on another planet in the universe pure logic prevails, but there has been no sign of a transfer of that ability here on earth (yet). People instinctively wrestle with the conflict between logic and emotion in virtually every human endeavor. The best we can hope for is a constant struggle to keep our emotions in check and out of our research, until well after the discovery of truth has been confirmed. Then, and only then, is it acceptable to become passionate about our discoveries and the implications they may hold for humanity.

It is essential, given the relevance of emotional intelligence to contamination of the research endeavor, to maintain sufficient checks and balances are in place to guard against the incorporation of emotional attachment to any decision that we make about discoveries. Great care should be taken not to initiate research in order to “prove” a point, nor is it appropriate to bend the interpretation of our results so that our findings confirm a suspicion that we have long held. Dispassionate objectivity is critical to assuring that we avoid the pitfalls of making an inaccurate conclusion. We have all seen instances where it appears evident that the researcher started out to prove their biased point and structured the research design so that it would achieve the objective of proving their suspicion. We have also seen instances where, despite the data, they have interpreted the findings incorrectly because of a passion that they held about the topic. In all such instances, we find that not only the results of the study are

without merit, but that the effort reflects poorly on the reputation and trustworthiness of the researcher. The same conclusions can be applied to people who manifest outlooks on non-scientifically oriented decisions. No matter the scenario, dispassionate objectivity and assuredness of avoiding contamination by a lack of emotional intelligence is vital in acquiring the truth.

Imagination and Vision

Imagination and vision are difficult qualities to define within this context, mostly because of the intangible nature of the subject, but research and logic thrive on both of these attributes in order to attain substantive breakthroughs in understanding. All too often however, scientific research is relegated to sequential advancement where the explorer is caught in the endeavor of incremental discovery because of their limited mindset and process orientation, as opposed to achieving quantum leaps of understanding that push back the veil of knowledge to a new level in a single stroke. Yet, those scientists throughout history that stand out as the true visionaries of their time, were those who (through their imagination and vision) constructed hypotheses that were demonstrably decades ahead of contemporary thought and offered insights that escaped the contemporary thinkers of their time. Often ridiculed for their perspectives and theories, these giants were not constrained by process-oriented thinking, and dared to imagine those things well out of reach and set course for a

lifetime of study to evaluate the evidence that would either confirm or invalidate their suspicions. Sometimes, they died before their theories were proven, but advances in technology occasionally confirmed their suspicions.

Many notable scientific philosophers and theorists such as Copernicus, Galileo, Bacon, Bohr, Newton, and Einstein all accentuated that region well beyond the outer edge of contemporary thinking and understanding in search of revelations in knowledge. They accomplished this, not by seeking incremental advances in understanding based on process-oriented methodologies, but through contemplation of the holistic framework of the phenomena and then applying their “mind’s eye” to seeing the entire realm of possibilities. Once they had envisaged the gamut of theoretical possibilities, they employed scientific principles and experimentation to prove the truth or fallacy of their theories, and in turn, validate or invalidate their hypothetical assertions.

We will discuss the importance of the “visualization” process in great detail within this text, but it’s important to note here the value of stepping back from the problem, seeing the whole board, and then examining the relationships and interactions of variables to ascertain the total picture of how things actually interact with one another. This is not an easy thing to teach. In fact, I am not certain that it is possible to teach it. Some argue (convincingly) that vision is a trait and not a learned characteristic, but I believe that everyone can advance

their skill in this area and that all of us have the ability to enhance our capacity to use the mind's eye to envisage the entire realm of possibilities. Such an approach makes it possible to hypothesize about all of the variables that could affect the outcome and to structure scientific experiments that seek to uncover the truth of each premise before coming to a conclusion.

The advantage in such an approach is that we avoid limiting our understanding to just one or two variables and (because of our broad vision) entertain the possibility that multiple factors may be exerting influence. Then, we can isolate each variable, formulate a research and null hypothesis, and test the truth of our beliefs. Ideally, our equations would contain all of the relevant influences and we could then assert with conviction that we had accounted for each and every meaningful factor in the equation.

As you progress through this book, you will encounter a series of references to this visual modeling approach and many of the tests for truth of the premises will be predicated on this mode of thinking. As previously stated, incremental minds make incremental discoveries and render incremental contributions, but it is the holistic thinker that sees the broad spectrum of possibilities and maintains the ability to step back and take note of all of the possible influences, which in turn renders the truly distinctive contribution to scientific discovery. Without such a methodology, the researcher merely gropes along

in their process-oriented world, in hopes of stumbling upon a meaningful discovery. I cannot imagine such an existence.

The Importance of Questioning

Many of us have been annoyed to our wits end by the constant questioning of small children, who habitually ask, WHY? No matter our answer to their previous question, they follow up with another shriek of, WHY? Finally, when we cannot take it for a moment longer, we resort to the age-old adage and respond, “Because I Said So”, which usually means that we do not know the answer, but it serves to quell their exclamations. We are a naturally inquisitive species from birth and it is during our formative years, where there are no social expectations for us to know the answer to any question, and during these times that we feel most comfortable asking the question, why. It is not until later in life that we seem to lose the passion for exclaiming this simplistic inquiry of others when engaged in conversation. What a shame. There are undoubtedly a plethora of social influences that curb our use of this word [why]. Some explanations are probably relative to the expectation that, by a certain age, we should know the answer. Yet others are based on the fact that we encountered people who told us that asking [why] was annoying and to be socially acceptable we stopped asking, but the fact of the matter is that we should never stop asking.

Within this chapter, we most certainly need to examine the importance of questioning relative to uncovering the truth. Questioning of one's own views about things is also important. So is questioning the logic offered in support of the premises and conclusions proclaimed by others as they express their positions. Without questioning we are relegated to the distinct possibility that we might sheepishly accept assertions that are unsubstantiated or even worse, accept arguments that lead us to a false conclusion. Likewise, failure to question our own beliefs, and those reasoning processes that we used to arrive at our point of view, is equally ill advised because it opens the door to delusion.

Questioning is a healthy component of the critical thinking and reasoning process. We should never lose sight of its value in the search for truth. We are obligated to question the logic of a premise in an argument, the relevance of a proven truth to the conclusion, and whether or not the argument presented is factual, relevant, and correct. This becomes increasingly important as applied to the more complicated logic processes that we will examine later in this book involving multivariate reasoning, vertical logic, and perpendicular influence.

In a simple argument (one that involves one premise and only one conclusion), it is a relatively simple matter to assess the proof of a premise and its relevance to the conclusion. Deductive reasoning often relies heavily on limited or simplistic logic to arrive at a conclusion.

Inferential logic on the other hand can typically involve a larger number of premises that require scientific testing, and then arrangement in an order to form a basis for the conclusion. In such equations, it is essential to question the basis for each premise, the merits of its inclusion within the equation, the truth or fallacy of its assertion, and then finally its individual and aggregate relevance to the conclusion.

Questioning is simply a critical and key component of the logic process. We should question our own beliefs, departure points, motives, bias, and intentions as well as subjecting another's argument to the same process of critique in order to assure that we have illuminated all of the pertinent factors in our decision making process. As scientists, we should not reject attempts to question our logic by others, but instead see it as an opportunity for us to present our research and logic to the scientific community for scrutiny. From this scrutiny, we will either be rewarded with validation of our ideas or we could receive insights relative to shortcomings of our methodologies. Either way, we are better positioned to attain the truth, which is, after all, the ultimate goal of our endeavor.

Let me provide a practical example of the importance of questioning in making a medical diagnosis. Several levels of questioning occur during the process of diagnosis and treatment. The purpose of such questioning is not only to narrow down the plethora of possible medical afflictions in

order to properly treat the patient, but also to assure that the physicians diagnosis is validated through objective peer review in order to confirm initial suspicions.

When a patient first arrives at a medical facility, the doctors begin the process of questioning by eliciting from the patient a description of the problem. They ask them to describe the problem and articulate the symptoms. From the information they receive, they begin the process of narrowing the broad range of possibilities contributing to their affliction. If, for example, a patient complains of chest pain, shortness of breath, and fatigue then the preliminary indications could be related to the cardiovascular system. It could also however be related to a problem in the respiratory system or a combination of several systems. More information is need before a specific diagnosis can be formulated. Information relative to the patient's age, prior history, family history, weight, medicines taken, illicit drug use, whether the pain is periodic or constant, whether it is more acute during exercise and other factors are also obtained through questioning. If for example, the patient is a male, sixty-five years old, with a history of family heart disease, who is overweight, not physically active, and who smokes two packs of cigarettes a day, the process of questioning tends to provide overwhelming evidence of a potential cardiac event. The diagnostic process continues and with each new piece of evidence obtained through questioning the list of all potential maladies is reduced to the most probable affliction. It's important to recognize however

that despite the physician's presumption as to potential causes of the problem, they do not treat the patient (typically) based on a suspicion. To assure their accuracy a series of tests are ordered for the expressed purpose of validating suspicions (or confirming the hypotheses) as to what might be causing the symptoms. Following the diagnostic process and formulation of a "best guess" that is based on observation, questioning, and testing, a treatment is prescribed.

Questioning, as you can see, is used to illuminate factual specifics and often leads to the eventual attainment of the truth. Without in-depth questioning, we limit our ability to discern important factors, ponder relevant considerations, and formulate the most probable conclusion. As stated earlier, questioning is an essential element of the critical thinking and reasoning process. We should embrace it as a valuable commodity that furthers the likelihood of the accuracy of our conclusions. Questions lead to suspicions and hypotheses, which can and should be scientifically tested to discern their accuracy. The combination of these two approaches (questioning and scientific testing) can provide validation of the truth, which as we see in the medical example, leads to a "best guess" that guides treatment options.

Judgment

You have undoubtedly heard countless references to the importance of judgment in human endeavor. After all, it is

a significant measure of one's ability to make rational decisions. Although a nebulous term that is not easily defined, it refers to a person's ability to render conclusions based on the objective review of pertinent information and then to render a decision that is commensurate with the conclusions they derived.

Like everything else in the world, a person's ability to judge is a multivariate issue. In other words, it should be based upon a variety of factors that combine to influence their proficiency at forming accurate and sound conclusions. To make it even more complicated, a person's judgment is not static. It evolves and grows (hopefully) over time as they collect, synthesize, and process new informational elements, experiences, and facts.

It is important (I believe) to call your attention to the notion that sound scientific practices for formulating hypotheses, collecting data, testing the merits of assertions, and then evaluating the results may NOT have direct relevance to judgment. Scientific protocols can serve to improve a person's decisions and chance of being correct, and sustained exposure to methodological processes can augment one's abilities to isolate potential factors that contribute to the outcome, but such familiarity and proficiencies do not necessarily mean that a person's judgment is better. Judgment (although related to this process) is an entirely different capability that is partially predicated on experience, mental health, wisdom,

ego, outlook, setting, option identification, skill at assessing reaction, and the ability to foresee the consequences of action. Simply because someone is schooled in the scientific methods and can use these abilities effectively to produce an argument, it does not necessarily mean that they possess sound judgment, nor judgmental ability that is preferential to others.

Many institutions of our society such as the legislature, appellate and supreme courts, universities, federal, state, and local commissions, as well as corporate boards and others recognize that judgment is not within the exclusive purview of a select few, but is significantly enhanced through collective assemblies. The reason for this is not that one person cannot render an effective decision, but rather that a collective assembly of minds (hopefully that are all well schooled in decision sciences) enhances the probability that all relevant factors are considered and a judgment rendered that considers all germane variables and possible consequences. Ideally, consensus relative to the facts and the implications of the events would be a product of collegial review and the decision/s rendered would be seen by all as the best course of action. Moreover, the formation of review bodies only enhances the likelihood that one of the members of the commission will possess sufficient wisdom to see the truth and guide the others in the right direction. It doesn't guarantee it.

We will talk more about judgment in future chapters, but remember the adage "with age comes wisdom" is not

necessarily true. There is no statistically significant correlation coefficient between age and wisdom. It is a multivariate equation.

Chapter 2

Argument Structures

Critical thinking and reasoning relies almost exclusively on argument structures to support the process of discovery. As you will recognize throughout this book, arguments are not mere disagreements between two people with opposing views, although that is often the most common synonymic explanation. Rather, as applied to critical thinking and scientific reasoning, arguments are expressed as a series of propositions that are fashioned into declarative statements (i.e., premises), which in turn, support a specific conclusion.

Whether you are engaged in formulating personal values and judgments about religion, politics, and social values, or whether you are venturing near the edge of scientific discovery and endeavoring to describe the most intricate interrelations of the universe, the same process of argument structuring should be invoked. That is to say, using a series of premises or statements of truth that possess independent accuracy and precision, and which (individually and collectively) lead to an objective conclusion. Stated differently, the argument structure is the mechanism that allows us to answer the question, [why] by isolating all of the relevant factors that contribute to the reason for variation in the thing we are examining.

The principle objective of logic rests in how the truth of independent premises, combine to support a particular conclusion. Essentially, arguments can be thought of as nothing more than a series of premises (or statements of truth) that lead to, and support, a specific conclusion. Arguments can be expressed verbally, in writing, or in an equation, but essentially the goal is the same and that is to identify those factors that contribute to the outcome of some specific area of interest in order to answer the question, why. As you will discover later in this book, arguments serve as the foundation for all reasoning and act as the building blocks for human understanding. The evolution of knowledge also depends, almost entirely, on the structure of our arguments, which disclose the discoveries made by previous generations and then help us to combine those truths with contemporary knowledge in order to form a greater level of understanding.

Although it sounds simple enough in theory, in practice it can be very a daunting task to fashion an effective argument that provides indisputable specificity of the premise and which also affords irrefutable accuracy of the conclusion. The reason for this (I believe) is that almost nothing in the world is univariate. Generally speaking, when faced with complex questions about the interaction of phenomena or while searching for an explanation about the cause and effect of things upon one another, the vast majority of people tend to dissect and interpret how the world is put together from a rather univariate perspective. The temptation to oversimplify things and to seek to

reduce a complex question to its simplest form is quite understandable really, due largely to the fact that contemplation of the multiple interrelationships that exists between variables is difficult to achieve, and as a consequence, most people naturally grab hold of the first reasonable explanation that occurs to them regarding how particular phenomena interact so that they can expediently articulate their conclusion.

The problem with this approach to problem solving is that people (once they have decided upon an explanation) tend to cling to their initial argument as though it were a reflection of their personal character, in spite of the introduction of new information that may either invalidate their assertion or better explain the situation. The natural byproduct of such an approach to problem solving (especially if challenged by another during a debate over the issue) is that the dialog typically degenerates into nothing more than a contest of wills, and the truth of the matter is never fully isolated by anyone. After all, it is hard to think up all of the possible reasons that something happens and then prioritize the potentially contributive factors into a coherent argument. It is extremely difficult for people to change who they are, how they think about things and seemingly even more difficult for us to withhold judgment about something until all of the possible alternatives have been examined. We all know that Who we are, our cognitive abilities to reason, the methods we employ to arrive at a particular conclusion, and the judgments we make about the world cannot possibly be

flawed, because that would mean that we are flawed, and this is simply not acceptable to us.

The most demonstrative difference between people who are trained in the scientific approach to problem solving and those practices employed by “normal people” is the ability of the former to recognize the innate complexities and interrelationships of the world and their conscious effort to employ a methodological structure to the problem solving process, which endeavors to assure that all potentially contributive factors are examined, prior to rendering a judgment. I think it is important to remember that Occam was wrong when he prescribed that all things being equal, the simplest explanation tends to be the right one. As you will recall, Occam also thought the world was flat. Yet, when people are at a loss to provide a specific (validated) explanation as to why something happens, they will occasionally invoke the concept of Occam’s razor, as though paraphrasing an ancient philosopher somehow lends credence to their position that a simple explanation is correct. It is probably not at all simple and it is probably not at all accurate.

We are well served to remember that there are a significant number of forces at work, at all times, exerting individual pressures and collective influence on the outcome of everything. Even for the most perceptively simplistic equation, the scientist must account for all the aggregated influences contributing to the outcome and withhold judgment until all the data are analyzed.

Arguments are the mechanism that we use to fashion this deconstructive process in order to isolate the variables responsible for exerting influence on the outcome. Arguments should specify the contentions and variables in our scientific equations and articulate the hypothesized relations that exist between the individual variables, as well as the eventual result.

An easy way to visualize such an argument structure can be seen below, where statements of truth (premises) are presented and ordered by perceived importance, culminating in support of the conclusion.

- **Premise 1** - The suspect, when arrested, was in possession of the gun that was used to kill the victim.
- **Premise 2** - Witnesses to the crime identified the suspect as the person who committed the act.
- **Premise 3** - Scientific tests indicated that the suspect had gunshot residue on his hands at the time of his arrest.
- **Premise 4** - The suspect was involved in a fight with the victim an hour prior to the shooting.
- **Premise 5** - Blood spatters of the victim were found on the clothing of the suspect at the time of his apprehension.
- **Conclusion** - The suspect killed the victim.

As you can discern, each factual premise in the example above is directly relative to the conclusion and the aggregate influence of all of the of the individual truths combine to support the overall conclusion that the suspect had motive, opportunity, and the means to commit the crime. Therefore, he is guilty of the crime. If only all criminal trials were this easy to prove, but you get the idea

that without a clear delineation of the premises, the conclusion is left to doubt. Remove anyone of these truths and the case gets weaker. Disprove any of these premises and the jury has a more difficult time arriving at a determination of guilt that is beyond reasonable doubt. This is precisely why we require unanimous consensus by a jury for criminal trials. If all of the jurors do not come to the exact same conclusion, then the accused is set free. This avoids the possibility of wrongful conviction based on a flaw in the logic of the prosecution's case and assures that not merely a preponderance of evidence is provided, but that the measure of "beyond reasonable doubt" applies. You'll be interested to know that this is not the case for civil trials. There, only a preponderance of evidence is needed for the juror to render a verdict. Even more fascinating is that civil trials do not require a unanimous verdict, which begs the question why not.

We could express an argument in a mathematical context as well, such as, $P^1 + P^2 + P^3 + P^4 + P^5 = C$. Although no quantitative values are assigned in this theoretical presentation, such a consideration sets the stage for hypothesis formulation for each individual variable in the equation. In scientific research this is precisely how we derive an equation that contains the independent variables (or premises) that we hypothesize may influence the dependent variable (i.e., the conclusion). In such efforts, we construct a method for quantifying the data, and then test each premise individually to assure the truth of the speculation, followed by measurement of the

individual and collective strength of all of the variables in affecting the value of the dependent variable.

You probably did not fully comprehend that explanation, but rest assured that by the time you read the entire book, you will have a better idea of how this process is accomplished. The point here is that there is no difference (structurally) in formulating an argument in support of a legal decision or for a scientific discovery. They are all based on an argument that presents a series of truths that individually (and collectively) have relevance to the conclusion and prove beyond reasonable doubt, the assertion offered in the conclusion.

Let me share with you an argument that I have used in my classes for over twenty years to profile how a person can construct a sequence of premises to support a conclusion. This is an intriguing argument that involves combining a series of apparently dissimilar premises into an argument in an effort to support the conclusion. The problem arises when you recognize that many of the assertions are not testable and therefore you cannot (as a scientist) prove the conclusion beyond reasonable doubt. Subsequently, you find yourself in a position of having to withhold a judgment about the argument, until such time as evidence is presented in support of the premise that removes all doubt about the truth of each assertion. It is an emotionally charged equation, which brings into play the perspective of the participant, as well as their imagination.

An Argument for the Existence of God

P1 - Flatlander's Perspective of 3 Dimensional Space applies equally to our ability to discern the truth about influences in our universe.

Plato's Flatland paradigm suggests that we are forced to formulate conclusions about the universe based on our limited perspective. A flatlander's description of a sphere (for example) passing through their 2D world would be based upon seeing the leading edge of the sphere only. All Flatlanders would describe the sphere as a line that grows, and then shrinks, as the two hemispheres pass through the flatland plane. Consequently, the conclusions that we draw about the universe may be totally incorrect, even though we employ logic and reason because of our limited perspective.

P2 – The Bible Code profiles specific events in human history through the incorporation of a process that is based on equidistant coding, the results of which are well beyond the results expected by statistical probability

Equidistant encoding of earthly events is repeatedly profiled in the Old Testament (Torah). Many historical and modern day events, the participants, the dates of occurrence, and the circumstances can be found within the Bible, in a (crossword puzzle) form of code structure. Depending upon the instance, the statistical odds of this phenomena occurring has been calculated well beyond probability (1:10million). Such codes do not occur with equal regularity when the same algorithm is applied to other texts of equal length. This begs the question, not that God would be clever enough to author a code, but more importantly, how could God know before it happens. This might be answered by suggesting that God is not bound by the same temporal limitations that we encounter.

P3 - Special Relativity and Quantum Mechanics confirm the relation between matter and energy, and reinforces the notion of multiple temporal dimensions

Albert Einstein, in his theories of relativity, postulated that matter and energy were forever phenomena of distinctively different realms of existence. Two distinct observations made by Einstein have direct relevance to this argument.

Law 1 $E = MC^2$. Matter can never attain the speed of light

Law 2.....Time Slows Down as we approach Light Speed

These two principles suggest that the realm of the physical universe and light (no matter its wavelength) are separated by a speed barrier (186,000 miles per second or the speed of light). To make this even more intriguing, Quantum Mechanics (an area of study in physics dealing with atomic and subatomic particles pioneered by Niels Bohr) suggests that multiple spatial dimensions and multiple time dimensions do exist within our universe but at the subatomic level. This multiplicity of time and space prescribed by QM, has been accepted as a necessary pretext for QM and validated within the mathematical proofs of quantum mechanics. String theory for example suggests that eleven additional spatial dimensions and six additional temporal dimensions exist.

P4 - The Shroud of Turin, not only accurately portrays the physical evidence of existence but also contains evidence of a conversion from matter to energy during the resurrection.

The alleged burial cloth of Christ (the Shroud of Turin) has been a matter of considerable debate. What is not in debate is that an image on the Shroud was found to represent a man that was scourged and crucified and when examined under VP8 analyzer, it reveals a 3D image. There are a variety of evidentiary proofs on this garment and the speculation is that the image was created by the transition of **Matter to Energy**, during the resurrection.

P5 – The Bacterial Flagellum is proof of Intelligent Design

In *The Origin of Species*, Darwin stated, “if it could be demonstrated that any complex organ existed which could not possibly have been formed by numerous, successive, slight modifications, my theory would absolutely break down”. A system, which meets Darwin’s criterion, is one which exhibits *irreducible complexity*. The Bacterial Flagellum, because of its simplicity of design, irreducible complexity, form, and function suggests an intelligent design and not chance or evolution. Several leading scientific authors (Behe) have changed their position about evolution based on this discovery.

Based on these five premises, the following two conclusions can be offered.

Conclusions =

C1 - *God is Light* (i.e., God exists as energy)

C2 - *God is Non-Temporal* (Time only appears linear to us)

Essentially, the argument contains five premises that present a series of explanations which (P1) delineate the potential limitations of our ability to understand because of limited perspective, (P2) articulate that specific events in human history are encoded in a biblical text using Equidistant coding seeing well into the future and beyond statistical probabilities, (P3) prescribes certain physics principles that confirm the difference between matter and energy to offer a recognition that time isn't necessarily constant depending upon how fast you're moving and as a result at the threshold of speed of light, all time is the same time, (P4) suggests that there is physical evidence for the transfer of matter to energy in the Shroud and that this may have transpired as a result of the resurrection that converted matter to energy, and (P5) presents an observation pertinent to a piece of physical evidence that cannot be explained and suggests that evolution (as prescribed by Darwin himself) does not account for the irreducible complexity of the bacterial flagellum and may hold evidence of Intelligent Design.

The conclusions presented take these five premises into account and provide a plausible explanation that God is

light (or energy) and because God is Light, God is not bound by temporal limitations, because God exists within the realm of energy or is pure energy. A number of additional premises could be added to this argument to support such a contention, but I think you get the idea.

The challenge here, as a scientist and critical thinker, is to recognize that although we cannot accept the conclusions presented in the argument because no conclusive proof has been provided that can be scientifically confirmed the conclusion, we are best advised to withhold judgment because the evidence presented clearly doesn't meet the conditions of being beyond a reasonable doubt. It is certainly a compelling argument however, and one which evokes that emotional intelligence quotient that I mentioned earlier in the book. The argument is without question an interesting approach at explaining the basic logic of such a proposition. The premises used here differ considerably from the traditional biblical arguments that are based on first hand observation of miracles, but which are unsupported by evidence. The Bible and religious dogma encourage belief based on faith. As scientists, we are bound by a different covenant and must insist on proof. This argument brings forth a series of assertions that prompt reflection about the plausibility of the conclusion. This argument calls to the reader's attention specific principles and assertions that may explain what God is and how God knows what is going to happen before it happens, so that it could be encoded into the Torah. However, without absolute proof of each premise and an

explanation of their relevance to the conclusions, we have to withhold judgment until such proof is provided. This does not mean we need to turn away from such attempts to explain beliefs or climb upon our scientific high horse in judgment of others. Rather we should embrace all prudent attempts at explaining the unknown and use these types of arguments to further our skill at evaluating logical propositions. We could also use the framework of such an argument to develop scientific tests to prove or disprove the hypotheses proclaimed.

One interesting test (that has direct relevance to this argument) and that can be scientifically tested involves a sixth premise to the aforementioned argument that asserts that when two photons are fired into a chamber with only one exit, one of the photons will exit ahead of the other. If light speed is constant at 186,000 miles per second with no acceleration or deceleration, then the photons should either collide or exit at the same time. The inference here (which is chronicled in the book *God at the Speed of Light*, T. Lee Baumann, 2002), suggests that light has intelligence and infers a consciousness so as to avoid a collision. Personally, I find the experiment and inference of considerable interest. The point here is that this is a scientifically testable assertion where an analysis can be conducted to isolate the variables and test for the accuracy of the claim. If it does not prove out, then you have dispelled the premise, but on the other hand, if it does prove to be true then we all need to consider the merit of the premise and search for the answer to the

question [why]. Is there intelligence at work, or does some other principle of physics apply that we do not yet know about. How does Plato's flatland apply to such a scenario? Are we looking at a sphere and because of our limited perspective, all we perceive is the leading edge? Inevitably, as scientists, we must defer judgment and wait until absolute proof is offered that can be scientifically confirmed before we manifest a conclusion. That is just the nature of who we are and how we think about the world, but isn't it an interesting and curious argument?

The final example that I will use to explain argument structures is predicated on a decision process that eventually becomes extremely important to most of us at some point in our lives, but which few of us rarely employ scientific reasoning to discern the right answer. This example involves the decision (or conclusion) about selecting the perfect spouse. One would think that this one decision (above all others) would be guided by our best efforts to make the right choice, but alas, it is not approached as a scientific equation by the vast majority of us. Instead, it is an emotionally charged decision where our EIQ comes into play and more often than not (as evidenced by the divorce rate) we fail to make the correct choice or even look at the relevant variables that might affect the outcome. Instead, we think with our heart instead of our mind, which is never a good idea.

This is an excellent example (I think), because we need to decide (what) constitutes "perfect", and how it is that we

quantify a measure for such an analysis. There are clearly a good number of possible alternatives, but let us assume that we choose whether the marriage ends in divorce as the ultimate empirical test of perfection. Certainly, other measures could be used, but for purpose of explanation, divorce should suffice.

When I use this example in a freshman college class, I normally start by seating the girls on one side of lecture hall and the boys on the other. It is a bit theatrical but it lends itself well to polarizing the group so there is a lesser probability of contaminating the experiment. Then, once that has been achieved, I asked the question. Okay ladies list for me the top three qualities of the perfect spouse. You can probably imagine the fervor that this question evokes and the qualities yelled forth in response from the crowd. Try it yourself before you read any further. What are your top three qualities?

After the ladies have spoken their mind, the men are asked the same question. Once again, we typically see an interesting litany of responses. Remember, these are college freshman so there's not allot of critical thinking going on in the crowd. From the vantage point of the women (who were intentionally placed in a segregated and protected grouping, variables such as money, physique, and loyalty are the first variables to be expressed. From the men's vantage point qualities such as culinary ability, physical features of the women and

submissiveness are often valued highest on the list. Did I mention that these were college freshmen?

What is interesting to note in this exercise in reasoning is that the decision about selecting the perfect spouse is not an easy one. There are, in fact, many critical level variables and qualities that should be assessed prior to making the final decision. Culinary skills, money, and physical features are all of value, but there are many more that significantly contribute to whether such a relationship would be perceived as perfect (as measured by whether the union ends in divorce) but which aren't often considered.

After the two groups have taken a deep breath from venting their hostility towards one another, I begin to point out those factors not expressed by the crowd but which are germane to the decision. They include;

1. Parenting Skill
2. Fidelity
3. Intelligence
4. Sense of Humor
5. Religious Beliefs
6. Social Status
7. Ethnicity
8. Responsibility
9. Future Promise
10. Judgment
11. Emotional Stability
12. Personal Habits

The list goes on and on, but as you can see, these factors, and the determination as to whether the qualities meet the minimum standards for acceptability between prospective mates, harbor a significant degree of influence in deciding the eventual outcome of the union.

We can actually construct a scientific experiment to test the individual and collective influence of each of these variables. The survey would be based on an instrument that elicits responses from two groups of people. One group constituting those who had experienced a divorce and the other group consisting of people who did not divorce their mate. Using scientific methods, we would use a statistical measure (which will be discussed later in the book) to prove the premise whether each variable possessed a statistically significant difference in determining group association. In other words, do the two groups of people have a distinctly different view as to whether their mate possessed each of the qualities hypothesized in the argument? Predicated on the results, we could then interpret the importance of each factor, as perceived by the respondents, in determining whether they believed that the quality was of importance to their decision to remain married to their spouse.

We will revisit this example later in the book to explain vertical and perpendicular logic, but I think you get the idea that even something as perceptively non-scientific as selecting a spouse depends greatly upon a significant number of (typically) qualitative variables that can be

measured and considered in the decision process. Within the initial paragraph of this chapter I alluded to the fact that an argument is not simply a disagreement between two people over an issue, but is also (within the educated community) a formalized structure that's used to articulate the premises used in support of a conclusion.

Now that you are familiar with the concept of argument structures, I want to return to the former definition of an argument at this point (that it is also a disagreement) to call your attention to the fact that arguments (or disagreements if you prefer) are an important part of the exchange of ideas because of the sharing of opposing views about a particular issue and the justifications for seeing things differently. It is within these "arguments" where people have an opportunity to express the rationale behind their viewpoints on the matter. It is important also to remember that these disagreements are the perfect medium for eliciting (from the person that is expressing their viewpoint), those premises they have used to arrive at their conclusion and gauge the merits of their contentions as they apply to the conclusion.

We should be careful never to attack or demonize the person making the claim, but we have an inherent responsibility to critique the truth of their premises, as well as the relevance of such assertions toward justifying the conclusion they have presented. We may just find, through this exchange of ideas, that they offer a variable that we have previously overlooked or that they possess a

slightly different interpretation of the data that may affect our ultimate judgment regarding the issue.

In such instances, whether they involve debates over social, political, or scientific issues we can employ the techniques of listing each premise, evaluating its accuracy and its relevance, and then assess the collective merits of the premises expressed to the eventual conclusion. By decomposing someone's argument (i.e., $P1 + P2 + P3 = C$), the merits and accuracy of their assertion are more easily recognized and subsequently, their assertions or claims can be more accurately evaluated.

The process of argument decomposition is a particularly effective tool in getting at the truth of each premise and then, in turn, assessing the relevance of each individual premise to the conclusion. Without such a process, it is difficult to discern the point of view being expressed and even more difficult to accurately gauge the relevance of the assertions being made.

Chapter 3

Introduction to Multivariate Reasoning

At this point in the book, it is time to translate all the prior (general) information provided into a practical framework, so that you can more effectively grasp the processes and methodologies, which support multivariate reasoning.

Over the years, it has been my experience that a striking commonality exists relative to many of the theoretical postulates and explanations provided for a wide range of academic disciplines used to explain why things happen. No matter which discipline you examine you will find that a significant number of the assertions offered to explain such things as human behavior, politics, economics, and even some matters relative to science fail to articulate all of the variability required to account for the value of the dependent variable. Expressed differently, a significant number of the authors of these contributions have put forth interesting notions (or theories) but which simply turn out to be nothing more than an unsubstantiated opinion about why things happen or why people behave in a certain manner. A significant number of authors never go to the trouble of conducting empirical studies to prove the truth of their speculations. Some of the reasons for this may be attributed to the scholar's lack of proficiency with empirical forms of analysis and scientific reasoning, or

perhaps it is predicated on their reliance on purely intuitive methods of analysis. No matter the reason, the consequence is the same in that the theories they prescribe (no matter how meritorious) often fail to succinctly account for the totality of influences of the factors involved, the relationship between the variables, or to offer conclusive evidence of their postulates.

I could write for hours about the shortcomings of primary, secondary, and collegiate education in adequately preparing students to think scientifically or to recognize the complexities of the universe. Suffice it to say, that these institutions (often times) do not provide adequate coverage of this most important skill. Students should never accept (at face value) the truth of a postulate put forth in any textbook, simply because it is in writing. Rather, they should question (excessively) the merits of the arguments and theories prescribed in these texts and demand empirical proof of the accuracy and relevance of such assertions. Students should also subscribe to a protocol in order to help them to discern the accuracy and truthfulness of any postulate, verbal or in writing. A protocol that relies on multivariate theory seems most appropriate given the complexity of the issues we encounter most often in our search.

In the subsequent chapters we will address (specifically) how scientific research designs are structured in order to assure correctness of the approach, pertinence of the data to the analysis, and relevance of the research structure to

actually [proving or disproving] the premises and conclusions of an argument. For now however, it is important that we concentrate on expanding your awareness and familiarity with multivariate reasoning principles and the structures associated with such deliberations. Toward this objective, I should relate that multivariate reasoning can be effectively defined as the examination of the interrelationships that exist between factors, in order to determine their effect upon one another. This form of analysis is commonly used to assess both the influence of a single factor upon another, or it can be used to assess the aggregate influence of multiple variables upon an isolated [dependent] variable. It is also important to point out that multivariate analysis is that mechanism in the scientific process where the truth or fallacy of an argument is tested.

Put more simply, multivariate theory suggests that, at any given moment, there are a considerable number of factors that combine to influence and alter the state or condition of the [dependent] variable. This dependent variable can also be thought of as the conclusion within an argument. By determining which variables (or premises) most strongly contribute to changes in the frequency of the dependent variable (i.e., the conclusion), the researcher is positioned to make judgments about the relationships that exist between these factors, and also which specific factors contribute to the outcome most influentially. Essentially, through multivariate reasoning and analysis, we are describing the evaluation criteria that will be used

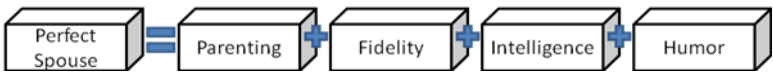
to evaluate the phenomenon. From this knowledge, judgments can also be made of how best to control and manage the fluctuations in the dependent variable. In other words, if the premises are represented as the independent variables in the equation and the dependent variable is the conclusion, multivariate analysis is the process where the individual and combined influences of these [independent] factors can be measured to discern their singular and collective impact.

Earlier I spoke to the importance of the decision maker stepping back and seeing the entire board. Multivariate analysis helps facilitate this objective. It is a process of contemplation, where all of the factors that could possibly affect the outcome are visualized and considered, as opposed to the ridiculous effort of trying to explain something complex (like criminal propensity or political disposition) based on a single theory.

From a multivariate deliberation, a decomposition diagram of the logic of an argument can be sketched out that specifies hypothesized interrelations for the multiple variables and factors involved in any phenomena. There are several steps in the process, but the end goals are to (1) visualize all of the possible influences ahead of the analysis, (2) to formulate hypotheses [i.e., premises] that support the inclusion of each factor within the equation, (3) which is followed by testing of each premise to discern their relative degree of influence. Once the truth of each individual premise is tested and confirmed, the final step is

to discern the proportional influence of each factor in the aggregate equation and then derive a conclusion.

Multivariate Logic Equation



The diagram above illustrates a classic multivariate design using some of the variables that we spoke of earlier that represent factors that potentially contribute to the selection of a perfect spouse. I would have included all of them, but it is not practical on a small sheet of paper.

As you can see from this example, using this multivariate approach, you immediately get an appreciation for the importance of “visualization” to the process of seeing the entire board. Through such a process it is a relatively simple course of action to contemplate and identify the dependent variable [or conclusion], as well as those factors that could [potentially] be exerting substantial influence. Through this visualization process, you are empowered to take a step back, contemplate the problem and establish a set of criteria that can be used to support the analysis. This would naturally involve the deliberation over a number of factors, each [potentially] exerting a degree of individual influence, as well as the combination of these factors in asserting a collective influence over the

outcome. As a side note I should relate that none of us should enter these types of analytical process with an objective in mind to prove one thought over another. By nature and design these are exploratory forms of analyses and as such we try to fervently avoid making prejudicial speculations about how the results are likely to turn out. Essentially we are engage in a venture that relies on the process of elimination to accept those premises that might possess influence and reject others that appear not to exert influence over the outcome. Stepping back and theorizing as to all of the possible factors that could account for changes in the outcome and keeping an open mind so that we avoid prejudicing our interpretations is a key part of the process. Multivariate theory facilitates our examination of all of the possible factors and helps us to remain objective, because we avoid the pitfall of favoring one explanation over another.

From the example provided (the perfect spouse) you can see that there are four hypothesized independent factors laid out horizontally and identified as (Fidelity-bX1, Intelligence-bX2, Age-bX3, and Promise-bX4) that are suspected to partially contribute to the dependent variable (the selection of the perfect spouse), which is designated as Y' (or the predicted value). Clearly (as we discovered earlier) there are many more "independent" variables or factors associated with such an equation, and as I also mentioned previously, as researchers we have a responsibility not to rush to judgment until we have accounted for all the variability in such an analysis. Stated

differently, we must withhold judgment until we have confirmed the truth of each premise and then determined its proportional degree of influence within the argument, as well as identifying all of the factors that exert influence to account for the variability in the dependent variable.

This will become more obvious in later chapters as to how (mathematically) to make such a determination, but before proclaiming the virtues and relevance of our discovery, we need to make very certain that all of the factors have been examined, measured, and interpreted. This assures that we know the correct answer to the question [why] and, unlike those writings I referred to earlier that offer a single explanation for a complex behavior or phenomenon, we aren't simply espousing our narrow minded opinion and misleading future scholars who might make the mistake of believing our unsubstantiated claims. Failure to make certain that we've accounted for all of the variability, leaves our research open to criticism, our methods suspect, and our reputation and standing in the scientific community open to censure.

Based on the identification of the four suspected independent variables, we are positioned to formulate a hypothesis that addresses each variable (separately) followed by an aggregate hypothesis that potentially explains the collective influence of all the variables on the outcome (provided each factor proves to be influential). The reason for this is based on the fact that before we can

examine the relevance of the multivariate argument, we need to make sure that each premise is truthful. Only after we have assured that each individual premise is correct can we analyze the truth and relevance of the multivariate equation (or argument) to the conclusion. In those instances where an insignificant correlation is computed, we continue the search for meaningful factors. It is (essentially) a process of elimination where the search for relevant and substantive factors are continually evaluated until all of the factors that influence the dependent variable have been uncovered.

We will discuss, at length, the concepts and processes for proving a premise using quantitative and qualitative analysis methods in the upcoming chapters. It is important to recognize here however the stepwise sequence of articulating the suspected influences, followed by data collection and testing of each independent variable or premise in the argument for truth, and that these steps precede the final effort involving analysis of the argument to discern the relevance of the assertion. If all of these steps happen in sequence, the result will be a scientifically supported position that provides specificity and conclusive evidence. With such an approach, we will also know the relative influence of each variable in determining the value of the dependent variable, which is a pretty important fact to understand.

Variable Identification

In order to fully comprehend the concepts of multivariate theory, it is essential that you become fluent with the terminology used to support this endeavor. Traditionally, the first term that you will encounter is referred to as the **dependent variable**. This term is used to refer to phenomena which exist in the world, but because this particular variable has been selected as a focus of the examination (or conclusion in the argument) we presume that it is [dependent] upon other phenomena for fluctuations in its state of existence. All shifts in frequency, and in some cases, its very existence, are presumed to rely upon the presence and influence of other “contributive” factors.

These contributive factors are classified as **independent variables** and are presumed within the hypotheses to influence or contribute to the dependent variable as it shifts, shrinks, grows, or winks out of existence altogether. In other words, these independent variables serve as the premises of our argument. Although they act independently, as they shift or morph, their individual variance is presumed to contribute to the dependent variable and fluctuations that it experiences, because of the relationship that is shared between them.

Empirically oriented scientists develop what are termed **research and null hypotheses** based upon a review of pertinent literature about these factors. By examining the

previous works of other scientists relative to the subject and through their own observations of the phenomena under study, they develop a theoretical postulate about these variables and any relationships which may exist between them. Subsequently, they develop research and null hypotheses regarding these factors.

The **research hypothesis** is always a positive statement about the potential relationship between the independent and dependent variables, while the **null hypothesis** represents a negative supposition and proposes that no relationship exists, whatsoever. Modern empirical methodologies reinforce this skeptical view, until proof is offered (by means of statistical verification) that a relationship does, in fact, exist. Accordingly, the scientist ALWAYS accepts the null hypothesis as their point of departure until they are offered proof to confirm that the relationship is not just suspected, but statistically verified and could not have happened by mere chance. As applied to argument construction, this approach assures that each individual premise is true before proceeding to the analysis of the collective logic of the argument.

Prior to measuring correlations between variables, we first need to test the truth of the premise which implies that a specific factor matters. To accomplish this, we structure a scientific test to check whether a statistically significant difference exists between the means or frequencies of two groups/categories relative to the factor or variable. The book will cover this process in depth in the next four

chapters, but proving the truth of the premise that a particular factor might be important in an equation rests at the heart of the argument formulation process. If we incorrectly assume that a particular factor matters, but fail to prove this supposition prior to testing its strength of influence using correlation analysis, then we stand a good chance of contaminating our equation with a spurious variable that, although mathematically correlates, possesses no theoretical association. To avoid this, it is always advisable to test the truth of the premise by using Z ratio, Student's T ratio, Chi-Square analysis or some other test of significance assure that differences between the means or frequencies exist. Such a preliminary check will lessen (not eliminate) the possibility that a spurious variable was included within the final correlation equation (or argument).

After proving the truth of each independent variable (or premise) we can use a powerful form of statistical analysis to prove the hypothetical relationships between the premise and conclusion. This process involves the use of the Pearson's Product Moment Correlation Coefficient (r). **Pearson's r** is used to measure the comparative movement between the X (independent), and Y (dependent) variables, the deviation values, the squared deviation values between the variables, and ultimately the sum of the squares. By dividing the sum of the squares by the number of observations in the sample, multiplied by the standard deviations for the X and Y variables, the degree of association between the variables can be

determined. In turn, by comparing the calculated value for (r) against a standard table of relative strength (.00 to 1.0), the degree of relationship maintained by the independent variable compared to the dependent variable can be measured. A Pearson's r of .80, for example, would indicate a relatively strong degree of relationship between the independent and dependent variables. This would indicate that as the frequency of X rises or falls, the frequency of Y rises and fall correspondingly a significant percentage of the time. Don't worry if you did not understand all of that. I will explain it in English later, but I felt compelled to show off a little.

After you have computed the correlation coefficient that exists between two variables, you are positioned to take the next step, which relates to determining how much of the variability, within the dependent variable, can be attributed to the change in the independent factor. To calculate the percentage of time that the variables fluctuate together, the **Coefficient of Determination** is computed. This is accomplished simply by squaring the value of a Pearson's r (let us say .80) to derive the coefficient of determination. To calculate the percent of time that the independent and dependent variables move in unison, simply multiply the coefficient of determination by one hundred. If the correlation coefficient is $r = .80$, then by squaring that value ($.80 \times .80$) and then multiplying the product by one hundred ($.64 \times 100$), you can determine that the two variables move in unison (both upward and downward) approximately 64 percent of the

time. This conversely means that 36 percent of the time (or $k = 1 - r^2 \times 100$) the variables do not move up and down in unison. This **Coefficient of Non-Determination (k)** indicates that some other factor/s is contributive to determining the unexplained movement observed within the dependent variable. In other words, something else has an influence as well and you need to find out what that is, and how (theoretically) it relates to the outcome before proclaiming that you know for certain the answer to the question [why]. We will go over this process in detail later in the book so you get more practice and have the ability to master the concepts.

To recap the process, multivariate theory requires the researcher to develop a well-defined hypothesis, which is supported by a sound theoretical premise, and then collect information and data according to a representative and unbiased sampling strategy. To accomplish this you first start out by visualizing the whole board and laying out potential explanations for the conclusion, which are expressed as independent variables (or premises).

Normally a random sampling process is devised to collect information pertinent to the study and which is representative of the population. After a sample is acquired, the data is examined using univariate tests to assure the truth of each premise and then a correlation analysis to determine if the associations expected in the hypotheses are relevant. If the correlation coefficients are strong enough and the level of significance is calculated

beyond the .05 or .01 (95% or 99% levels), then a measure of association (coefficient of determination and non-determination) pertinent to the observed relationships of the sample can be surmised and inferred back to the total population. This process is called inferential analysis for a reason. That reason rests in the fact that most of the time we use sampling to test our theories, because it is not physically possible to evaluate data for an entire population. As a result, we are extremely cautious in using established standards for determining statistical significance and only after we conclude that the difference or correlation is beyond the .05 or .01 levels, are we safe in “inferring” that the results we received in the sample “probably” apply to the total population.

Expressed differently, inferential analysis is used to test for significance as represented by a sampling of the total population and then the results are inferred back to the total population, provided the statistical odds indicate that it is safe to do so. What most researchers never fully understand is that inferential analysis is not necessary if you are examining the entire population. In such a case, you can reply simply on “descriptive” forms of analysis because you are already looking at the total population, so why infer anything. This does not mean that correlation analysis and tests for significant differences between means or frequencies is not important in determining relationships, only that there is no need to infer the results to the total population, because you are already looking at the total population.

As I have mentioned previously, the natural temptation when conducting such research is to oversimplify the number of potentially contributive variables in the equation and to treat all independent variables as primary sources of influence upon the dependent variable. In fact, some variables may turn out to have a second-order effect upon the dependent variable. In other words, they affect a variable that actually does have a primary influence, but which may not be included in the model. Since there are few researchers that fully understand this notion and even fewer quantitative tools necessary to make this distinction, it is difficult to differentiate primary from secondary influences. We will talk more about this in later chapters, but suffice it to say some variables have a direct effect, while other variables have a secondary influence, and which themselves are influenced by tertiary factors. Distinguishing primary variables from secondary or even tertiary variables should become an integral part of the hypothesis development phase of the study, not the methodological phase. We will address this in specific in the chapters dealing with vertical and perpendicular logic structures, but for now it is important to recognize that not everything has a direct relation on the outcome.

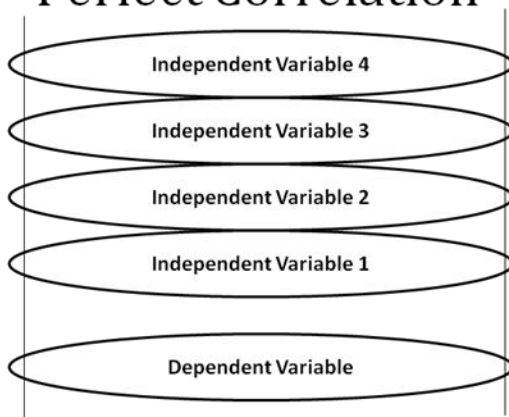
Additionally, there may exist, temporal adjustment factors, which are necessary to include within such models. A common mistake made by a significant number of researchers is that they forget to consider that, a lag adjustment might be required to the dataset prior to the actual computations in order to accommodate for delay in

the influence of one variable upon another. Let me provide an example. If you were to theorize that toxicity affects fish populations and then formulate a research hypothesis that stated “there is a statistically significant difference in the mean fish populations between lakes where toxicity levels were above a specific level and those below that same level”, then the temptation might be to rush off and collect data for these two factors and compute your statistical test. If your data collection method reflected the measure of toxicity for Day 1 alongside the fish population for the same day, your research design would not take into consideration a very germane consideration, which is, that toxicity may not be immediately threatening to fish (depending upon the type of toxin). To account for a delayed effect, the researcher would need to lag the raw data so that it represents and provides for the examination of near-term and longer-term consequences. We will talk more about this later, but I thought it might help clarify the importance of assuring that your data collection and processing methods needs to **exactly** match the possibilities that exist in the real world in order to avoid making a mistake that leads to an erroneous conclusion.

An Academic Example

In the diagram below you can see a hypothesized multivariate example that involves a theoretically, spatially, and temporally perfect correlation between four independent variables and one dependent variable. I should first point out that this will probably never happen in the real world, because God did not create such a perfect world. I use this example merely to point out that the independent variables occupy the same space, at the same time, as the dependent variable and that a degree of theoretical relation is presumed between all of the factors.

Theoretically and Spatially Perfect Correlation



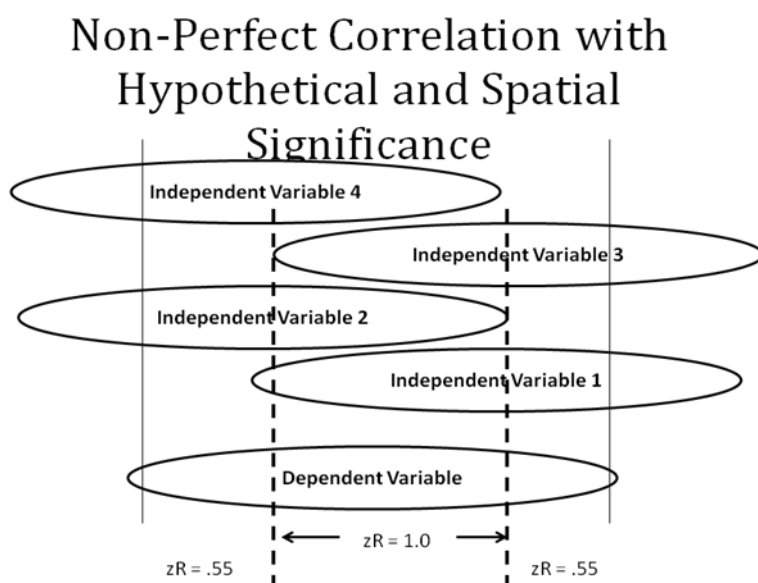
Given that sufficient theoretical association between these phenomena exists, you would conclude from this diagram, that you have isolated four potentially contributive factors that possess a strong correlation to the dependent

variable, and that there is sufficient spatial overlap and temporal congruity to conclude that this could not have happened by mere chance.

You could also use the previous visualization (the sequence of boxes) method to illustrate the argument but in this case, we are referring to natural phenomena and measurement of the quantitative factors should also be accompanied by assessment of the geographic proximity of each independent variable to the dependent variable prior to rendering a decision. In such cases, the researcher can employ geographic information systems (GIS) along with standard statistical analysis methods to assure a greater level of assuredness. GIS has become a very powerful tool to support scientific analysis. When you combine standard statistical analysis that is used to identify the hypothesized independent variables and test for the truth of the each premise along with the correlation between all factors, followed by an analysis of the spatial and temporal considerations, your analysis will be about as comprehensive as possible.

We cannot forsake the importance of good science simply because we have access to advanced tools like GIS, but combining both statistical analysis and spatial analysis into our protocol gives us an added advantage in discerning the theoretical logic of our argument, as well as assuring that the spatial and temporal criteria meet the standards needed for acceptance.

The previous diagram showed a theoretically, spatially, and temporally perfect relation between variables, but what is more likely to occur is the pattern illustrated in this next diagram. As you can see, some degree of spatial overlap exists between the independent variables and the dependent variable, but it is not at all a perfect correlation. Although they may be hypothetically related, seldom would the data suggest that the relationship is mathematically perfect, nor would you expect it to be spatially perfect.



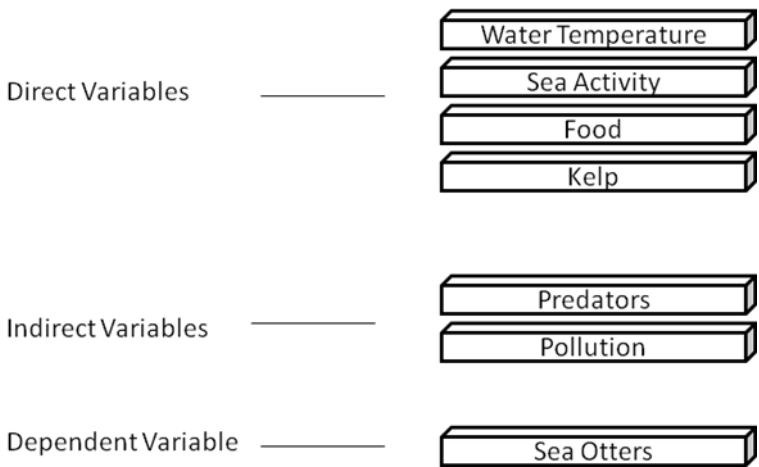
As you can see from the overlap, there are areas where all of the independent variables line up with the dependent variable and it is within this vertical region, where a perfect spatial correlation is achieved. Such a result would

indicate that it is this specific spatial region where conditions are most optimal to support the dependent variable. Additionally, if these IndVars fell into the category of “controllable” factors, then they could possibly be manipulated to alter the value of the dependent variable. Such a situation would provide a good degree of utility if several of the independent variables used in the model were directly controllable. Such control would provide a mechanism for manipulation of the independent variables so that a desired shift in frequency can be realized within the dependent factor. This ability to control your environment and manage resources and physical systems is the ultimate goal of such efforts. Simply understanding the interdependence of phenomena lends itself to the realization that certain variables are directly controllable by people and subsequently, if we know that changes in Independent Variable #1 will cause a desired increase in the frequency of the dependent variable, then we can proactively assure that we maintain the desired levels. The problems caused by such manipulation rest in the fact that we rarely clever enough to figure out that changes to one IndVar will not only have the desired effect on the DepVar, but will have second and third order consequences to some other physical phenomena. If you recall the effect that DDT had on the Red Tailed Hawk you will understand that killing bugs may be the desired first order effect, but the bug frequency was just a small portion of a greater and more complex model. By increasing the parts-per-million of this pesticide, we changed the balance of nature, crossed over into other

multivariate models, and nearly wiped out the entire species of hawks. Care must be taken to insure that we do not ignore such complexities and manipulate the wrong independent variables.

A Real World Example (of sorts)

This next graphic reflects a more real world oriented example of a multivariate model. In this case, the hypothesis would suggest that sea otter frequency is determined by the presence of P1-kelp beds, P2=the existence of an adequate food supply, P3-a relatively calm mean sea activity level, and P4-a water temperature range that is tolerable by the otter colony. ($P1 + P2 = P3 + p4 = C$)



These variables serve as the direct IndVars and without the existence of these factors within a tolerable range and spatial proximity, the likelihood of finding a sea otter colony is remote. Also illustrated however are two additional variables, which serve both as IndVars, but also as Chaos Inducers. Pollution levels can be hypothesized as having a substantial influence over the well-being of otters, but the typical state of the pollution level is well below the intolerable range of the otters. If this variable dramatically increases to a point where it threatens the otter colony, then over a measureable period of time, the members of the sea otter colony will die off.

The term *chaos inducer* also means that this single variable may result in a dramatic or disproportionate level of influence upon the other IndVars and potentially damage the entire ecosystem. I have included the IndVar described as Predator Population within the chaos inducer layers to signify that although normally within tolerable limits, a substantial shift in this variable could have a correspondingly negative effect upon the well-being of the colony. It would not necessarily affect the other variables in the equation, but an increase in predator frequency for example might dramatically reduce the number of colony members over a brief period of time. Recovery time from an increase in this IndVar might be substantially less than for the pollution variable, but near-term impact might be devastating. You will notice again in this illustration that the layers used in the model lay spatially perfect upon one another. This again, would probably not happen accept in

this academic example. The real world is much less predictable, so you would probably notice a substantial degree of overlap, but not nearly as perfect as shown here.

What you could anticipate to see in a properly constructed GIS representation of this type of study, is a high degree of overlap between the kelp bed and sea activity variables. The water temperature variable would surround all of the variables used and the food supply layer would extend well beyond the kelp and otter observation range layers. Where you would expect to find the highest concentration of sea otters would be where these IndVars overlap. Within this range, the conditions are ideal and would serve to support and promote otter populations.

Summation

Remember that the ultimate goal of this multivariate analysis procedure is not simply to observe the correlations between these phenomena, but based on this recognition, to identify controllable variables, which can serve to facilitate the effective management of environmental resources. If your focus is not on proactive control, then you must at least conclude that such knowledge facilitates an increased ability to direct disaster recovery efforts. Before you can employ either process however (control or recovery), you must understand the relationships that exist and then determine those things that can be manipulated by humankind.

Chapter 4

Vertical and Perpendicular Logic

As mentioned previously, perhaps the most demonstrative difference between people who are trained in the scientific approach to problem solving and those practices employed by "normal people" is the ability of the former to recognize the innate complexities and interrelationships of the world and to employ a methodological structure to the problem solving effort which roots out the influential factors of any situation under study (or at least we scientists like to think so). The byproduct of these analyses are the development of an awareness and the construction of quantitative formulas that can be employed to describe and exploit the relationships discovered and which can extend the analysis to the formulation of strategic policies that are based on this understanding and that hopefully provide effective control over the outcome of efforts to manipulate the environment in which we live.

This technique is not hard to master, but it does require practice and self-discipline. The pace of today's world exacerbates the temptation to cling to the univariate model of problem solving because decisions must be made quickly. However, those people who are most effective at policy formulation, I believe, tend to realize that decisions about how forces interact, what actions would be most

effective at achieving the desired results, and the consequences of such actions, do so from an informed perspective rather than a quickly acquired univariate determination. This means that they do not make decisions in haste. Rather, they take the time to examine the complexities of the issues under study and render decisions based on their assessments of what is best for all concerned based on the volumes of information that they have painstakingly assembled and analyzed (or at least that's the way in which it is supposed to happen).

As many researchers have discovered, using computer programs such as statistical analysis software or spatial analysis systems can greatly aid in the development of complex analytical models, but it all rests on the researcher's ability to visualize the equation. The key to using multidirectional analysis to facilitate such scientific examinations rests with the awareness (on the part of the researcher) that things are not as simple as they may appear to be at first glance. It is normally the combined effect of multiple factors (in multiple directions) that is responsible for fluctuations in the observed behavior of a dependent variable or resultant conclusion.

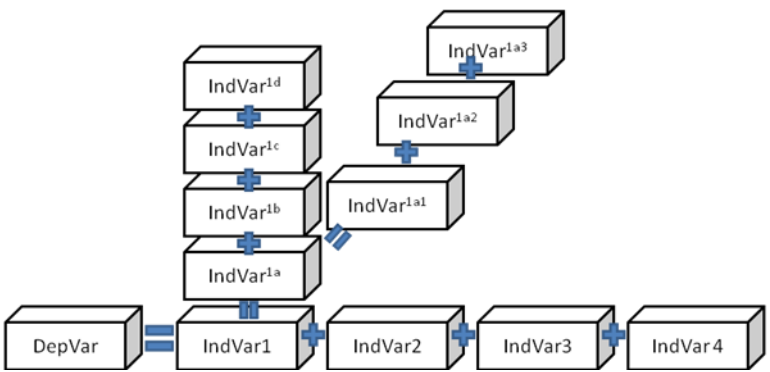
Even those people with more advanced knowledge and skill about how to construct scientifically structured methodologies for problem solving, I find, tend to fall victim to the temptation of limiting their examination of phenomena to a very linear form of multiple correlation equation that often fails to fully illustrate the

multidirectional intricacies and associations that also exist between phenomena. Because of our training, we rely heavily on the standard quantitative design methodologies we learned in graduate school to help us in developing such models, but we frequently fail to extend those analytical processes along both the vertical and perpendicular axes in order to attain the most complete explanation possible as to why something happens. This again is certainly understandable because of the difficulty associated with building complex models and the time required to construct such predictive equations.

With the advent of spatial analysis systems it has become “somewhat” easier to construct advanced models that account for not only a linear form of correlation but also which address the multidirectional relationships and expanded hypothetical interrelations that exist. It still comes down to the researcher’s ability to envisage such potential complexities however, in order to properly build an analytical environment inside the GIS software that accounts for such events. The inherent design of these types of software, as applied to the analytical process, is extremely conducive to replicating the scientific approach to problem solving. Within these systems, data are stored into separate relational tables according to some logical design structure and represent information in a manner that considers the value of these attributes. Some values will be purely quantitative, while others can be represented using qualitative values. No matter the strategy used to assign values to these data, the analytical

process must contend with not only the horizontal logic, but also the vertical and perpendicular connections. When using such systems, the identification of these causative/influential factors must be achieved prior to rendering a judgment if the researcher has any hope of isolating those variables which are not only causative in nature, but which also can be controlled or manipulated to achieve the desired changes in the dependent variable. Below you see an example of how the multivariate equation used to theorize the relationships that exist between the dependent variable and multiple independent variables can be extended along both the vertical axis as well as the perpendicular axis to expand the equation. This isn't done simply to build a better looking model, but gives us the ability to factor in dependency along a variety of theoretical dimensions.

MultiDimensional Logic Equation



If you consider our example equation that used selection of the perfect spouse as the dependent variable and

fidelity, intelligence, age, and promise as the principle independent factors, you can see how each of these “independent” factors is also “dependent” in a more fully articulated logic model. Expressed differently, fidelity may be a quality that affects the outcome of whether or not someone is considered as a candidate, but a person’s level of fidelity also depends on a variety of other factors. These might include parental role model, religious beliefs, whether or not the person was a victim of infidelity themselves, the presence or absence of opportunity to engage in such behavior, the person’s sense of obligation, and a plethora of other potential factors. Each of these influential factors in turn become dependent themselves on even more factors. For example, if you begin by looking at fidelity’s role in spousal candidacy and then consider that fidelity might be influenced by the person’s sense of obligation to the partner, it’s not hard to imagine an entirely new axis of variables that could affect a change in sense of obligation. Obligation might be (theoretically) influenced by any number of variables including belief that the union is mutually beneficial, a role predicated upon a sense of support, financial bonding, equivalent belief values, and so on and so forth. As long as all of the primary (horizontal) variables, secondary (vertical) variables, and tertiary (perpendicular) variables remain constant in the equation, no substantive change is likely to occur in the relationship. A sudden change in one of the variables in either direction can have a cascading (domino) effect upon the behavior, perception, and final judgment

about whether the spouse still meets the condition of “perfect”.

Clearly this was a relatively easy academic example contrived to help you understand the concepts of multidirectional logic but I think you can see how consideration of subtle level variables in the vertical and perpendicular axes are important to identifying all of the potential contributive influences. Real world equations can become extremely complex and move out in multiple levels (as depicted in the next illustration) but the point here is that in order to fully comprehend the dynamics of the equations you are studying, it is important to think in multiple directions. If you simply follow the process of identifying major independent variables, followed by vertical level influences, and then perpendicular influences you will significantly enhance the predictive accuracy of your model in illustrating all of the potential factors that can exert influence over the outcome.

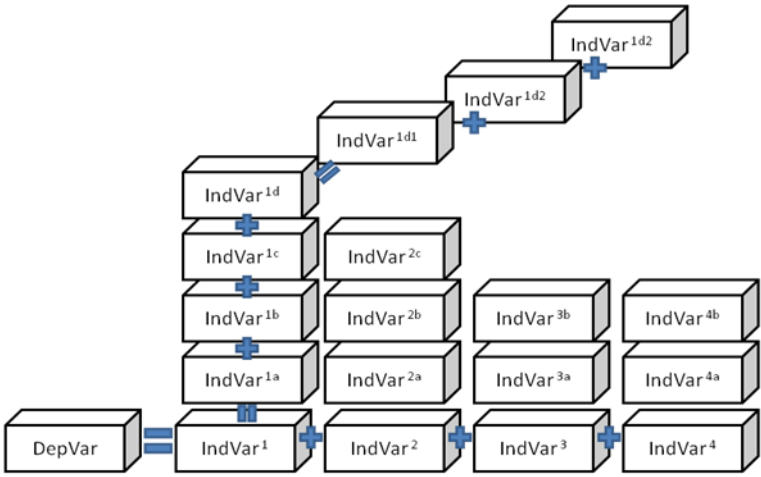
The research design strategy presented within the graphic below depicts just how complicated this process of multidirectional hypothesis formulation can become and also depicts the degree of comprehensiveness required in order for the researcher to build-in such multidimensional considerations (which probably explains why most of us never really build these types of behemoth equations in the first place). As can be observed, the relative spatial position of each variable, along with its defined hypothetical interrelation to all of the variables contained

within the multidirectional correlation matrix, is prescribed under such a design as well as the relative value of each of these variables.

From a practical perspective, in order for us to use spatial analysis systems as a mechanism to study such complex interrelations, it is imperative that we subscribe to the rules that govern interaction and relation, which are simply that (1) nothing exists separate and apart from everything else in the universe and (2) everything occupies space and time (even variables). In order for us to fully examine the complexities of the universe, we must build our computer models according to the same manner in which they really exist within that universe. It really doesn't matter whether you're building a model to explain a complex equation in physics, biology, natural resources, or the social sciences, the fact remains that everything exists within space and time and everything has either a dependent or independent status depending upon which way you look at it. This particular graphic shows the relational nature of things to one another as well as to and offers a unique look at how such interrelationships can be envisioned by researchers who are attempting to fully account for the interactivity of variables upon one another.

As presented, each independent variable contained within the traditional multiple correlation model (the lower/horizontal equation in the graphic) possesses a contributive influence over the value of the dependent

variable, but each of these IndVars is also subject to relative influences itself by other more subtle factors, and in essence becomes a dependent variable itself when considered in a more fully articulated model design. Those subtle variables contained in the perpendicular equations, in turn take on a dependent relationship as you consider even more perpendicular dimensions to the same equation. The most advantageous aspect of all of this however is not simply the infinite number of perpendicular relationships that may exist in such an equation (although that is entertaining to consider as well), but rather the degree of subtlety that one can achieve in identifying and controlling first, second, and third order variables.



MultiDimensional Logic Equation

After we have identified the perpendicular relationships that might exist between multiple correlation equations

and the interrelation that exists between perpendicular equations, we are free to further examine the controllable and non-controllable properties of each IndVar and use this recognition to experiment with the impact that subtle manipulations of each variable may have on the greater principle equation. This process of recognizing, identifying, measuring, and manipulating variables, in multiple directions, can be highly effective at disclosing the subtleties that exist between variables and at providing researchers with a tool for creating models and subsequent policies based on these models that maximize the degree of prerogative that exists in manipulating our universe.

From the strategist's perspective, such research designs not only offer a comprehensive examination of phenomena, but also further provide for the ability to construct advanced equations that offer a highly refined degree of subtlety over potential courses of action to manipulate the environment. This approach can not only maximize effectiveness in achieving the desired results to control phenomena, but can also provide for the expanded consideration that minor adjustments in the values of the IndVars contained in the perpendicular equations will have upon the more principle aspects of the model. Spatial analysis systems, in concert with the more traditional statistical analysis processes, makes it much easier for us to engage in this form of complex empirical modeling and afford researchers with a relatively inexpensive tool that can be used to produce extremely complex exploratory

and predictive models. The key to success in these endeavors is to simply remember that we cannot fall victim to the perils of oversimplification, nor can we forsake or ignore the principles of good science when engaged in such endeavors. In order to create the real world inside of a computer, it is necessary to pay attention to the rules under which it exists outside of the computer.

The world in which we live, is complex and in order for us to develop empirical models that accurately describe the interrelations of the world and the prospects that exist for securing positive change, we must build our equations with this complexity in mind. Accordingly, we must develop models that afford the ability to exploit the subtleties that exist within the universe and take advantage of our understanding of the subtle influence associated within our equations. In other words, Occam was wrong when he prescribed that the simplest explanation tends to be the correct one. The world is not simple at all. We simply choose (frequently) to look at it from a simplistic perspective which fails to account for the interrelations that exist and the subtle influences that contribute to virtually everything that exists within the universe.

Continuing with our examination of the importance of vertical and perpendicular variables within the multivariate logic process, I have included a most interesting real world example for your consideration. This particular application, named the Interactive Opacity

Neutralization Array project, was derived based on a research effort that I undertook some years ago to assess the viability of developing technology that would render military vehicles invisible to the human eye. As farfetched as this notion sounds, the results were extremely promising and the use of multidirectional analysis played an important role in determining the primary, secondary, and tertiary factors associated with the concept. The results of the study were presented to the Defense Advanced Research Projects Agency (DARPA) and then reviewed by the U.S. Air Research Laboratory. To properly articulate how perpendicular logic applies to this particular technology, it is important to provide a comprehensive description of the issues and considerations. The next several pages provide a detailed explanation of the concepts, considerations, and logic used in the project.

Using multivariate logic with vertical and perpendicular sequences to design a computer generated camouflage technology

The Interactive Opacity Neutralization Array

Today's complex modern battlefield incorporates a wide variety of complimentary strategies, tactics, and technologies that have been integrated to assist military forces in achieving an overwhelming level of supremacy during land, sea, and air combat actions. Recent confrontations have repeatedly proven the value of many

of these approaches and have disclosed that the most effective technology based systems used by conventional U.S. military forces during combat engagements have been those that not only involved technologies designed to provide an enhanced ability to detect, target, and deliver weapons with extreme accuracy, but also those technologies used to assure the survivability of military assets by making them harder to detect and subsequently more difficult to destroy.

Most of the technological advances made to date in the area of defensive systems have involved the concealment of military assets from detection by enemy forces and have been engineered to reduce, conceal, or mask the associated radar, acoustic, and thermal signatures of deployed military assets. Few, if any, advances however have been achieved in the area of reducing or eliminating the visual profile presented by military hardware within the visible light portion of the spectrum. The reason for this lack of progress in developing systems that can effectively reduce the visual signature of military equipment as presented within the visible light portion of the spectrum (400nm –700nm) rests primarily with the fact that it has been considered, by most, to be virtually impossible to render an object that has a definable mass and subsequent level of opacity, as visually transparent. What is not impossible however to achieve is to create the illusion that the object is transparent, no matter its size, or shape, or surroundings.

The answer to the question of how best to achieve this goal is believed to reside in the premise that transparency (as perceived by humans) is achieved when one can look through any object, to what is located on the opposite side of that same object, without detecting the shape, edges, and texture of the object, while it is within the foreground field of vision. It has long been the practice of those engaged in the development of traditional methods of camouflage to simply paint equipment or military clothing with a color that generally resembles the anticipated battle sphere and, through the use of traditional materials, to attempt to eliminate the enemy's ability to easily detect the edge, surface, texture, color, and contour of any object used for military purposes. Although moderately effective, the limitations of these conventional camouflage methods have been their ineffectiveness at achieving the goal of providing concealment within a dynamically changing combat environment.

The Interactive Opacity Neutralization Array concept was specifically conceived to address the issues of concealing military hardware and weapons in plain view and to explore the viability of using modern technologies, such as projected or embedded digital video arrays, to overcome the limitations of past camouflage methodologies. The system concept was also conceived to examine the viability of creating a new interactive system of concealment, which can be used by all branches of the military to assure asymmetric lethality by providing an

enhanced level of protection for mission critical assets during both daylight and night time hours.

In order to achieve the objective of developing a system that could potentially be incorporated by all military units to conceal the presence of ships, airplanes, tanks, helicopters, and even infantry units, the development of a functional interactive opacity neutralization array would need to concentrate on bypassing the impediments to achieving physical transparency and focus singularly on achieving the objective of creating “perceived transparency”. The solution that was been envisioned, and which embodies the concepts presented within this explanation, postulates that a technological means of achieving the result of presenting a transparent visible light signature is, in fact possible through an amalgamation of existing technologies in order to render objects imperceptible to the human eye and consequently reduce their exposure to potential enemy threats. The design criteria envisioned for an IONA system, would naturally incorporate the use of a variety of computer graphics technologies, so that all exposed surfaces of a combat vehicle would reflect imagery gathered from cameras mounted from six opposing perspectives that provide IONA equipped vehicles with the ability to camouflage themselves against enemy forces, no matter their position or visual perspective. Since an IONA system would need to provide simultaneous interactive video displays (taken from multiple directions) that are in turn presented to the opposing surface, concealment and camouflage would be

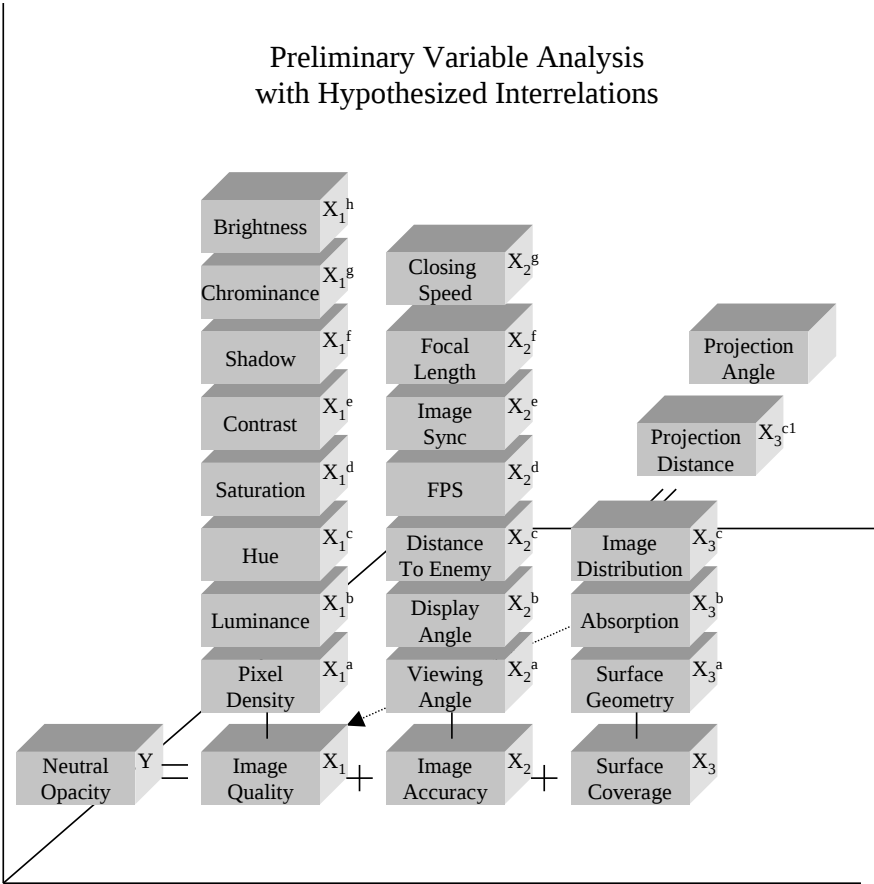
provided against all angles of attack. This capability is seen as demonstratively beneficial over conventional practices of camouflage in the fast paced battlefield of modern warfare.

During the initial stages of research and development for an IONA system it was surmised that it would be necessary to identify all of the pertinent variables in such an equation and to additionally discern their individual strength and contributive power, interrelation, and interdependence as related to achieving the presentation of an opaque visual signature. In order to attain an optimal level of opacity neutralization or to achieve the complete elimination of a vehicle's visual signature, it is also surmised that an examination of the relative variables must first be conducted, using a research design strategy similar to "discriminant function analysis", but where each major variable is not only examined from a linear perspective, but where second order variables are hypothesized and examined in order to determine their degree of subtle relevance to the overall equation. This process would allow for the development of perpendicular correlation equations that could provide an enhanced level of control, manipulation, and refinement to the opacity neutralization model. The results of such a scientific analysis could then be directly applied to the creation of a functional IONA system, using the various technologies identified during the research. A very preliminary form of such a research design can be seen in the graphic, where opacity neutralization (Y') is specified

as the overall dependent variable in the equation and image quality (X1), displayed image accuracy (x2), and vehicle surface coverage (X3) are initially designated as the principle discriminant variables. The main horizontal equation is then extended vertically to include all potential second order independent variables (as measured against the primary discriminant variables) for an expanded view of the interrelations that occur. This process is continued (multidimensionally) until all relative variables have been identified and correlated to the achievement of opacity neutralization. Each discriminant variable contained within the traditional linear segment of the model (the lower horizontal equation in the graphic) theoretically possesses a degree of contributive influence over the value of the dependent variable (in this case opacity neutralization), but each of these IndVars is also subject to relative influences by other more subtle factors, and in essence becomes a dependent variable itself, when considered in a more fully articulated model design.

Those subtle variables contained in the perpendicular equations, in turn take on their own dependent relationship as you consider even more perpendicular dimensions to the same equation. The advantage of using this scientific approach, along with traditional engineering methods, is not simply the number of possibilities that can be examined and applied to the perpendicular relationships that may exist in such an equation, but also the degree of subtlety that the engineering team can potentially achieve in identifying and controlling first,

second, and third order variables pertinent to deriving opacity neutralization under different combat scenarios.



The basic tenant of the IONA concept rests in the presumption that the visual signature of any combat vehicle, whether stationary or mobile, can be eliminated by reducing its level of detectable opacity through the use of synchronized digital video displays on its surface that mirror the terrain and objects, which are located on the opposing axis of the vehicle, as viewed by an enemy.

Essentially it is not the hypothesis itself, which suggests that if one can reduce the visual signature of any vehicle one can render it harder to see that is at the heart of such research, but rather it is the exploration of how best to achieve this objective which would need to be fully studied in order to design a fully operational opacity neutralization system.

With specific regard to the examination and testing of various technologies that may be contributive to the creation of a functional IONA prototype, preliminary analysis and investigation indicated that there presently exists a variety of commercially available technologies and resources which hold considerable promise for direct applicability and inclusion within such a system. It is anticipated that significant modifications would be required in order to adapt these technologies to create a fully functional IONA system that can withstand the rigors of the modern battlefield, but research and discussions with selected engineers (from multiple disciplines) suggest that the current state of the underlying technologies themselves hold considerable promise, as applied to such an endeavor.

The IONA, as presently conceived, could potentially be constructed using several different, but complimentary technologies. These include but are not limited to (1) the incorporation of a specially designed, tapered fiber optic array, which is affixed along the surface areas of any vehicle that can be used to display captured video images

received from a multi-axis display computer control panel, or (2) the use of a lateral image projection methodology that incorporates either white light laser technology or direct high intensity video projection to create an interactive display surface, and (3) the use of a non-reflective, flexible and non-flexible light-emitting diode panel technology that could either be used independently, or combined with other materials, to create a heat resistant laminate composite array that is capable of displaying video imagery throughout a vehicle's surface area. The LED configuration appears to hold the greatest degree of promise for unilateral system application.

In all of these configurations, digital cameras mounted on gyrostabilizers would need to be strategically positioned on, or within, the vehicle to capture real-time video images of battle sphere conditions. Depending upon the general direction or threat axis, estimated distance to the enemy, and the computed visual perspective axis, onboard computers would then adjust the display of the video images onto the surface of the vehicle that is facing the opposing force in order to match the surrounding terrain, thereby removing the edges, surfaces, contours, color, and texture of the vehicle. Camera control and surface displays can be coordinated through the input and use of radar data, GPS coordinates, or approximate enemy location provided through direct visual observation.

Each of these technologies would need to be fully examined, tested, and refined (where appropriate) in

order to deduce their individual and collective relevance and contribution toward the creation of a fully functional prototype IONA system. The display, resolution, and array configuration as presented through either; the tapered fiber optics method, the embedded LED display method, or the lateral projection configuration, along with the degree of digital video interactivity and display precision as presented to the enemy, are considered to possess a significant degree of potential capability to maintaining effective concealment. Under combat conditions, the most optimal level of display precision for concealment (despite the projection method employed) would naturally occur when the IONA array is provided with the known location, direction, and elevation of an opposing force (through GPS, radar, or observation data). But, even when only the general direction of an enemy force is known, such a system is anticipated to significantly outperform traditional methods of camouflage.

Under such a scenario, dynamically changing video images could be continually relayed to the leading edge and surface display areas of the IONA as the vehicle approaches an enemy's general location and these surface areas would subsequently reflect the topography, color, lighting, and textures of the terrain that are located on the opposing axis of the vehicle, as provided by correlated camera movement. This interactive display technique would theoretically create a near transparent edge and surface profile for all exposed surfaces of the vehicle that would not only provide an enhanced level of concealment,

but will also afford a significant element of surprise for military forces.

In addition to mitigating visual signature, interactive opacity neutralization would also maintain a significant capability to mitigate an enemy's ability to accurately direct hostile fire upon friendly forces (even after detection) by diffusing edge surfaces, color, lighting, surface contours, and textures. The concealment and disguise of these types of variables, in turn, makes it more difficult for enemy forces to discern those specific secondary variables necessary for effective fire control such as speed, direction of movement, and angle of attack. By reducing an enemy's ability to detect friendly forces until they are within range to engage in offensive action and by further reducing that same enemy's ability to calculate the necessary fire control solutions to effectively return fire once the battle has begun, such a system as conceived within this paper could have a demonstrative level of benefit to operational forces.

Given the number of remarkable technical advances that have occurred during the past ten years in the area of computing power and display systems, it seems reasonable to conclude that this IONA concept is well within reach and that it's just a matter of time until we see this concept put to the test on the battlefield of tomorrow.

As you can see from this extensive articulation, there are a considerable number of practical applications, not just scientific or decision based processes, for multivariate

theory. The importance of including vertical and perpendicular axes within the logic becomes extremely apparent when engaged in the analysis of comprehensive studies. This particular example isn't that far removed from the traditional scientific analysis or social science study, or decision based application of logic. It may seem extraordinarily complex, but no more so than most decision processes we engage in everyday. The difference is the level of contemplation that we elect to employ. Needless to say, it would be easier to avoid the consideration of vertical and perpendicular variables, but the end result would be a product that is less than effective in achieving the goal of invisibility. So too would be our efforts to author a comprehensive public policy that was based on only a few critical level variables, but which failed to consider the vast realm of contributive influences that factor in to such an equation. Scientific discoveries that attempt to explain universal phenomena but fail to take in to account the complex dynamics of interactive agents, compounds, and physics would also fall well short of providing a thorough interpretation of how the world is put together.

Summary of the Multidimensional Process

I think that it should be relatively easy to see from this example that rendering a definitive explanation as to these influences that can affect science and technology almost exclusively depend upon higher reasoning and logic processes. Although it is certainly possible to render

decisions without fully examining all of the possible factors that influence the outcome, the results are more likely than not to be either less than accurate or limited in their ability to provide a definitive account for the factors that contribute to the outcome. Yet, people do it every day.

The proper way, in my opinion, to conduct multivariate analyses is to engage in the painstaking process that examines, not only the major level variables within the horizontal equations, but to push back the edge of discovery by considering and testing the vertical and perpendicular equations. This process affords the most comprehensive review possible of the primary, secondary, and tertiary influences and positions the researcher to not only isolate principle level variables, but also to expose subtle level variables. This can become very important in determining how manipulations of lesser variables might effect, or even contaminate, the situation. It also positions us with the ability to focus our attention on controllable factors in order to affect change. Naturally, some variables we will be able to exert influence over, while others may not be directly controllable, but might still exert an influence. By discerning those variables that lend themselves to manipulation, we further position ourselves to model how changes in the frequency or proportion of individual factors might alter the state and condition of the dependent variable or conclusion. Should changes to a variable that we exert no control over happen, awareness of the fully articulated logic model (i.e., a thorough understanding of all the variables contained within the

horizontal, vertical, and perpendicular equations) would support our efforts to either counteract adverse changes in non-controllable factors, or permit proactive change by manipulation of controllable influences. The bottom line here is simply that until we account for every ounce of meaningful variability, in all directions, we don't really fully understand the equation. As a result of this limited comprehension we are more prone to make mistakes. To limit this possibility, we simply need to extend, to the degree possible, our familiarity and knowledge of all of the factors that hold either a primary, secondary, or tertiary influence over the phenomenon we are trying to understand. Once we achieve this comprehensive model, we can then anticipate consequences of changes to the equation before they occur, or we can model the results of planned changes before we implement them.

Now it's time to turn our attention to the scientific protocols and processes associated with structuring a research design, proving a premise before we include it within our problem solving equation, and rendering a decision about why things happen. The next few chapters will provide an informative articulation of the processes that we use to discern the truth and assure that we fully understand why.

Chapter 5

Scientific Evaluation Process

Organization of any research design, even where we are merely trying to prove the premise within an argument, requires structure, clarification, and language precision. It is important to remember that confusion and misstatement are often the product of the poor organization of one's thought processes. As a result, it can lead to uninformed judgments or the failure to recognize distinct differences between variables pertinent to an argument. To preclude this from occurring, we use a formalized structure that prescribes the essential elements of our deliberation and research. This structure breaks down the thought processes we are using and assures specificity in our delineation of the issues and factors we are examining, as well as the hypotheses we have formulated to explain our postulates.

The traditional scientific protocol includes the following:

- Problem Statement
- Research Question
- Research Hypothesis
- Null Hypothesis
- Probability Statement
- Presentation of the Data
- Examination of the Null Hypothesis
- Interpretations

Problem Statement: this is the basis or reason for conducting a research study. These are statements that delineate a difficulty or shortcoming, or which articulate the (general) premise of an argument.

Research Question: the research question inquires about the specific relationship between the independent and dependent variable, or which postulates an expression of inquiry with respect to a decision in favor of one perspective over another. Essentially, the research question captures the essence of the research study, but does so in a question format, which leads to the actual collection and analysis of data germane to the study.

Research Hypothesis: These specific statements and expressions tentatively propose a difference between groups or relationships between variables. The RH states the expectations of the researcher in positive terms. The research hypothesis should evolve from a literature review where pertinent theory is identified and serves to justify an explanation for a difference or correlation between the variable/s. A Research Hypothesis is a declarative statement that (again) is always expressed in the affirmative. It is based on sound theoretical postulates prescribed in the literature or is predicated upon an observation of a potential explanation to the phenomena under study. However, the RH is NEVER accepted as a statement of truth until conclusive proof has been attained because of the research effort. At the outset of the research, we always cling to the Null Hypothesis. The

reason for this is that it assures that we never make the mistake of accepting a fallacy (like the earth is flat) and contaminating our future judgments based on an error in the equation. The data selected for inclusion within the study must also accurately reflect the variable being studied and the results of the analysis must specifically answer the research question.

Null Hypothesis: The Null Hypothesis is always expressed from a negative perspective on the issue and we cling to these perspectives in the face of conclusive proof offered to the contrary. In other words, the NH must always indicate that NO statistically significant difference exists between the arithmetic means for the two groups, relative to a particular variable, so there is no significant difference in behavior between these two groups. This negative assertion precludes any possibility that we might falsely embrace a presumption without proof.

An example of a correctly stated Null Hypothesis might be, “NO statistically significant relationship exists between unemployment rate and crime rate for the population sample examined”. This is the presumption we would cling to in our research design. If this were found to be true, you could eliminate unemployment as a variable that affects criminality. If, on the other hand, it were determined (after exhaustive scientific study) to be significantly influential) then you would conclude that such a variable does influence behavior and might be relative (in part) to explaining human behavior. In a Z ratio form of analysis of

such a hypothesis (which is used to measure significant difference between arithmetic means between two groups), the null hypothesis would be stated; “there is NO statistically significant difference between the mean unemployment rate for those who committed criminal acts versus those who were surveyed and found never to have committed a crime”. As you will discover later, Z ratios are used quite commonly to measure whether a difference exists between averages and to mathematically test whether such differences are so large that they could not have happened by mere chance. In those instances where the Z ratio is larger than 1.96, the difference is presumed to be so large (factored against the standard deviation of both means) that it could not have happened by mere chance, and subsequently indicates a difference in values beyond what would be expected as a result of simple probability.

An NH that expresses that no difference exists between these two groups is tested by the researcher. If the NH is found to be tenable (in other words proven to be true and accepted), then a statement is made that evidently no difference exists between the samples (relative to the mean values). This finding consequently would also infer that the Research Hypothesis cannot be accepted or supported. On the other hand, if the Z ratio computed is greater than 1.96, then the difference between the sample means is so large that it could not have happened by mere chance and the Null Hypothesis (that said there was NO difference between means) is rejected. The RH would then

be accepted as truth. You will be provided with many more examples of how this approach is applied to critical thinking and decision making in later chapters, so do not despair if you do not fully comprehend the application at this early stage of your studies.

Probability Statement: These are statements of the values needed for the critical ratio or statistical test in order to be able to reject the H_0 at either the .05 or, .01 levels of significance. Probability statements indicate whether the test is one or two-tailed (we will talk more about these later too), and the degrees of freedom involved, if appropriate. There are no degrees of freedom (df) used for Z ratios, but **df** will become important in Chi-Square, Student's T ratios, and Pearson R evaluations.

As applied to a two-tailed Z ratio test for difference between means, a Z of plus or minus 1.96 is needed for significance at the .05 level (always), and plus or minus 2.58 at the .01 level (always). The level of statistical significance required for the study should always be established prior to the evaluation and data analysis.

Presentation of the Data: In most studies using ratio analysis, the number of participants (N) included within the sample, the Mean (M) value computed for each sample, and the Standard Deviation (S) computed for each group are presented in tabular form. This approach reveals the data used to compute the value of the statistical

measure and allows for replication of the scientific study by others.

Mean Unemployment Rates per 10,000 Population			
	N	M	S
Group A (Criminals)	35	3.5	1
Group B (Non Criminals)	35	1.2	.90

Examination of the Null Hypothesis: In this section, you provide a statement of whether the NH is accepted or rejected and at what level of significance, based upon the Z ratio computed.

For Example: $Z = 10.11$

The NH is rejected at the .01 level of significance; $P < .01$

Interpretations: This is a discussion of the statistical significance and inference of the findings. For the example cited above, the Z ratio that was computed was so large (10.11) that it exceeded the value required for embracing the null hypothesis (1.96 at the .05 level and 2.58 at the .01 level) and could not have been based on mere chance. Subsequently, the difference detected between the mean arrest rates for the two groups is confirmed for the sample study and will likely be present in a study of the entire population as well. Accordingly, we are safe in assuming that there is indeed a statistically significant difference

between mean arrest rates of those who were employed versus those unemployed. Based on this empirical confirmation, we can reject the Null Hypothesis and accept the Research Hypothesis.

Now, a word of caution, in this analysis we only proved the premise that there is a difference in the unemployment rates for each group. We did not explain “why” there is a difference. That assessment comes well after we test the truth of the premise that a difference actually exists.

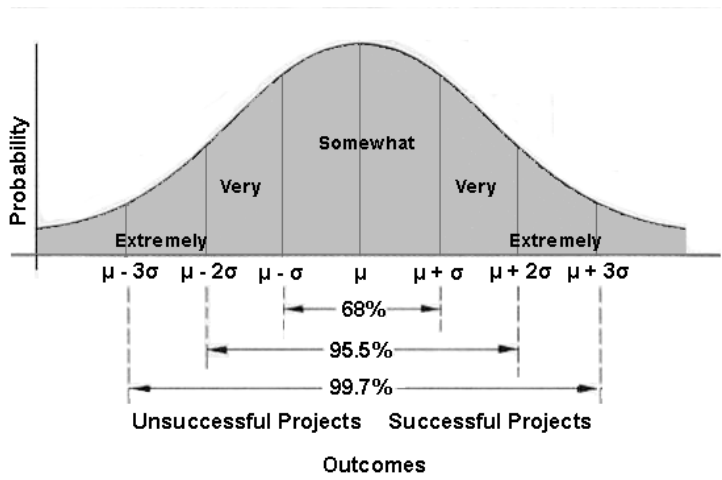
The statistical findings are related back to the theoretical structure from which the hypothesis was derived. Conclusions may be drawn about the findings relative to the research question only. At this point a more comprehensive search for a plausible explanation as to why there is a difference between these groups is needed, but (based on the Z ratio computed) you have proven the premise that a significant difference exists in the unemployment rates for the two groups sampled.

In the next chapter, we will spend more time discussing the specific process of proving a premise using the Z ratio method, but at this point you should have a pretty fair idea of how statistical measures can be used to prove or disprove the truth of an assertion.

Now it's time, to turn our attention toward some basic statistics to refresh your memory.

Normal Distributions

As we discovered in the previous example, statistical significance is important to hypothesis testing and eventual assessments of the truth of related assertions. To refresh your memory the graphic below is provided. In this illustration, you will see that [in a normal distribution] the arithmetic mean (or average) is represented by the vertical line in the center of the diagram as well as two tails of the distribution that are based on variant values for the population sampled. Since not everyone will likely have the same average value, a degree of variance will likely be seen for all participants in the study.



As part of the analytical process, we compute an arithmetic mean for the sample that reflects the average value, as well as a standard deviation for the sample. Under statistical theory, a normal distribution would possess several distinct cut off points (or standard

deviations) that we could use to analyze our data. As depicted in the graphic, under statistical probability we would presume that 68% of the participants would possess a value that was plus or minus one standard deviation. Likewise, it is reasonable to assume that 95.5% of the population would possess a value that was between plus two and minus two standard deviations. Finally, we would expect 99.7% of the population to reflect a mean value that is between plus three and minus three standard deviations from the mean for the group. This “normal distribution” is the keystone for all statistical analysis and is the basic presumption for relevant interpretations.

When we conduct our Examination of Null Hypothesis, we are using this normal distribution to measure the value of the statistical test we employed, as compared against the values needed for the experiment to assure confirmation. In other words, as scientists, we never accept the assertion of a research hypothesis unless we discover a statistical computation that places the value beyond the 95.5% percent (.05 level of significance). Similarly, a statistically significant difference measured at the .01 level would be based on a calculation of the statistical test that was beyond the 99.7% (.01 level of significance) for the population. We can also interpret such a finding to mean that if we conducted this sample 100 times, we would expect similar results 99.7% of the time.

Let me provide a practical example to lessen the brain pain you are experiencing right now. If we observed one hundred sprinters who were over the age of 50 and one hundred sprinters below that age and then computed the average times for the two groups in the 100-yard dash, we would discover that not all the runners would have the exact same speed. Both categories of runners, would have participants who ran faster and slower than the average for the group.

In the over 50 category, some would run faster and others slower, but the average for the group might be 15.9 seconds. Based on the variance for the group (which is derived by calculating all the times and then mathematically determining the standard deviation – let us say 2 seconds) we could construct a normal distribution chart similar to the one in the example. If there was a normal distribution for the group, we would expect 68% of the people in the over 50 category to have run times that were within one standard deviation or not faster than 13.9 seconds and no slower than 17.9 seconds. Likewise, we would expect 95.5% of the sample to post times no faster than 11.9 seconds, and no slower than 19.9 seconds. At the third level of measure we would expect that 99.7% of the population would post run times no faster than 9.9 second and no slower than 21.9 seconds (or plus and minus three standard deviations away from the mean of 15.9 seconds).

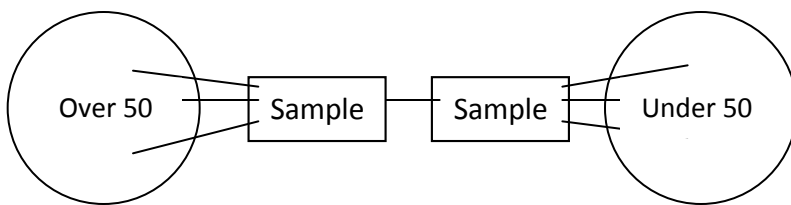
Inference and Normal Distributions

Within an inferential analysis we use these values (N , M , and S) to compare one set of findings for a particular group to the scores of another group in order to discern (and thus infer) the principle that [age] matters relative to speed in the 100 yard dash. If there is a statistically significant difference in the mean time between the two groups in the samples, then we would infer that the difference measured within the samples also (probably) applies to the entire population.

The Z ratio (which is commonly used for large sample sizes or samples above 30) is an excellent example of the inferential method for testing the truth of a premise and is computed by comparing the means and standard deviations between the two groups. We will discuss this technique in considerable detail in the next chapter, but it seems prudent to set the stage here and explain how such a statistical measure applies to the notion of inference and normal distributions.

If a Z ratio was calculated that was larger than 1.96, we could conclude that the variable “age” is relevant to the average speed of a runner in a 100 yard sprint. If not, then we accept the H_0 and look for another variable to explain the difference. The illustration below provides a reference to the sampling process that might be employed in support of such an analysis. As you can discern, representative samples are drawn from two “total”

populations. One representing the population of people over 50 years of age, and the other consisting of the total population under 50 years of age. The inferential comparison that is conducted is based on the N , M , and S for each sample and if the results are deemed statistically significant (meaning the difference computed in the average run times was beyond the .05 level) then the results observed between the two samples could be “inferred” back to the total populations. This not only proves the premise that “age” influences the average speed in the 100, but also sets the stage for inclusion of the premise within a comprehensive argument.



Later in the book we will discuss a process known as “Discriminant” analysis, which is used to discern group association, based on values which influence propensity or association. This is an extremely complex form of multivariate evaluation that uses a variety of iterative computations to disclose the strength of association, group polarity, relative influence of predictor variables, and finally an equation that provides the researcher with a method for determining the probability of group association. We will, in great depth, discuss this method later, but at this point it might be helpful to reiterate the importance of illustrating the sampling and evaluation

processes in order to achieve that step backward we've been referencing in order to see the whole board and assure that our logic is sound.

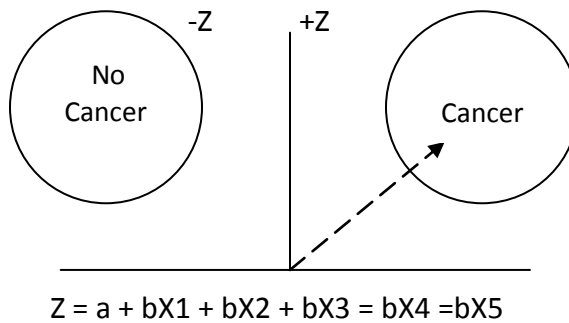
When we are involved in comparing the difference between means of two groups in order to determine the univariate relation between the groups based on a single variable, our only concern is drawing samples that reflect the total populations so we can discern whether the difference between means is so large that it could not have happened by mere chance. If it is determined to be large enough not to have happened by mere chance, then we are safe in inferring that the difference we discovered in the sample probably applies to the total population and interpret the difference as being significant in our search for factors that might explain the phenomenon we are studying. In the previous example the variable [age] was found to be a significant factor relative to influencing human athletic ability, as measured by running speed in the 100 yard dash. There are times however when you might want to not only examine the univariate influence of such a variable to see if it is a credible factor, but also to study whether it is influential as part of a larger equation (or argument). In such an effort [age] would still be treated as a premise, but also part of a greater equation that seeks to determine group association.

For example, if we were searching for an explanation as to why someone contracts lung cancer as opposed to those who do not, we could start by structuring our sampling to

isolate these two distinctively different groups of people. In a number of books that speak to the topic of research methods, you will find references to the importance of random samples to prevent contamination, but unfortunately, many of these texts do not also focus on the merits of purposive sampling, so that you can isolate specific populations. In our research designs it is important to determine not only the focus of the study, but also the sampling and quantification strategies that we need to employ to accurately reflect the target populations we plan to analyze. When involved in trying to answer the question regarding why someone ends up with lung cancer or not, it is often advantageous to use a multistage stratified random sampling approach, as opposed to a purely random, so that we isolate specific sub-groupings of people/cases that most closely represent the populations pertinent to such a study. The point here is to not get caught up in the scientific purity of the process, at the expense of the objective. Sampling strategy matters a great deal to assuring that the cases contained within the study not only reflect an accurate cross-section of members of each grouping, but it is equally important to assuring that the data drawn from the samples is reflective of the premises within the argument. It does no good to include participants in the sample if the values gleaned do not accurately match the scientific criteria being used to conduct the study and evaluate the postulates. To avoid this, we often resort to multistage stratified random sampling to assure that we are targeting populations that most closely represent the

groups we wish study and then we randomly sample from inside these stratifications in order to obtain specific data we require to conduct our analysis of the hypotheses. We will talk more about this concept in later chapters, but I thought it important to include a mention within this section.

Using such a multistage stratified random sampling approach as in the lung cancer study, we would still draw upon two distinctly different populations (the first consisting of those who were diagnosed with lung cancer and the second consisting of a broad cross section of people who have not been diagnosed with such an illness). From these two samples we would draw a representative number of observations that are pertinent to those variables that we presume might influence whether you contract such a disease or not. Naturally, this listing or variables would contain such things as diet, air pollution exposure within working conditions, history of tobacco use and other drugs, family medical history, and a host of potentially pertinent factors, each supported by a sound theoretical basis for inclusion. For the univariate comparisons of such variables (i.e., to discern whether there is a statistically significant difference in the mean number cigarettes smoked daily by members of the two samples) we would use the exact same process as previously identified in the preceding pages, but for the multivariate argument we could graphically illustrate the logic of our argument as seen in the next illustration.



The **b** factors (as you will discover later in the book) refer to the unstandardized discriminant function coefficients assigned to each variable based on the analysis of the data for each case in the two groups. The X factors represent a multiplier for each individual case. In other words in the X part of the equation, the researcher could place a value for a specific individual in order to discern the mathematical probability of their particular group association. This approach is extremely valuable in determining the likelihood of a patient who wasn't part of the original study and is outside the sample in order to determine their individual probability of falling into one group or the other.

The point here is that graphic illustrations not only help to achieve the visualization of the sampling processes required, but also can be used to articulate the logic structures for the study as well as the relevant equations. By using this form of visualization within the scientific research design, you can communicate effectively how the

data are to be obtained and how they can be applied. This helps to express the scientific methodology that was used to perform the analysis, as well as fostering a greater understanding by those reviewing the report in discerning the logic supporting the effort. Such an approach also helps you as the researcher visualize the approaches you have decided upon to support the study so you can evaluate the appropriateness of your own logic.

Quantification

Quantification is an important component of inferential analysis and multivariate reasoning, because it determines how we assign values to variables in order to test them. As you have learned in your statistics classes, there are several different data types used to quantify information. The four basic types of data are categorized as:

Nominal Scales

Ordinal Scales

Interval Scales

Ratio Scales

Nominal scales of measurement typically involve the assignment of values in order to classify or categorize responses and traits. They provide extreme flexibility to the researcher and facilitate the assignment of values to support comparisons. They are considered a qualitative form of analysis, yet they often are used to yield a substantive level of insight into behavior, characteristics, and traits that could not otherwise be effectively

evaluated. Using such a categorizing procedure, the researcher must be careful to include all cases in the sample and also to assure that no single case is assigned to more than one category. Nominal values are often used to measure qualitative factors in argument analysis. Some examples of nominal value quantification might involve assigning a value to the presence or absence of a trait.

Non Smoker = 1 Smoker = 2

Not Married = 1 Married = 2

Alive = 1 Dead = 2

Low Risk = 1 High Risk = 2

As you can see from these examples, a nominal value is assigned to each category and as you evaluate each case, you assign a value based on the category. I have found it advantageous to avoid the use of [zero] as a value because many of the computations used in advanced forms of analysis will encounter a problem when they try to divide by zero. Subsequently, since these are arbitrary values, one and two work just fine. If you have multiple categories for responses, you can quantify each case using an advancing numeric value (one, two, three etc). Remember that the value you assign does matter in analyzing and interpreting the data. For example, if you plan to conduct a Discriminant analysis, a Z ratio will be computed to determine probable group association and the value for each case will either have a positive or negative Z value. Accordingly, you need to make sure that the quantification

matches the logic you are using in your study, along some line of intuition, so that you avoid misinterpreting the product computed in the equation. By assigning a lower number value to the lower risk category for example might prove easier to interpret than vice versa.

Ordinal scales are somewhat similar to nominal scales in that they also use mutually exclusive and exhaustive categories, but in these instances, the categories are ordered along a continuum according to hierarchy. Using ordinal scales, data are ranked and ordered from low to high or high to low depending upon the strategy. The key however, is that the categories used and the distance between the categories (numerically) have no meaning or inference.

An example of ordinal data might be:

Educational Attainment as

- 0=less than H.S.
- 1=some H.S.
- 2=H.S. degree
- 3=some college
- 4=college degree
- 5=post college

You can use ordinal scales to effectively measure and represent socio-economic status, political party affiliation, rank in a hierarchy, agree-disagree responses, and similar

categorical measures. Ordinal scales support qualitative analysis.

Interval Scales are considered one of the higher forms of quantification. They are more precise than nominal or ordinal as data are distributed along a continuum with established distances between points. Interval scales can be used to represent a variety of standard measures such as temperature in Celsius, highest grade completed, and gestation period as measured in days. Interval data is continuous data where differences are interpretable, but where there is no natural zero. Interval scales are used when distance between the data points is meaningful. Tests and evaluations of data using interval scale sets are considered quantitative in nature.

Ratio scales are very similar to interval level data but are considered the highest or strongest and most precise level of measurement. Data are distributed along a continuum with established distances between points, but in ratio scales, there is an absolute zero assigned to the scale. Analysis using ratio data is considered quantitative in nature. Examples of ratio data might include age in years, weight in pounds, time in hours, and money in dollars.

Despite the various categories of measure or scales, the important point is to recognize that both qualitative and quantitative data can be used to support analyses. Scientists occasionally get caught up in debates over mathematical purity, but knowing how and when to use these various types of data can work to your advantage. Certainly each has limitations, but all support testing of a premise and can be used to evaluate the conclusion put forth in an argument. Your greatest task will be to determine which measurement scale best suits the analysis you have chosen to undertake and how to analyze the data in order to make the correct determination.

Chapter 6

Proving a Premise with Z ratio

Inferential statistics serve as the basis for much of the scientific thinking and reasoning that you have been reading about in this text. This is especially true with regard to **proving** a premise before including it within an argument.

Inferential statistics involves drawing samples from populations, quantifying the data pertinent to each participant or phenomenon, and then making inferences about the entire population through the examination of the data contained within a sample. One such technique involves drawing two samples from their respective populations, computing the means and standard deviations relative to a designated quantitative or continuous variable (such as age, or education) and then testing to see if there is a statistically significant difference between the means of the two samples.

If the difference is so significant that it could not have happened by mere chance, then the inference is made that there is a significant difference between the means of not only the sample, but probably the total population as well. In other words, we "infer" the findings of the sample to the total population, based on probability.

Thus, the approach includes one variable measured along a quantitative scale, and two sample groups. A critical ratio test (Z ratio) can be used to determine if the difference between the means is statistically significant or not. As applied to critical thinking, before we state in an argument that age is a factor in involvement in driver safety and involvement in traffic collisions, we should test the hypothesis by collecting two samples at random, one for a group of people who have been involved in a collision over the past three years and another group of people who have not been involved in a traffic collision. If we find that the mean age for the accident group was 22 and the mean age for the non-accident group was 32, we would see a visible difference between the mean ages. If the Z ratio was computer beyond a level needed for statistically significant difference (say 2.65), then we could conclude safely that the same level of age disparity observed in the sample probably exists in the total population as well and thus our assertion that collision participants tend to be younger would be a safe presumption. Conversely, if we draw two samples with the same mean age, and test the hypothesis to discover there is no statistically significant difference in the mean number of traffic collisions, we have furthered our understanding that [age] might play a role in determining propensity for safe driving. Again, such a comparative analysis does not prescribed [why] this occurs, only that there is a measurable and statistically significant result.

Statistical significance (based on the laws of probability) is a conventional way of stating whether a difference between groups or a relationship between variables has occurred simply by chance, or not. When comparing means of two samples, a Z ratio helps the researcher decide whether the difference would be expected to occur by chance, or whether it would not be expected to happen by chance. This also facilitates our determination of whether the difference may be attributed to random sampling error, or whether the difference between the means was so large that it overcame expected sampling error. The .05 and .01 levels have been accepted by the scientific community as the standard cutoff points of acceptability for such measurement. In other words, distribution away from the mean of plus 2 and plus 3 standard deviations. WHAT DID HE SAY!!!!

Think of it this way, our goal is to measure whether the difference is so significant that it could not have happened by mere chance. To accomplish this, we have to rule out that the difference between the means we found in our samples was caused by mere chance or sampling error. Accordingly, we need to conduct an empirical test to assure ourselves that the premise (that age is a contributive factor) is correct. We accomplish this by incorporating a quantitative analysis that examines the means of two samples to see if there is a [significant] difference between the mean ages of people involved in traffic accidents. To be absolutely sure that sampling error wasn't the cause of the difference that we noted between

the two groups, we start out by accepting the null hypothesis (which would state; There is NO statistically significant difference between the mean age for those who are involved in traffic accidents and those who have not experienced such problems) until the Z Ratio proves, beyond reasonable doubt or mathematic certainty, that such a difference exists. The .05 and .01 levels of significance correspond to the 95% and 99% (or plus 2 and plus 3) standard deviations on a normal distribution bell curve. If a difference between means produces a Z ratio that is large enough (plus or minus 1.96) so that it would be expected to occur by chance in less than 5% of the cases (.05 or 5 times out of 100), then the difference between means is said to be significant at the 5% or (.05) level. If we see a Z ratio equal to or larger than 1.96, on either side of the curve, we can safely "reject" the null hypothesis and accept the research hypothesis, which says, "there IS a significant difference between the two groups relative to mean age", subsequently we can also safely presume that age is ONE of the factors that influences driver safety.

Z ratio is a statistical test that can be used to examine the difference between means of samples drawn from populations. It is considered a critical ratio test for sample sizes larger than 30 ($N > 30$). Z ratio assumes normality of the distribution. Areas under the normal curve may be examined to determine within that level of probability the difference observed between means would have occurred by chance.

Proving a premise before including it within an argument is a requirement if you stand any chance at all of being accurate. To say that age matters in determining driver safety might be an interesting presumption, but certainty requires that the arguer take the time to establish the research protocol that specifically tests the accuracy of such a presumptive statement. Proof is essential before forming a conclusion and creating policy. There may be (and usually are) many variables that influence the outcome of things. Age might be one of 50 variables that contribute to driver safety. The Z ratio lets you prove the truth of the premise, but does not disclose what the other remaining 49 variables might be in such an equation. To identify all 49 factors, you have to study the problem, identify potential influences, test each one of them individually to discern whether they have any contributive value, and from this aggregate analysis, you can formulate a comprehensive argument about the factors affecting driver safety.

To further our proficiency at conducting this form of premise analysis let's examine it in context to the scientific evaluation process discussed in the previous chapter. We start out by articulating a problem statement, followed by a research question, the research and null hypotheses, the probability statement, the data, then the examination of the null hypothesis, and finally the interpretation.

Problem Statement: Rehabilitating Juvenile Offenders

Research Question: What is the relationship between the recidivism rates for juvenile offenders who received treatment as part of a diversion program versus those offenders who were strictly confined in a youth detention facility?

Research Hypothesis: There is a statistically significant difference in the mean re-arrest rate per client for those treated in the diversion program as compared to those who simply receive incarceration.

Null Hypothesis: There is no statistically significant difference in the mean re-arrest rates between these two groups.

Probability Statement: Upon initial arrest, juvenile offenders were randomly selected and either assigned to undergo diversion treatment as part of their sentence or they were placed exclusively within juvenile detention. Re-arrest records were maintained for a one year period for all juveniles included within the study. For a one tailed test, a Z ratio of 1.64 is needed for significance at the .05 level, and 2.33 for significance at the .01 level.

Presentation of the Data:

Mean Re-Arrests and Standard Deviations Of Juvenile Offenders			
	N	M	S
Incarceration	400	1.2	.90
Diversion	400	.75	.25
Z Ratio = 9.63, P<.01			

Examination of the Null Hypothesis: For this analysis a Z ratio of 9.63 was computed. The NH is rejected at the .01 level of significance, P<.01.

Analysis and Interpretation: Predicated on the results of this analysis, it appears prudent to conclude that diversion programs do have a favorable impact in preventing recidivism rates among juvenile offenders. As applied to the samples, we see that for those juveniles who were simply incarcerated, the mean re-arrest rate was 1.2 arrests during the subsequent year, while for those who received diversion, the mean re-arrest rate was only .75 instances. Clearly diversion isn't the only factor associated with juvenile delinquency, but based on these data; treatment of offenders (as opposed to mere incarceration) does appear to have a positive effect in influencing future behavior.

As you can see from this example, the premise that treatment has a favorable impact on future criminality was proven to be true. Not exclusively true, but you can safely infer that treatment facilitated by placement in a diversion program does have a favorable impact on future criminality. The results of the analysis within the sample were proven true beyond reasonable doubt and as a consequence of computing such a large Z ratio (9.63) we are safe in inferring the results experienced in the study of the samples back to the entire population. That's all there is to it.

If this (diversion) were a variable in a larger argument, that included a number of premises associated with criminality, then the premise that there is a significant difference in the mean re-arrest rates for those who receive diversion versus incarceration, and that such an approach might be more effective in dissuading future criminality, would have been deemed a true presumption, leading the way for inclusion of this factor within the larger argument. This analysis didn't prove that it is a correlated variable, or even influences human criminality, only that there is a difference between the mean re-arrest rates for those who receive treatment versus those who are merely confined. This discovery does however lead the way to the next iteration of analysis, which would involve testing of the variable's relevance within a correlation model.

Let's look at another example of how Z ratio might be used to prove the truth of a premise.

In this example the postulate is made that formalized education in critical thinking and reasoning maintains a direct relation to student ability to perform such processes. Such an assertion would naturally be contained within a comprehensive multivariate argument that asserts that education in critical thinking is [one] of the components associated with student understanding and ability in the area of reasoning. In addition to formal training might also be a number of other variables that have relevance to increasing awareness, affording insights to processes, familiarity with the concepts, etc. But, in order to prove the truth that formal education in critical thinking might be a relevant factor in determining critical thinking proficiency we would need to test the truth of such a premise before including it within in a multivariate argument.

Below is an example of how we might structure a research design, using Z ratio, to test whether there is a statistically significant difference in the mean test scores for an examination administered to measure knowledge of the subject, between people who have completed formal instruction in critical thinking and those who have not completed such training. The underlying theoretical framework for such a contention would naturally be oriented toward an assertion that training improves ability.

Problem: Determining the effectiveness of critical thinking education in undergraduate programs.

Research Question: What is the relationship between mandatory critical thinking and reasoning education at the undergraduate level and test scores measuring proficiency in this academic area?

Research Hypothesis: There is a statistically significant difference in the mean test scores between students who were required to complete a class in critical thinking during their undergraduate education and students who did not complete such a class.

Null Hypothesis: There is no significant difference between mean test scores for these two groups.

Probability Statement: Prior to graduation, students were randomly selected to participate in an experiment that endeavored to measure their knowledge and familiarity with the concepts of critical thinking and reasoning. Students were grouped according to whether they were required to successfully complete a class in critical thinking during their undergraduate program versus those who did not complete such a course of instruction. For a two-tailed test, a Z ratio of ± 1.96 is needed for significance at the .05 level and ± 2.58 at the .01 significance level.

Means and Standard Deviations of Students Test Scores

	N	M	S
Completed Critical Thinking	361	94	5
No Critical Thinking Class	324	83	3

Z Ratio = 35.31, $P < .01$

Examination of the Null Hypothesis: The NH is rejected at the .01 level of significance; $p < .01$

Interpretation: Analysis of the data indicate that there is a statistically significant difference between the two groups, measured at the .01 level, suggesting that college students who are required to complete a course in critical thinking during their undergraduate program score higher than other students given the same examination. These findings suggest that completion of a critical thinking class during the undergraduate years serves to produce a more highly refined level of reasoning in college students (as measured by the average scores on the examination) and it can be subsequently argued that the experience better prepares students to engage in reasoning and logic.

To prove this later assertion however, we would need to conduct a more comprehensive analysis, but based on the findings of the study, there does seem to be cause for further investigation into the situation. You can probably see where else such a conclusion might lead an academic administrator or faculty committee. If such a hypothesis is proven true, and the goal is to assure that all graduates possess skills at critical thinking and reasoning then based on the findings academic policies might be changed to assure that all students are required to successfully complete a critical thinking class as part of their undergraduate studies, no matter their major.

This is precisely how statistical analysis and validation of premises serve to support policy change. By utilizing this type of information, policy makers can influence the outcome, predicated on science rather than intuition. If the variable being examined is deemed to be within the control of the policy makers, then it can be manipulated to affect the outcome. This approach is used in all forms of government enterprises to “engineer” change or in purely scientific areas where controllable variables (such as above ground biomass levels) can be influenced to affect an associated variable such as species population. We will probe further into controllable and non-controllable variables later in the book, but it seemed particularly relevant at this point to mention it.

Continuing on in our examination of proving the truth of premises using Z ratios, you will find this next example enlightening.

Problem: Determining the impact of clear cut forestry practices on soil erosion.

Research Question: What is the relationship between clear cut forest harvesting versus selective cut forestry practices as applied to soil erosion?

Research Hypothesis: There is a statistically significant difference in the average number of inches of soil lost to erosion between regions using clear cut harvesting practices versus regions where selective cutting is used.

Null Hypothesis: There is no significant difference in soil erosion as measured in annualized loss as measured in inches between these two harvesting practices.

Probability Statement: To facilitate this analysis, a representative sample of ten (separate) forest regions was used for each category of harvesting method. Data pertinent to soil erosion was collected for each type of region (cut clear and selective cut methods), and assembled. The data represents average soil erosion as measured in inches of top soil lost for the five preceding years. For a two-tailed test, a Z ratio of +/- 1.96 is needed for significance at the .05 level and +/- 2.58 at the .01 significance level.

Means and Standard Deviations of Soil Erosion			
	N	M	S
Clear Cut Regions	10	1.2	.35
Selective Cut Regions	10	.95	.31

Z Ratio = 1.69

Examination of the Null Hypothesis: The NH is accepted at the .05 level of significance; $p > .05$

Interpretation: Analysis of the data indicates that there is not a statistically significant difference between the two samples. Although the descriptive information clearly reflects that the mean soil loss rate for the clear cut regions was larger (1.2 inches per year) than for those areas where selective cutting was utilized (.95 inches per year), the difference computed was not so large that it could have happened by mere chance. Subsequently we must conclude that there is no statistically significant difference in the mean soil erosion rates for regions harvested using these two forestry practices.

As you can see from this example, the analysis does not meet the criteria established for rejection of the null hypothesis. This happens when the difference between means between samples is small and when the standard deviation values (as computed in this case) are somewhat congruent. If, for example, the standard deviation of the selective cut regions had been .03 instead of .31, a Z ratio of 2.25 would have been computed. This would have exceeded the value needed for statistical significance beyond the .05 level and we would have interpreted the results as significant enough to reject the null hypothesis and accept the research hypothesis. The means for each sample would have still been the same, but the standard deviations would have affected the resultant Z ratio calculation. We never really can be certain as to how the analysis of the data is going to turn out and we should never hope for a particular result. This is exploratory research and as such, we are not out to prove a particular premise. Our goal is the search for truth and our

interpretations must be based on a dispassionate, objective, and thorough review of the data. This is a particularly interesting example because of the fact that the study didn't reflect that a "statistically" significant difference exists between the two samples. This could well mean that the variable either isn't a good predictor for soil erosion or that some other variable needs to be included within the multivariate equation to affect the outcome. Logic suggests that if you remove all of the vegetation in an area that it becomes more susceptible to erosion than if you selective cut the region, thereby leaving a percentage of above ground biomass to act as a barrier to soil erosion. So what's missing? The answer might be found in the structure of the predictive equation and this precisely why we use multivariate theory to support such analyses.

As you can see below, soil erosion is a dependent variable in an equation that must take into account not only the forest harvesting practices used, but also the slope of the area that was harvested, the rainfall that occurred during the year, and the composition of the soil. These multiple factors [combine] to affect the outcome as measured in inches of top soil lost per year due to erosion. If for example you clear cut an area thereby reducing the biomass that would serve as a barrier and if that clear cut region is located on a slope of seven percent or greater, combined with a substantial rainfall amount for the subsequent year that is sufficient to initiate downward movement of soil, then mean top soil loss would be substantial. Conversely, if you examine a sample that is flat, then the biomass quantity removed and the rainfall might not have a significant impact on soil erosion.

This model brings up another interesting point which relates to assuring that the (1) the reality of the logic necessitates that the samples represent the risk factors (meaning you should set up categories such as flat, gentle, and steep slopes) and then measure the clear cut versus selective cut approaches are all three types of zones to determine the effect and (2) you need to assure that the sample region bear similar characteristics so that the study is comparing similar events. To draw a sample that was clear cut on a steep slope and then compare those results to a selective cut on a flat slope would induce a substantial level of contamination within the study. The Z ratio that would be computed would likely be enormous, but the sample procedure used would invalidate the comparison. You've heard the term Liar's figure and figures lie? This is exactly what that means and to remain credible you must, as the researcher, assure that your sampling strategy, data collection processes, evaluation protocols, and analyses conform to the logic of the study. Failure to assure any one of these can result in not only false discoveries, but condemnation from the scientific community. One final note of caution about proving premises with Z ratio before we move to the next chapter and that is be very careful not to overstate your findings. In the last example (clear cut v. selective cut) even if we did note a statistically significant difference, the difference was only .25 inches of soil per year. The temptation is to proclaim "Eureka I've Found it!!!" when in fact although the difference might be "statistically" significant, it is only a quarter of an inch of top soil per year. Is that really significant? Overstating the enthusiasm of your discovery might become the point of contention whereas a dispassionate proclamation might better serve your interest.

As you know by now, Z ratio analysis only helps you prove the truth of the premise that a difference exists and that the variable “might” be influential. More study is needed to discern the degree of influence that this factor may have in determining the outcome of soil erosion. Even if harvesting method were discovered to possess a significant correlation to soil erosion in the final analysis, the difference is only .25 inches per year. Your interpretation should reflect the realism of the phenomenon and not be used as a source of criticism. If the mean loss were two to three inches per year, then you could argue with vigor that the results of the study have a potential for dramatic reductions and that changes in policy to forestry practices are an absolute requirement. Such a finding would also set the stage for a follow on study where your findings that such a large percentage of soil erosion is prevalent annually that it may well be a critical factor in related environmental models such as river and stream sedimentation, which might be related to fish population equations, and even river flow models that are affected by closures to the mouth of the river that result in flooding. As you can see from this extrapolation, there are a god number o related chains in the logic and one model is a small portion or branch of an even larger model, and so on and so on. We will discuss these interdependent logic arrays in the chapters dealing with vertical and perpendicular logic, but the processes are you learning here have a direct relation to such complex logic arrays.

Chapter 7

Proving a Premise with Student's T

It should be evident by now, that guessing and speculation are not necessary when involved in trying to discern the truth of an argument and the relevance of specific factors to your suppositions. There are a plethora of scientific testing mechanisms available that can provide insight and confirmation regarding the truth or fallacy of an assertion. Use of such methodologies is important when formulating your own thoughts about the world, as well as when evaluating the arguments of others.

The use of Student's T ratio to support inferential analysis has long been held as an advantageous mechanism for comparative analysis. This is especially true for proving the truth of a premise prior to its inclusion within an argument. Similar in nature to the Z ratio form of analysis, Student's T provides for the comparison between means and standard deviations, but it is especially effective for small size samples. As you discovered in the previous chapter, Z ratio analysis is used for larger samples (30 and above), however Z ratios can be ineffective for analysis of smaller samples. Student's T was specifically designed to meet the challenges of conducting comparative analysis for smaller samples and provides a stricter level of measure before rendering a verdict as to whether or not to accept or reject the null hypothesis.

Student's T is considered a critical ratio test for significant differences between means and was originally developed by William Gosset, writing under the pseudonym

“student”. Although it can be used for any size of sample, Student’s T is especially well suited for analysis of very small samples (less than 30). The T sampling distribution is more leptokurtic than the Z ratio distribution (meaning that it possesses a higher peak valuation) and as a consequence, the value required for attainment of statistical significance is located further out on the scale.

As you will recall from your previous experience, Z ratio requires at least 1.96 for significance at the .05 level, and 2.58 for significant difference between means at the .01 level. Student’s T, because of its sensitivity, typically requires values beyond these levels for attainment of significance. If the total sample is approximately 1000, then the values required for significance for Z and T will be about the same, but in lesser samples, Student’s T requires a greater level of difference before confirming statistical significance. Student’s T is far stricter and accordingly requires a greater value in the resultant computation in order to achieve statistical significance, than we see using Z ratios.

Another important difference between Z and Student’s T analysis is that the T ratio requires the use of Degrees of Freedom (df). When working with quantitative data, variance must be considered, which is determined by the examination of cases or values away from the mean. Student’s T takes this variance into consideration. In order to calculate the degrees of freedom for Student’s T all you need to do is remember that $df = N_1 + N_2 - 2$, or, the total of the first sample, plus the total of second sample, minus 2. If for example there were 24 in the first sample and 18 in the second sample, the degrees of freedom calculation would be $24 + 18 - 2$, or $42 - 2$, or $df = 40$. Once you

obtained the DF you can then look up the required values for statistical significance at the .05 and .01 levels in the appendix of any statistics book. In this case (where $df = 40$), a Student's T ratio of plus or minus 2.021 is needed for significance at the .05 level and plus or minus 2.704 would be required for the .01 level. That's all there is to it.

The formulae for Z and T ratio bear a striking similarity. The major difference between the two forms of analysis rests in the fact that T ratio is a more sensitive measure and has considerable application for small sample sizes and incorporates a measure for degrees of freedom in computing the required values for statistical significance. Z ratio (as we already learned) has a standard measure for significance at the .05 and .01 levels (1.96 or 2.58), whereas T ratio value requirements are ever changing, depending upon the degrees of freedom. Analytically however the processes are the same for structuring the research design, formulation of the research and null hypotheses and interpretation the results.

If you examine the appendix of this text you can see how the values needed to for attainment of statistical significance in order to reject the null hypothesis. These values change when using Student's T ratio depending upon how many degrees of freedom are observed and computed. It is important to note that Student's T ratio values for one tailed tests differ from two tailed tests. In other words, a different value for statistical significance must be obtained depending upon whether you are conducting a one or two tailed test. Generally speaking one tailed tests require a lesser value in the T statistic for significance than are needed for two tailed analyses.

Probably the easiest way to see how Student's T ratio is used to prove a premise and evaluate a hypothesis is by looking at it in action. The example below profiles exactly how this is accomplished.

Problem: Assuring improved harvest yields per acre based on fertilizer application.

Research Question: What is the relationship between harvest yields for produce where minimal fertilization is applied, versus acreage where greater levels of fertilizer is used?

Research Hypothesis: There is a statistically significant difference in the mean volume of plant matter harvested for fields treated with exceptional levels of fertilizer versus those fields where only a minimum quantity of fertilizer is applied.

Null Hypothesis: There is no statistically significant difference in the mean volume of plant matter harvested between acreage treated with minimal levels of fertilizer versus acreage that receives excessive treatments.

Probability Statement: To facilitate this empirical analysis, a 1000 acre farm was used for the experiment. Half of the acreage received the normal level of fertilization, while the other half received twice the quantity of fertilization. Random samples were then drawn from each half comprised of 24 acres that received the double application of nitrogen fertilizer, and 18 acres that received the normal application of annual fertilization. Therefore, $df =$

$N1 + N2 - 2$ or $24 + 18 - 2 = 40$. For a two tailed test, a Student's T ratio of plus or minus 2.021 is needed for significance at the .05 level, and plus or minus 2.704 is needed for significance at the .01 level.

Presentation of the Data:

<u>Means and Standard Deviations for Crop Yield</u>			
	N	M	S
Double Fertilizer	24	42.9	14.2
Normal Fertilizer	18	33.5	13.8

Student's T ratio = 2.10, $P < .05$

Examination of the Null Hypothesis: The NH is rejected at the .05 level of significance, $P < .05$

Interpretation: As indicated by the data, the mean annual harvest yield for those acres that received twice the quantity of nitrogen fertilizer was computed at 42.9 bushels per acre, versus a mean annual yield of 33.5 bushels for acres that received only the minimal prescribed coverage. Predicated on these values, a Student's T ratio of 2.10 was computed which indicates that there is a statistically significant difference in the mean quantity of plant matter harvested between the two sample categories. Based on this computation, it appears evident that fertilizer quantity is a critical factor in

determining plant growth and harvest yields. The premise that plant growth is significantly different predicated on the amount of fertilizer applied is therefore valid and we can reject the null hypothesis and accept the research hypothesis. Based on this information it appears safe to conclude that fertilizer is one of the variables that could potentially exert influence over crop yield. More study will be required to discern exactly what degree of influence fertilization has on plant growth and which other variables might be within such an equation, but based on the analysis conducted, it is safe to infer that there is a statistically significant difference between crop yields based on the amount of fertilization applied.

As you can see from this particular example, the variable (fertilizer quantity) did turn out to have significant influence in determining crop yield. The study split a large farm into two sections, and then used two different forms of fertilizer application. This was followed by a random sample drawn from each section. The means and standard deviations for crop yield measured in average bushels per acre were calculated and the results were compared. The analysis indicated that there is a difference in annual crop yield levels and that the difference is so large that it could not have happened by mere chance. The T ratio verified this discovery (2.10) which was more than required for the .05 level (2.021) and subsequently we are safe in concluding that fertilization levels do play a role in determining crop yield. See how easy that was? Nothing to it.

Let's try it again using an example from medicine. In this example the exact same evaluative process will be employed, along with a similar research design. The only difference will be the variable being examined. Now keep in mind that nothing in the world is univariate. Never have I seen just one factor alone account for all the variability in an equation. Consequently, when we conduct these types of analyses, we never depart from a perspective that endeavors to prove our point and we never presume that the factor we are evaluating is the only influential variable in the equation. Rather, we step back, reflect on all of the variables that might exert influence, and then measure each one individually for its univariate relation. After each factor has been examined, we would be safe in including it within a multivariate correlation study to determine the proportional degree of influence it exerts over the dependent variable. We will talk a great length about this correlation process in the later chapters, but it is important to remember that, at this point, we are concerned with proving the premise and not in determining predictive strength of the variable.

To set the stage for our medical example using Student's T ratio, let us assume that we are desirous of assessing the impact that a specific inoculation has on mitigating future incidents of patient contraction of a virus. Clearly in such a study we would presume that there are a significant number of factors that combine to increase or decrease a person's likelihood of contracting a virus. Accordingly, we would examine the literature, synthesize the postulated

variables, and then compose a theoretical equation. Our goal would be to lay out all of the possible variables that could increase a person's susceptibility to the virus and then determine what proportional level of influence each factor has on increasing one's likelihood of contracting the disease. Before we get to that step however, we need to prove the truth of each premise individually before examining the correlations. Since one variable seldom accounts for all of the variability in an equation and since there are a relatively few examples (outside of the Polio Vaccine) where inoculation has a one hundred percent favorable affect on eliminating contraction of the disease, it is necessary to assess all of the potential factors that could combine to influence a patient's susceptibility.

To accomplish this we might assert that sufficient presumptive evidence exists based on our review of the literature and through our observations to theorize that the following factors (hypothetically) have may an influence in whether or not someone contracts the A234 virus. What our goal would be by examining these factors is to rule in or out each variable as a partially contributive factor. So essentially, inoculation would not be the only variable that affects reoccurrence of the illness, but other factors combine with inoculation to increase the probability of avoiding future contraction of the virus. Accordingly our research design would be conducted from a before and after perspective where all patient's included within the study were initially treated for contraction of the virus and inoculated. One year later, two samples are

drawn to measure whether each of the factors below were influential in determining instances of reoccurrence.

- 1) Diet
- 2) Body Weight
- 3) State of Health
- 4) Frequency of Hand Washing
- 5) Age

To conduct our multivariate study, each of these factors would have to be examined (individually) to see whether there is a statistically significant difference between patients who did reacquire the virus, versus those who did not contract the virus again, within the year. For our example we will select the Body Weight variable and analyze whether there is a statistically significant difference in the mean body weight for those who contracted the disease a second time and those who did not require the virus. You will notice that we are not offering an explanation yet as to why body weight might be a factor, only that our observation has been that slightly built people seem to be more prone to re-contracting the virus than those who weigh more.

Because of the before and after design of this analysis, we would use a one tailed Student's T ratio due to the fact there is no chance of a negative T value being computed.

Problem Statement: Determining the affect that body weight has on contracting Virus A234.

Research Question: What is relationship between average body weight and susceptibility to infection by Virus A234.

Research Hypothesis: There is a statistically significant difference between the mean body weights of people who contracted A234 versus those who were equally exposed to the virus yet did not contract the disease.

Null Hypothesis: There is no statistically significant difference in the mean body weights for these two groups of people.

Probability Statement: To facilitate this empirical analysis two groups of people were identified and a random sample of twelve participants from each group, were selected for inclusion within the study. The mean body weight and standard deviation were computed for each group and then compared using Student's T ratio. Since $df = N_1 + n_2 - 2$ or $12 + 12 - 2 = 22$, then a Student's T ratio of 2.07 is required for significance at the .01 level. A T ratio of 2.82 is needed for significance at the .01 level.

Presentation of the Data:

<u>Means and Standard Deviations of Body Weight</u>			
	N	M	S
Re-infected Group	12	113.4	25.4
Not Re-Infected	12	103.3	32.1

Student’s T ratio = .85, $P > .05$

Examination of the Null Hypothesis: The NH is accepted at the .05 level of significance, $P > .05$

Interpretation: As indicated in the table, although there is a difference in the mean body weight between the two groups, the difference is not so significant that it could not have been associated with mere chance. Consequently, the study failed to provide conclusive evidence that body weight is a potentially influential factor in determining whether or not someone who was exposed and contracted the A234 virus would like become susceptible again. Based on this finding we must accept the Null Hypothesis which suggests that there is no statistically significant difference between the mean body weight of the two groups of patients and their susceptibility to reacquiring the A234 virus.

As you can see from this example, we were not able to provide conclusive evidence that a “significant” difference existed between the mean body weight and susceptibility to the virus. Even though both groups received the inoculation, the mean weight difference between those included was only about ten pounds. What is interesting to note was that the standard deviations for the two groups was 25.4 pounds for the group that reacquired the disease versus 32.1 pounds for the group that did not reacquire the disease. Since Student’s T ratio calculations examine not only the difference between means but also take into account the standard deviations, this could partially account for the Student’s T ratio of .85, which failed to exceed the minimum value necessary for rejecting the null hypothesis (2.07). Consequently we are forced to accept the null hypothesis and exclude body weight as a potentially contributive variable in determining susceptibility to the A234 virus.

The important point here is not the outcome, but the fact that we employed a scientific protocol to the process of determining which factors contribute or not to disease susceptibility. Discovering what does not matter in such an equation is every bit as important to identifying those variable that do contribute to disease. Such knowledge helps us not only reject false claims or conjecture, but serves our ambition of focusing on only those factors that have a demonstrably influence. Ruling out a variable helps this process and as I mentioned previously, the process of elimination is an important component of the logic

function. No one should expect that we should know (in advance of our research) all of the right answers. Rather, it is a process of elimination that is facilitated by thoughtful scientific reasoning that differentiates our efforts at isolating the truth that puts us in a different category from those who would simply rely on speculation and guess work. This scientific approach to hypothesize, collect data, formulate hypothetical assertions, and then test the results is what puts us in the position of knowing for certain, as opposed to guessing and hoping that we are right.

Chapter 8

Proving a Premise with Chi Square

As we learned in the previous lecture, inferential statistics is the basis for much of the scientific thinking and reasoning you have been reading about thus far. This is especially true with regard to **proving** a premise before including it within an argument.

Inferential statistics involve randomly drawing samples from populations, and making inferences about the total population by examining the summarized data contained within the sample. One very effective technique involves drawing two samples from the respective populations, then splitting the sample into groups, and computing the frequency of responses within separate categories. If the difference is statistically significant then the inference is made that there is a significant difference between the frequencies of not only the sample, but probably the total population as well. In other words, we "infer" the findings of the sample to the total population based on the probability.

A critical frequency test (Chi-Square) can be used to determine if the difference between the frequencies of two groups, relative to two or more responses, is statistically significant, or not. As applied to critical thinking, before we state in an argument that a person's membership in a gang is a factor in committing criminal behavior, we should test the hypothesis. Although it makes sense that such a presumption may be true, we do

not really know for sure because we have no empirical evidence to support our premise.

If we find that gang membership is a significant contributor to criminal propensity within a randomly drawn sample, and the level of statistical significance is beyond that needed to reject the null hypothesis, then we can safely infer the results to the total population and thus accurately assert that gang affiliation is a factor in criminality. The manner in which we would collect data is similar to the Z ratio sampling methodology, but instead of collecting mean and standard deviation data, we would bifurcate the sample into two groups (gang members and non-gang members) and then assess their past criminal history.

If the Chi-Square was computed beyond a level needed for statistically significant difference (3.83 at the .05 level or 6.64 at the .01 level)), then we could conclude safely that the same relationship between gang affiliation observed and criminal propensity observed in the sample probably exists in the total population as well. Thus, our assertion that gang members have a higher propensity for engaging in criminal behavior would be a safe presumption.

As I mentioned in the prior lecture, statistical significance (based on the laws of probability) is a conventional way of stating whether a difference between groups or a relationship between variables has occurred simply by chance, or not. When comparing frequencies of two samples, a Chi-Square helps the researcher decide whether the difference would be expected to occur by chance, or whether it would not be expected to happen by chance; whether the difference may be attributed to

random sampling error, or whether the difference between the frequencies was so large that it overcame expected sampling error. The .05 and .01 levels have been accepted by the scientific community as the standard measures of acceptability for such measurement.

Chi-Square analysis differs from Z ratios in that Chi-Square does not use a standard measure. Instead, it uses a changing scale of needed values for the .05 and .01 level so of significance based on the degrees of freedom. These degrees of freedom are computed based on the number of rows and columns (or variables) within the analysis. A two by two table (two vertical categories and two horizontal categories) requires 3.83 and 6.64 at the .05 and .01 respectively. In other words if a Chi-Square is computed at 3.85, then you would reject the null hypothesis at the .05 level which gives you 95% certainty that the same or similar difference between frequency measured in the sample would apply to the total population. Subsequently, you could safely assume that criminality is associated with gang membership universally and contributes (partially) to criminal behavior.

As with the Z ratio test, we need to conduct an empirical test to assure ourselves that the premise (that gang membership is a contributive factor to criminality) is correct. We accomplish this by incorporating a quantitative analysis that examines the frequencies of two samples (those who have and those who have not committed criminal acts based on whether they are or are not a member of a gang) to see if there is a significant difference between the frequency of gang affiliation and criminality. To be absolutely sure that sampling error was not the cause of the difference we noted between the two

groups, we accept the null hypothesis (which would state There is NO statistically significant difference between gang affiliation for those who commit crimes and those who do not) until the Chi-Square proves beyond reasonable doubt that such a difference exists. The .05 and .01 levels of significance correspond to the 95% and 99% (or plus 2 and plus 3) standard deviations on a normal distribution bell curve. Chi-Square however is different from Z ratio in that it only uses a one-tailed curve. The reason for this is that there cannot be a negative number.

If a difference between frequency of criminal involvement produces a Chi-Square that is large enough (plus or minus 3.83) so that it would be expected to occur by chance in less than 5% of the cases (.05 or 5 times out of 100), then the difference between frequencies is said to be significant at the 5% or (.05) level. If we see a Chi-Square bigger than 3.83 we can safely "reject" the null hypothesis and accept the research hypothesis, which says, "there IS a significant difference between the two groups relative to gang affiliation and criminality", subsequently we can also safely presume that gang involvement is ONE of the factors that influences criminality.

Chi-Square is a statistical test that can be used to examine the difference between frequencies of samples drawn from populations. It is considered a critical frequency test for samples. The larger the samples the greater the degree of representation or inclusion of an adequate cross section of the population and the less likely the sample was not representative of the total population. Chi-square assumes normality of the distribution. Areas under the normal curve may be examined to determine, within that level of

probability, the difference observed between means would have occurred by chance.

As we have learned, the presentation format for any research design, even where we are merely trying to prove the premise within an argument, requires structure, clarification, and precision. The scientific protocol for such formats is as follows:

The Influence of Gangs on Criminality

Problem: The growing issue of gang affiliation and its influence over people to commit crime.

Research Question: What relationship exists between gang affiliation and criminality?

Research Hypothesis: There is a statistically significant relationship between membership in a street gang and the probability that persons within that gang will commit criminal offenses.

Null hypothesis: There is no statistically significant difference between gang affiliation and propensity to engage in criminal behavior.

Probability Statement: Sixty randomly selected participants were included with the study. Thirty participants were selected based on their gang affiliation and thirty were selected based on their non-involvement in gangs. For a 2 x 2 table a Chi-Square of 3.66 is need for significant at the .05 level and 6.64 for significance at the .01 level. An analysis of past criminal history was conducted which yielded the following results.

Presentation of the Data:

Category	Criminal History	No History
Gang Members	25	5
Non Gang Members	5	25

Examination of the Null Hypothesis: Based on the information collected, a Chi-Square of 15.50 was computed. The null hypothesis was rejected at the .01 level of significance; $P<.01$

Interpretations: Based on the extremely large Chi-Square computed in this analysis, the null hypothesis was rejected at the .01 level and the Research Hypothesis was accepted, indicating that in fact there is a statistically significant difference between criminal propensities for members of a gang, versus those who do not hold such an affiliation. Predicated on these results, the premise that gang affiliation influences criminality is confirmed.

Proving a premise before including it within an argument is a requirement if you stand any chance at all of being accurate. To say that gang affiliation matters in criminality might be an interesting presumption, but certainty requires the arguer to take the time to establish the research protocol that tests the accuracy of the statement. Proof is essential before creating policy, or determining the type of reaction a society should employ. There may be (and usually are) many variables that influence the outcome of things. Gang affiliation might be one of 50 variables that contribute to criminality. The Chi-Square lets

you prove the truth of the premise based on a comparison of frequencies versus means. Such results may in fact confirm your premise that gang participation contributes to criminality, but does not disclose what the other remaining 49 variables might be. To identify them all you have to study the problem, identify other potential influences, test each one of them to discern whether they have any contributive value, and from this, you can formulate a comprehensive argument about the factors affecting criminality.

As you can see the use of Chi-Square as an inferential tool to measure for statistically significance differences between the “frequency observed” and the “frequency expected” is a straight forward process that results is definitive measure. The use of this computation to test the proof of a premise, as expressed in the hypothesis statement, is particularly useful.

Let’s look at one final example so that you can see how to use this technique as it applies to an application that doesn’t try to compare two differing groups, but rather that uses the statistic to examine categorical placement of data. This example is a one-variable case that measures goodness of fit.

Examination of Transportation Patterns

Problem: Given the plethora of transportation options available to commuters in modern day society, it appears safe to presume that there would be a balance between the options selected by commuters.

Research Question: What mode of motor vehicle travel do commuters use to get to work?

Research Hypothesis: There is a statistically significant difference in the frequency of observations relative to the types of motor vehicle transportation used to get to work between by members of the community.

Null Hypothesis: There is no statistically significant difference in frequency.

Probability Statement: Since four alternative types of vehicle transportation options are available with the City of Serenity for commuting, degrees of freedom would equal the # Columns – 1, which would equate to $4 - 1$, which equal $df = 3$. For $df = 3$, a Chi-Square of 7.82 is needed for significance at the .05 level, and 11.43 at the .01 level.

Presentation of the Data: Responses of 60 randomly selected commuters to the question: How do you get to work?

	Alone In Car	Car Pool	Public Transportation	Other
Fo	25	15	10	10
Fe	15	15	15	15

Examination of the Null Hypothesis: Based on the information collected, a Chi-Square of 10.0 was computed. The null hypothesis was rejected at the .05 level of significance; $P < .05$

Interpretations: As evidenced by the data contained within the table, more commuters travel to work alone in their cars than was expected based on the laws of probability, while fewer commuters than expected used public transportation or other (alternative) forms of transportation. This frequency distribution would calculate to a Chi-Square of 10.0 and accordingly would provide statistically significant differences between the frequency expected to occur from the frequency observed, thereby meriting rejection of the Null Hypothesis and acceptance of the Research Hypothesis. Essentially, this confirmation would place us in a position to prove the truth of the premise that asserts that commuters have preferences in modes of transportation or stated

differently, there is a statistically significant difference between individuals and their preference for the method of transportation they elect to employ to commute to and from work.

Summation

As you can clearly see through the examples provided in this chapter regarding the use of Chi-Square, it is possible to prove the truth of a premise using “frequency” information. This process can be applied in a univariate argument, or the methodology can be used to provide proof of a premise in a multivariate argument. The point here is that there are a plethora of statistical techniques available for providing statistical proof of a premise prior to including it within an argument. Chi-Square serves the need to conduct such a test where frequency data is available, as opposed to means, medians, or modes.

Chapter 9

Proving a Premise with the Correlation Coefficient

I must confess that my all time favorite statistical test is the Pearson's Product Moment Correlation Coefficient. Besides having the coolest name in the arsenal of statistical measures, this one test provides a wealth of applicational benefits to the sciences and social sciences. Pearson's R (as it is commonly referred) is used to measure the correlation that exists between two variables. I should probably clarify a bit that by saying also that correlation measures the degree of association between things. Did that help? Perhaps not.

Expressed more simply, Pearson's R is used to determine whether two variables (such as Unemployment and Crime) are related in their movement up or down, using whichever quantifying approach you choose to apply. It doesn't attempt to provide a causal association between two factors, such as saying that unemployment causes crime, but simply measures the association that exists between these two variables. Correlation analysis is the hallmark of scientific research and is used for a variety of purposes, but it is essentially applied to research endeavors that attempt to measure what degree of relationship exists between two distinct variables.

Because of the flexibility of this statistical test, it can be used to measure statistical significance or strength and direction, the amount of shared variance accounted for, and its value in terms of prediction. It is a ratio of how things vary together, divided by how they vary separately. The correlation coefficient is a derivative computation that normally ranges from -1.00 through 0.00 to +1.00.

The closer the Pearson's R is to either -1.00 or + 1.00, the stronger the correlation between the two variables examined. If the correlation is -1.00 then the correlation between the variables is said to be perfectly inverse. This (essentially) means that each and every time one variable increased in value, the other variable decreased. A positive correlation of +1.00 would suggest just the opposite relation and indicate that every time variable X increased in value, variable Y also increased. Normally such measures are done over time or they can be predicated on measures in differing geographic locations, or perhaps they are products of measures where one variable was manipulated to see what consequence it had on the other variable. Notice that I cleverly avoided the use of the word effect in that sentence? I'll explain later.

If a Pearson's R was calculated at 0.00, it would be safe to conclude that no relationship whatsoever exists between the two variables. As you increase the value of variable X, variable Y might go up one time and down the next time. Essentially, in such a case, no relationship exists between the two variables at all and no inference can be made.

The research hypothesis for such an analysis would be somewhat different than those we've been using thus far in the book and rather than asserting that there is a statistically significant "difference" between variables or groups, we need to craft the statement to express that there is a statistically significant "relationship" between the variables included within our analysis. Actually this makes a great deal of sense, because relationship is what Pearson's R is measuring, so subsequently our hypothesis statements should reflect that measure applied to the examination. In addition to providing an outstanding measurement of association, correlations also provide us with an excellent platform upon which to base our conclusions in an argument.

Consider for a moment that up until this point we have been engaged in crafting premise statements that we endeavor to prove using tests for statistically significant differences between means and frequencies and now we discover a test that measures correlation or association between two variables. Pearson's R then can be applied [after] we test the truth of a premise to see if there is a difference between the two factors and to measure whether the factor itself has any association to something else. Let me provide an example. If we establish a premise in our argument that asserts that unemployment logically influences the level of crime in the city, we can first test that premise by using a mean or frequency statistic that measures (for example) the mean unemployment rate for a high crime city versus a low

crime city. If we discovered that there was indeed a statistically significant difference in the mean unemployment rates between these two categories, we could reject the null hypothesis and accept the research hypothesis. This (essentially) would offer proof of our premise that unemployment has an influence as applied to the instances of crime. The next logical step would be to prove the truth of the argument by employing Pearson's R as a measure of correlation for crime and unemployment. In this example we would not be comparing two distinctly different geographic regions, but rather we might collect data for unemployment (X) and crime rate (Y) over time for one specific city. Once we have a representative sample, we can compute the correlation coefficient between these two variables to measure whether or not crime increases and decreases as unemployment fluctuates. If we detect a strong positive correlation (.89 for example) between these two variables, then we could say that there is indeed a "relationship" between the two factors.

Obviously, this example is univariate in nature and uses only one dependent variable {crime} and only one independent variable {unemployment} to explain the process simplistically, but I think you probably get the idea of how the correlation coefficient can be used to support analysis of the conclusion in an argument. It can also be used to examine the truth of premise within an argument, should you determine the correlation is the best measure.

Examining Statistical Significance Using Correlations

Pearson's R is similar to other statistical tests for significance in that an analysis of the significance is conducted using the .05 and .01 levels. The correlation needed for significance is computed by calculating the degrees of freedom in a given study and then consulting a table of values (see the appendix).

The formula for computing the degrees of freedom is based on the fact that there are always two pairs of observations and is expressed as $df = N - 2$. The N value is the number of pairs of data in the sample and the minus 2 is applied because of the two variables within the sample (i.e., unemployment and crime). If there were 10 sets of data for both variables in the sample then the degrees of freedom would be $10 - 2$ or $df = 8$. With eight degrees of freedom the chart would prescribe that a Pearson's R of .707 would be required for statistical significance at the .05 level, and .834 at the .01 level. So, to reject the null hypothesis that expresses that no statistically significant correlation exists between variable X and variable Y , you would need a Pearson's R of .707. If the Pearson's R was equal to or greater than .834, you could safely reject the null hypothesis at the .01 level, meaning that if you did such a test 100 times, you would expect that in 99 out of the 100 times, a Pearson's R of .834 or greater would (according to probability) be generated based on the analysis of the data.

That is a pretty powerful conclusion when you stop and think about it which places the researcher in a position where they can, with confidence, assert that the relationship was so strong that it has significant implications to not just the study conducted, but also to other similar studies of the same variables in different geographic regions. Naturally, there isn't any validation of this assertion yet, so you'd want to qualify your exuberance, but you get the idea. Using Pearson's R we have a tool that not only can be used to test the truth of a premise, but a tool that can be applied to assessing the truth of the conclusion of our argument. That's a pretty important discovery to make, don't you think?

I'm not planning on spending allot of time explaining the math formulas or the analytical procedure here, but rather it is my goal to not overwhelm you with methodology or procedure, so you can focus on how to apply the results of such computations to your analysis. This doesn't mean that you shouldn't take the initiative to thoroughly study the mathematics behind correlation coefficients, but this isn't a statistics book. I recommend that you spend some time refreshing yourself in the mathematical procedures associated with all of the statistics we've explored thus far, but particularly the Pearson's R so that you can spot problems in the computations and results. It isn't hard to learn this procedure and in fact most software packages have a built in correlation coefficient algorithm, but to be sure that it was done correctly, it helps to have a masterful grasp of the mathematics.

Magnitude and Direction:

Any correlation coefficient should be interpreted only within the context of the particular study where it was generated. With that said, we violate this rule all the time because we are engaged in “inferential” analysis and at some point we find ourselves in the position of extrapolating the findings of our studies. To help you understand the process of interpreting correlation coefficients, the information below is provided to assist you.

Guide for Interpreting the Pearson’s R Value

.00 to .20	Weak
.21 to .40	Weak to Moderate
.41 to .60	Moderate
.61 to .80	Moderate to Strong
.81 to .90	Strong
.91 to .99	Very Strong
1.0	Perfect (which usually means your data is wrong or you’ve made a computation error)

This guide is applicable to both positive and negative correlations, but a negative correlation coefficient indicates that an inverse relationship exists.

Coefficient of Determination:

The correlation coefficient, as you've seen, is an extremely powerful statistic that has utility in assessing the degree of association between two variables and can be used within a standard research protocol to assess whether or not the RH or NH is accurate. By itself that is an important accomplishment, but many find the coefficient a bit nebulous when it comes to extending the interpretation much beyond accepting or rejecting a hypothesis statement. The good news is that the Coefficient of Determination is also available and can be used to discern the percentage of time that the two variables under study vary separately and in unison.

The Coefficient of Determination is computed by simply squaring the Pearson's R value. For example a Pearson's R of .80 would compute to an R squared value of .64. If you multiple .64 x 100, you would discover that variable X and variable Y move up and down together 64% of the time. In other words, as Unemployment (X) moves up or down, Crime (Y) moves in unison 64% of the time. This can also be expressed a bit differently by suggesting that unemployment constitutes 64% of the variability in the crime equation. Although not immediately obvious, that's a pretty important statement, especially when applied to an effort that endeavors to extend the measure of correlation to the prediction of future values. Essentially, by knowing the correlation coefficient and then computing the coefficient of determination, you can draw inference

about the degree of dependency between the variables, which may be of further value in a predictive equation. If, for example, you account for all of the variability between independent variables and the dependent variable, you stand a much better chance of predicting the future value of Y given changes in all the X variables.

In the example just provided, we used .80 as the Pearson's R value, which computed to $.64 \times 100$ or 64% for the Coefficient of Determination. Conversely, there must be 36% of the variability in the dependent variable that isn't accounted for yet. This is called the Coefficient of Non-Determination or [k]. You compute this value by simply subtracting R squared from 1 (or $1 - R^2$) $\times 100$. In our example, $1 - .64 = .36 \times 100$, or $k = 36\%$. This means that there is still 36% of the variability out there that we haven't discovered or cannot be accounted for by the one independent variable included within our correlation analysis.

Remember that I mentioned earlier about the multivariate argument? This is exactly what I was talking about. Almost never, do we as scientists encounter a perfect correlation between two variables. Accordingly, we cannot conclude that only one thing influences another thing in totality. It's a multivariate equation and as such, it requires us to identify ALL of the hypothetical variables [premises] that might assert influence over the value of the dependent variable [conclusion] before we can assert with a degree of confidence about the answer to the question [WHY].

As applied to argument analysis, again almost never can we assert that only one premise is responsible for the conclusion, because the world essentially isn't a simple place and the dynamics of the world we live in are not easily discovered. Speculation gets us into trouble. So does over simplification of the complexities of the universe. Almost never do we encounter perfect correlations, which would confirm that one variable accounts for all of the variability in an equation between two variables and if we do, we probably made a mathematics error or we forgot to select the correct set of data for the X and Y variables so that the computer could compute the Pearson's R statistic. I've actually made that mistake myself several times, which is a startling revelation, which is soon dashed. This isn't to say that we won't come close to finding correlations that equal 1.0, but rarely will we see perfect correlations in our study of the world with just one independent variable.

Let me now provide an example that will help you understand how the correlation process works. To make things simple, we'll continue with our unemployment (X) versus crime (Y) example. As you will see in the research design, we use the same format that we have for all the previous examples in the book. The only difference is that this time we are looking for statistically significant relationships, instead of differences.

Example of Pearson's R Using Unemployment and Crime Rate

Problem: Determining the relationship that exists between unemployment and crime.

Research Question: What relationship exists between unemployment and crime rates in the State of Serenity?

Research Hypothesis: There is a statistically significant relationship between unemployment rate and the crime rate experienced, per 100,000 residents, in the State of Serenity from 2000 to 2009.

Null Hypothesis: There is no statistically significant relationship.

Probability Statement: Data were recorded for the years 2000 to 2009 for both unemployment rate and crimes reported per 100,000. Since ten years of data were included within the study for two distinctly different variables, $df = 10 - 2$ or $df = 8$. A Pearson's R of plus or minus .707 is needed for significance at the .05 level, and plus or minus .834 is required for significance at the .01 level.

Presentation of the Data:

Year	X Unemployment Rate	Y Crime Rate per 100,000
2000	3.5	297.6
2001	4.9	348.4
2002	5.9	385.9
2003	5.6	374.6
2004	4.9	382.7
2005	5.6	441.3
2006	8.5	465.0
2007	7.7	420.2
2008	7.0	404.9
2009	6.0	417.0

Associated Computations:

N = 10

Sum of X = 59.6

Sum of Y = 3937.6

Mean of X = 5.96

Mean of Y = 393.8

Sum of X Squared = 19.6

Sum of Y Squared = 20551.5

Sum of x times y = 532.7

Standard Deviation of X = 1.4

Standard Deviation of Y = 45.3

Pearson's R = .84

R Squared = .72

K = .29

Examination of the Null Hypothesis: The NH is rejected at the .01 level of significance; $P < .01$

Interpretations: The very strong positive correlation coefficient of .84 indicates that, as unemployment rate increases in the State of Serenity, the crime rate per 100,000 residents also increases. As reflected by the Coefficient of Determination of .71, these two variables vary together 71% of the time, and are independent of each other 29% of the time ($k = .29$). The findings suggest that when people find themselves unemployed, stimulation occurs in crime rate, which supports the economic theory of crime. If unemployment rate appears to be increasing, criminal justice planners might well predict a corresponding increase in crime rate.

As you can clearly see from this example, unemployment as a dependent variable might serve as an effective predictor of future crime trends, given the exceptionally high correlation that exists between the two variables. It isn't perfect, nor should it be given the multivariate nature of things in the world, but it does provide a good start in identifying the influential factors associated with crime. Just out of curiosity, what level of strength did that correlation coefficient possess? Did you remember to use the table provided earlier in the chapter to determine the relative strength of the coefficient?

Unlike the other forms of statistical analysis, correlation coefficients possess a much higher degree of utility. They not only explain the degree of association between two

variables, but they can also be used within a regression equation. Obviously understanding the dynamics at work is an important part of the scientific process, but certainly we don't want to be relegated to mere spectators. Having a comprehensive understanding of the dynamics at work is, by itself, fascinating but using this knowledge within a predictive equation that can forecast likely outcomes given changes in the value of the independent variable is an added benefit. With such information and capability we can not only answer the question why, but we can generate a pretty fair estimate of what is likely to happen should the predictor variable rise or fall in value. Such ability has application to science, economics, criminology, engineering, and virtually every form of human endeavor. The good news, it is hard at all.

Linear regression provides for the prediction of a value of one variable from the value of the other variable and can be used to predict Y based on the estimated value of X, and vice versa. The **Independent Variable** or predictor variable is referred to as X, while the **Dependent Variable** is given the designation of Y. Using a regression equation, values are substituted from the correlation analysis and then a hypothesized value is assigned to the independent variable to facilitate the calculation of the dependent variable. That probably hurt a little, but you'll get the idea after you review the next example.

Using the Pearson's R derived from our Unemployment (X) and Crime Rate (Y) example, we can compute the two necessary values required for the regression equation. The first value would be the regression coefficient. We designate this point as [b] and the formula to compute the value is the Standard Deviation of Y, divided by the Standard Deviation of X, multiplied by R (.84). As you will recall, those values were 45.3 and 1.4. So, first we simply divide 45.3 by 1.4 and derive 32.35, then multiply that product by .84 for our [b] regression coefficient which equals 27.1.

Next we need to compute the slope intercept or the [a] part of the formula. This is an equally simple calculation that uses the Mean of Y, minus the Mean of X, multiplied times the slope intercept or [b] value that we computed previously. In our example we would substitute the values already computed for the correlation coefficient and insert them within the formula for deriving the regression coefficient. The Mean of Y we calculated was 393.8, while the Mean of X was 5.96, and [b] we just discovered was 27.1. That's all the data we need to compute the regression coefficient.

$$\begin{aligned} a &= 393.8 - (5.96) \times (27.1) \text{ or} \\ &393.8 - 161.5 \text{ or} \\ a &= 232.3 \end{aligned}$$

Now that we have the two iterative computations necessary to complete the forecast, we can build our regression equation, which is represented as;

$$Y' = a + bX_1 \text{ or}$$
$$Y' = 232.3 + 27.1 X_1$$

This may look a bit confusing because of X_1 symbol, but it simply means whatever number you want to put in there to represent the hypothesized value of the independent variable. Since our example involves predicting crime from fluctuating unemployment, we simply insert a probable unemployment rate that could occur next year. This value might come from the economics community or be based on government projections.

For this example, if we used a projected unemployment rate of 7%, we would see that our formula would look like this;

$$Y' = 232 + 27 (7) \text{ or}$$
$$Y' = 232 + 189 \text{ or}$$

$$\mathbf{Y' = 421}$$

The projected crime rate per 100,000 people for the State of Serenity would compute to 421, which is fairly close to the historical trend that we have seen inside our data table that we used to compute the correlation coefficient.

Let's recap for minute. We just used correlation analysis to test for significant relationships between two variables (unemployment and crime rate), and discovered that there was indeed a strong positive correlation between the two variables. This proved the premise that there was a relationship and then we used the data and the derivative computations, along with two additional computations to derive a prediction equation that allowed us to forecast the future crime rate, based on an anticipated change in the unemployment rate for 2010. Do you feel empowered now? Stop and think about it. You not only proved a relationship, but you predicted the future based on a theorized change in the independent variable, twelve months before it actually happens. That's pretty cool stuff.

If you apply it further to an argument, you'll see that you not only proved the premise, but proved the conclusion as well. That is exactly what we have been talking about throughout this entire book. Not guessing that you're right, but proving that you are correct through the application of statistical methods to test the truth of the premise before we insert it within an argument, and then testing the truth of the conclusion in the same manner by using correlations. The same technique can be applied to vertical and perpendicular logic equations to provide an estimate of each variable. I suspect that the light bulb just went on over your head and you now see how all this comes together in determining the answer to the question of WHY.

Chapter 10

Multivariate Correlation and Discriminant Function Analysis

In my opinion, the most powerful analytical tool available to researchers in support of multivariate-multidimensional analysis is a methodology entitled Discriminant Function Analysis. Available in most high-end statistical software, discriminant analysis combines a series of related techniques and computations into a single methodology that affords the researcher with a holistic view of each variable, as well as an aggregate assessment of the individual and contributive value of the equation needed to determine the probability of an outcome.

Developed principally to aid in the differentiation of variables to determining group association, discriminant analysis uses six distinct statistical measures that aid in the evaluation of centroidal separation between groups, the collective strength of predictor variables, the degree of residential discrimination of the predictor variables in estimating group placement beyond the sample elements, verifications of predictor accuracy, the relevant discriminating power of each predictor variable, and finally a regression equation that is useful in computing the probability of group association.

Discriminant analysis can be applied to all types of evaluations. The methodology can be used to determine horizontal influence of major variables in determining the probability of group association or it can be used to assess the individual and collective strength of vertical and perpendicular equations. The process itself is not hard to accomplish, nor is the application of the findings to the process of elimination and the search for truth, but it does require some practice. DFA essentially allows you to lay out the logic of an argument, then quantify variables (you can even use dummy weightings for those things that don't lend themselves to a linear form of quantification such as yes or no, true or false, gender types, or positions on an issue), and then analyze the information to determine those factors that contribute to a decision, outcome, group affiliation, perspective, etc.

Discriminant function analysis is used to determine which variables discriminate between a particular behavior or phenomena such as in the case where a researcher seeks to investigate which variables discriminate between college graduates who decide to go on to graduate school, as opposed to those who elect not to pursue graduate level education. To support such an analysis using discriminant analysis, the researcher could collect data about numerous hypothetically related variables that might influence such a decision. Discriminant Analysis could then be used to isolate those variable(s) which appear to have the greatest (and least) degree predictive power over such a deliberation.

Discriminant analysis is also useful in medical research and can be used to differentiate predictor variables that hold contributive influence relative to the category of afflictions contracted by members of the community. An example might be where the researcher seeks to isolate those variables most influential in determining whether someone is likely to contract head and neck cancer. In such a research design, discriminant analysis can be employed to measure the influence of behavioral, genetic, environmental, and other factors relative to their individual and collective power in increasing the odds of whether one is likely to contract such a form of cancer or not. Like all other logic endeavors cited in this book, the researcher must conform to the requirements of avoiding presumptions and speculations about the factors included within the study until such time as the data support a conclusion. Literature reviews are used to derive scholarly observations and theoretical postulates that support the formation of a reasonable hypothesis. Variables are identified and quantified in such a manner to support the assessment of the truth or fallacy of the claim made within the research hypothesis. Based on the analysis of each individual premise in the scientific argument and the truth of each individual claim, the analysis of the argument process moves forward to the next step whether the aggregate influence of the equation is measured using discriminant analysis, finally culminating in the formation of a predictive equation that can be applied to not only the sample data, but to cases outside the sample.

The incorporation of discriminant analysis, as I mentioned previously, provides the researcher with the most powerful analytical capability within our arsenal for determining the statistical probability of an event. Underlying the utility of this evaluation methodology is the fact that discriminant analysis doesn't rely exclusively on linear data. Instead, the researcher has a greater degree of latitude in assigning values to each variable and case by assigning "dummy" variables.

Discriminant function analysis is used to classify cases into the values of a categorical dependent, usually a dichotomy. If discriminant function analysis is effective for a set of data, the classification table of correct and incorrect estimates will yield a high percentage correct.

Key Terms and Concepts

Discriminating variables: Independent or predictor variables. These are akin to the independent variables in our correlation example and represent those factors that influence the behavior of the dependent variable.

The criterion variable: This is the dependent variable or that outcome we are endeavoring to measure. It can be group association or an outcome depending upon the structure of the research. Typically in discriminant analysis, we group factors into two or more distinctly separate groupings and then test to see how the discriminating variables influence the outcome or placement.

The **Eigenvalue** of each discriminant function reflects the degree of group separation reflected within the study sample based on that variable. The Eigenvalues are associated with percents of variance explained in the dependent variable, cumulating to 100% for all discriminant functions. The greater the Eigenvalue, the farther apart the two groups are in theoretical space and the more effective the variable is at predicting separation.

Standardized discriminant coefficients illustrate the degree of relevant discriminating power possessed by the predictor variables used with the analysis. But they are limited in their ability to provide needed elements for the regression equation and consequently allow only for the assessment of future group classification.

Unstandardized discriminant coefficients are used in the discriminant or regression formula and serve as the multipliers against which the hypothesized values are compared.. Similar to the b coefficients that are used in regression equations in making predictions, the unstandardized discriminant coefficients are used as multipliers in the discriminant equation [b]. The product of the unstandardized coefficients, with the observations, yields the discriminant score.

Discriminant scores are a bit different than the standard multiple regression equations we are used to seeing insofar as the Y' is actually a z ratio, and not a linear product. A positive Z score for example would depict placement in one group, while a negative Z score would place a case in the opposite group. The greater the Z score,

the greater the likelihood is of group association (remembering that Z uses 1.96 and 2.58 as the standards for significance).

Let me provide an example so that you are better able to understand all of this mumbo jumbo. Let's presume that we were endeavoring to build a predictive equation using discriminant analysis to discern the probability of escape propensity for prisoners that can be used during screening. The resultant equation would be used as a classification tool that would take the guess work out of the classification process with regard to this possibility. It may not correctly predict each and every prisoner's behavior, but it would provide a scientific estimate based on past history of such events and the analysis of the factors influencing such behavior.

The variables we might include in such an analysis would be based on our review of the literature to discern those factors that possess influence in human behavior in such situations and we would then construct a sampling protocol that examines past cases of escapees, as well as data for those who did not elect to escape while in custody.

The variables we might choose to include within our study might include;

Height
Weight
Age
Marital Status
Education
Employment
Parole Status
Residency
Previous Escapes
Previous Arrests
Active Holds
Pending Court
Ethnicity
Confinement Period

You'll note that some, if not many, of these predictor variables aren't easily quantified using interval data and must be approached using a nominal or ordinal form of measure. This consideration actually affords a tremendous degree of flexibility in that the researcher can assign what are termed "dummy" values to each case. If we collected data for two hundred escapees and two hundred non-escapees, we could classify each as either a 0 or 1 (depending upon their past action) and then treat each predictor variable in the same manner. Marital status might result in (1 for yes, 2 for No). Parole status might be reflected as (1 for on parole at the time of arrest and 2 for not only parole). It is really important to remember that no

matter the number you assign, unfavorable should hold a lower value than the number you elect to use to apply to favorable. This makes it easier to interpret the data after you compute all the associated statistics involved in discriminant analysis. If you forget to use such a quantification process, then you will adversely affect the outcome because for one variable you will have assigned a positive value to represent the presence of status, while for the next variable you assigned a lesser value to show a positive status. This will mess up the mathematics of the process and you'll be left wondering why you didn't discover a significant discriminant score as a byproduct of your equation.

If all goes according to plan, and your analysis yields favorable results in determining the standardized and unstandardized discriminant functions of each variable, you might end up with a highly accurate predictive equation that can evaluate the values for each new inmate that requires classification and the concomitant ability to place that individual within a housing location that corresponds to the level of escape threat that they present. The formula might look something like this;

Discriminant Equation for Predicting Escape Propensity:

$$Y' = 8.47 + .137x_1 + .0079x_2 + .0173x_3 + .0273x_4 + .0065x_5 - 1.168x_6 + 1.049x_7 + 1.23x_8 - 1.268x_9 + .589x_{10}$$

Such an equation sets the stage for not only the evaluation of every member of the two groups of prisoners that you included within the study, but it can be applied to future arrivals in order to render a scientific determination of their escape propensity based on the variables deemed statistically significant. In other words, each of the variables or premises within the multivariate argument have been statistically tested to measure whether there was a significant difference between the means or frequencies of the members of the two groups and then once the truth of the premise was established, the multivariate argument that stated that all of these factors contribute individual influence and aggregate influence in the determination of whether or not someone is likely to present an escape risk was validated.

All you have to do is substitute the X values of any individual within the predictive equation and out pops a Z value that you can use to determine which of the two groups they will most likely fall in to. That's a pretty incredible thing to know and as you can see, all we did was follow the steps prescribed throughout this book in how to structure a multivariate argument, test each individual premises to discern whether or not it matters, and then combine all of the variables into an aggregate equation in order to see which was more influential than another, followed by the formulation of a comprehensive equation that combines all of the statistically significant factors into a predictive equation to forecast an expected outcome.

Chapter 11

Applied Spatial Modeling

Now that you are familiar with the basic conceptual theories and practical skills associated with multivariate-multidirectional reasoning it is time to examine the potential applications of these tools, along with Geographic Information Systems, as a mechanism for determining real world studies and subsequently using this knowledge to formulate effective strategic and tactical level policies. In fact, it is also advantageous at this point to discuss how this GIS, along with advanced empirical methods, can actually change the way in which you think about the world and the judgments you come to relative to these matters.

The inherent design of Geographic Information System software is extremely conducive to replicating the scientific approach to problem solving. Data are stored into separate tables according to some logical design structure and represent information in a manner, which considers attributes of value, space, and time. Similarly, the methods employed to represent geographical areas force the process toward a layered approach which necessarily subcategorizes each individual layer into its own unique file format, while all the time paying attention to the spatial integrity of each layer, when aggregated. When combined with the ability to develop horizontal and vertical level queries about data values relative to space and time, GIS systems clearly become the most pragmatic tool available for developing complex scientifically

oriented examinations. The key to using this tool to facilitate scientific examinations rests with the awareness, by the researcher, that things are not as simple as they may appear at first glance. The combined effects of multiple factors are usually responsible for fluctuations in the observed behavior of a dependent variable and the identification of these causative factors must be achieved if the researcher has any hope of isolating those factors which are, not only causative in nature, but which also can be controlled or manipulated toward achieving desired changes in the dependent variable.

To complicate matters, the more experienced researchers realize that they must also look for linkages between apparently dissimilar equations, which have a determinate degree of influence over one another, as well as a collective degree of influence to the equation under primary study. To illustrate this concept it is first necessary to establish a hypothetical situation that can be used for reference. Let us assume that we are engaged in the process of determining the answer to the following question.

How do we increase the number of Sea Otter colonies along the California coastline?

Our earlier discussions about the univariate mind set and its inappropriateness in dealing with the dilemma that we have created for ourselves here would probably yield a conclusion that involved the arbitrary relocation of breeding pairs of Sea Otters to other locations along the California coast as a solution. Such a policy decision would presume that by simply relocating a sufficient number of otters to other locations, humankind could satisfy its

desire to have a larger number of these cute furry creatures scattered along the coastline. Once relocated, the otters would continue to mate, subsequently producing offspring who would find other mates, and in short order, California would be awash in sea otters from one end to the other. The simplicity of this notion would be comical if it were not that this approach has already been tried. As you can imagine, the only thing that happened was that they ended up with a considerable number of (dead) Sea Otters.

A more appropriate methodological approach would have been to isolate those independent variables which exist within established Sea Otter colonies and which were suspected of having a demonstrative impact. Hypothetical inferences should then be made between the presence of these influences and the existence of the otter colony. Data relative to each of these variables is then collected, quantified, and stored for later examination. In the example used to teach you about multivariate correlation modeling, I identified several independent variables, which theoretically possessed some degree of influence over the population of the colony. For this example, we will limit the number of IndVars to three. X_1 = Kelp Bed Volume, X_2 = Otter Preferred Food Supplies, and X_3 = Sea Activity Levels for the area occupied by the otter colony. A straight forward multiple correlation and regression model would quantify the resultant formula as follows:

$$Y' = a + bX1 + bX2 + bX3$$

Where:

(Y') represents the quantity of Sea Otters within the colony

(a) represents the slope intercept of the multiple correlation equation

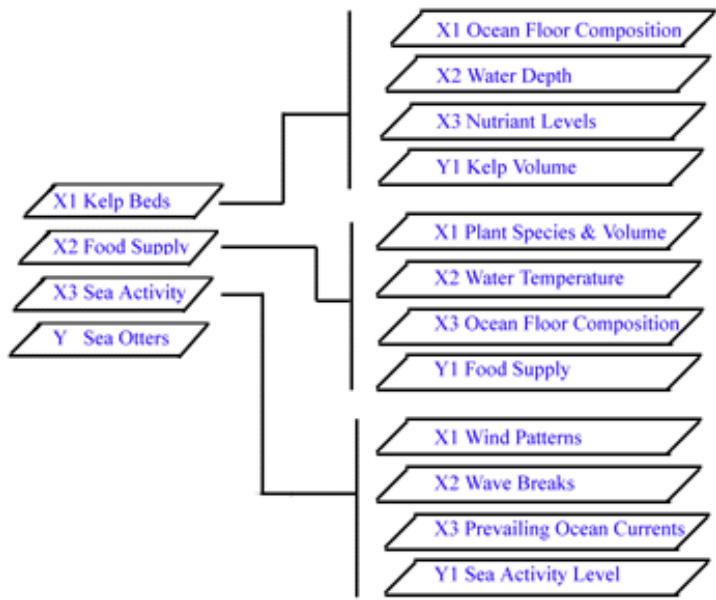
(bX1) represents the regression coefficient and multiplier for Kelp Bed Volume

(bX2) represents the regression coefficient and multiplier for Preferred Food Supplies

(bX3) represents the regression coefficient and multiplier for Sea Activity Level

In keeping within the theme of complexity in modeling, a more appropriate approach would be to use the capabilities of GIS to expand the model to its most fundamental components. This effort then facilitates the assessment of those IndVars might be first order influences upon the otter colony and in turn, which of them might be humanly controllable. In the graphic below, I have expanded the analysis to illustrate the collective empirical model, which relates directly to the otter population and also profiled three indirect models. These indirect models have relevance to isolating the most predominant influences on each of the IndVars used in the otter survivability equation. They can be used to conduct a micro-level examination of the primary influences over

kelp growth, otter food supply, and sea activity levels. In turn, each of the causative influences in these models can be assessed relative to their degree of controllability and subsequent suitability for manipulation.



As you can see, the hypothetical model used to predict kelp growth includes variables such as the composition of the ocean floor, water depth, and nutrient levels as its primary causative factors for determining the volume of kelp, which will grow within a geographic region. A great deal of theoretical support can be offered to support the contention that these phenomena exert some degree of influence, but for this example, let's just presume that strong positive correlation coefficients have been realized through traditional empirical analysis methods and that these figures have been validated through component file

aggregation and vertical comparisons. The second subset reflects that otter preferred food supplies (crustaceans & mollusks) are primarily dependent upon water temperature, plant volume and diversity, and ocean floor composition. The collective influence of the correct amounts of plant species, combined with ideal water temperature levels and a rocky ocean floor theoretically provides the ideal environment for crustaceans and mollusks to thrive, which in turn serve as the primary food source for other species including the sea otters. The last subset has reference to sea activity levels and suggests that prevailing wind patterns, combined with the presence or absence of physical formations, which inhibit water flow, combined with prevailing currents account for the most predominant factors in determining the strength of sea activity levels. Surge, wave action, and extreme tidal flows are all influenced by these three phenomena.

Perhaps the most challenging aspect of using Geographic Information Systems toward spatial modeling is the recognition that each of these IndVars included in this study must be recreated within the GIS environment before they can be examined. More traditional forms of empirical analysis would simply rely on sampling methods and approximation schemes to quantify information relative to these phenomena, but with GIS, we can use air or satellite photography to serve as the base layer and then create each IndVar layer through a combination of photo interpretation, GPS based ground truthing, and transect/sub-regional data collection. The aggregate effect

of this multifaceted approach creates analytical environment, which allows researchers to employ a host of tools to determine hypothesized relationships. The most prominent of these tools would be the use of thematic analysis to profile sub-regional areas, which display the optimum singular conditions, combined with SQL, based overlap analysis. This combinatorial approach would zero in on those sub-regions which demonstrated the most likely combination of influences and collectively, the research effort would most likely delineate that ideal sea otter habitat is a delicate balance of factors.

From the strategist's perspective, several controllable influences would have been identified as a byproduct of the micro-level analysis which was conducted and which could be applied to answering the initial question of "how do we increase the number of sea otter colonies along the California coastline"? Our research would indicate that relocation of breeding pairs is but one variable, which must be considered. A site identification effort, which located those regions that maintain similar water temperature ranges, wind patterns, prevailing ocean currents, and bottom topography, must also be found, if we are to assure any degree of success. In addition to these "non-controllable" factors, there are several other variables, which could be manipulated within these regions that could increase the probability of a successful transplant. Prior to the relocation of our breeding pairs, it might be necessary to create artificial reefs, which could control wave action, and which would also minimize the

adverse effects brought about by radical fluctuations in the ecosystem that the otters might be unable to cope with. Additionally we could manipulate the composition of the ocean floor to provide a more stable environment for potential food supplies and transplant indigenous species of underwater plant life from the host area to the experimental site.

GIS systems make it much easier to engage in this form of complex empirical modeling. Prior to the development of these types of systems, a good deal of labor-intensive handwork had to be done. This involved assembling large teams of field researchers who would map areas by hand and then teams of analysts would examine the information collected, construct transparencies from the data, and develop interpretations and conclusions. As you can imagine, this was a very expensive process. Today's GIS systems make these types of advanced research efforts commonplace. They afford research principles with relatively inexpensive tools that can be used by empirically oriented team members to produce extremely complex models. The key to success in these endeavors is simply that those people using these tools cannot fall into the trap of oversimplification.

Chapter 12

Conclusion and Final Thoughts

We have reached that point in the book where I get to ramble on about anything I want to in order to draw closure on the materials presented. This is my favorite part of the book, as I suspect it will be yours as well, knowing that you're near the end. If your professor made you read this book, great. Tell them I said thank you. If you elected to read the book on your own, even better. In either case, I would like to leave you with these final thoughts.

The Merits of the Educated Guess

As I have stated over and over, the world in which we find ourselves is not at all simple and almost nothing in it is univariate in nature. Virtually everything is multivariate in its origin and in its state. I'm not sure if this was by design or whether it was the product of accidental interaction, but the fact remains that it is the way it is.

None of us are blessed with divine insights into all things and each and every one of us is engaged in a lifelong journal to discover the truth. Some people gave up long ago because it was too hard, or they weren't interested, or they couldn't comprehend the process, or they just didn't care, but apparently you're still engaged because you read this book.

I believe it's important to remember not to be overly confident in your ability to render a decision, because as you've discovered, there are a plethora of factors associated with the behavior and dynamics of everything in the universe. Whether it is a physical property, or an evolutionary trend, or even something as perceptively simple as rendering a decision or perspective, it's all multivariate in nature. To make it even more complicated, it's multidirectional as well as multivariate and the combinations are almost limitless and the correlations and pathways that need to be examined are beyond the abilities of most of us to compute.

Despite the enormity of the task in answering the question [why] an educated guess is better than no guess at all. You'll find as I have, that some "educated" people venture a guess (sometimes with conviction) because they acquired an education and not based on an in depth analysis of the issue or its variables. Those people are easy to spot and who the guy was talking about when he said you can't argue with ignorance. You can and should put them on the spot and identify the shortcomings of their logic, but you probably won't persuade them of the inaccuracies they put forth in their unsubstantiated supposition. You on the other hand, now know better, so you will be held to a higher standard of being accurate. As I mentioned if you aren't certain, then qualify your assertions, and if you have no idea, then say so. Educated guesses are the next best thing to being right, but only if

they are truly educated guesses. Opinion is simply opinion, even if it is offered by someone with an education.

Absence of Evidence is not Evidence of Absence

Despite what you may have heard, just because someone doesn't have evidence doesn't mean it ain't true. We can employ all of the methodologies we have at our disposal, and design even the cleverest, multivariate-multidirectional equations that we can ponder, but there are some things that simply are beyond our present ability to deduce. This isn't new. It's been happening to us since we arrived. We thought the earth was flat, the stars were holes in the black curtain of space, that the atom was thing of science fiction, and an unending list of knowledge that was just beyond the reach of previous generations. With discovery came new thoughts and revelations about the world and we used these new perspectives to build an even more comprehensive level of understanding of the universe. We are still today engaged in this process of discovery and will most likely be involved in the process until we cease to exist. It's simply human nature to look over the horizon to see what's on the other side and then peer over an even more distant horizon in the search for truth.

Why does an atom or molecule not require the absorption of energy from another atom or molecule to sustain its existence, yet as soon as molecules combine to form a cell, a constant struggle ensues to find sufficient energy to maintain the cellular state of existence?

Why is it that all living creatures require sleep? Who could they be communicating with while they're asleep? If they don't sleep, they'll die. If they die, all those cells they've acquired revert back to molecules and atoms. Interesting.

Scientific reasoning can be accurately described as the ongoing process of theorization and hypothesizing regarding why things happen. It is supported by a continual process of elimination and reconsideration in order to isolate the influential factors pertinent to the conclusion. This, as you've discovered is an endless journal that requires continual re-evaluation and re-assessment. The fools among us forget this requirement.

I would like to share with you one final thought as we near the conclusion to this particular journey of discovery and that is just because someone is older and in a position of power doesn't mean they no longer need to articulate all of the variables and findings that support their conclusion. Yet, this often happens as people aspire and attain positions of authority. It may stem from the belief that because they are in charge, they no longer are required to defend their position on an issue. Or perhaps it is related to the fact that the person in authority either hasn't done the math or can't do the math and as a result they are relying on purely intuitive processes to arrive at a conclusion. Whatever the reason, as you have discovered by reading this book, each and every one of us is bound by the covenant of assuring that our conclusions are above reproach and that we are equally required to provide

thorough explanations of our arguments so that others have the opportunity to review the logic we used in forming our conclusion and either confirm our observations or point out its deficiencies. Without critical reasoning and peer review we rarely achieve validation of our conclusions and we run the risk of limiting perspective or overlooking important variables. By explaining our logic and reasoning to others (even those who are junior in the organization) we stand a better chance of assuring the accuracy of our equations and passing along our reasoning approach to others who may not be aware of how to construct a scientifically verified argument. This alone merits our remembrance that we all have a duty to not only employ sound logical practices, but to share our insights and processes with others to help them understand our motivation for seeing the world as we do. Failure to share our insights with those around us, who have committed themselves to the same enterprise, can breed resentment, contempt, and disdain. By simply taking the time to assure that we go through the process of thinking through the multivariate equation and then sharing those insights with others, we help them understand our position, teach them how to be a more effective thinker, and elicit their continued respect for our views of the world because they understand how it was that we derived the answer to the question [why].

Appendix

Formula for Z ratio

$$Z = \frac{M_1 - M_2}{SE \ M_1 - M_2}$$

Formula for Student T ratio

$$T = \frac{M_1 - M_2}{\sqrt{\left[\frac{S_1}{N_1} \right]^2 + \left[\frac{S_2}{N_2} \right]^2}}$$

Formula for Chi Square

$$\chi^2 = \sum \left(\frac{fo - fe}{fe} \right)^2$$

Formula for Pearson's R

$$R = \frac{\sum xy}{N \sqrt{SX SY}}$$

Values for Student's T

Degrees of Freedom	p=0.05	p=0.025	p=0.01	p=0.005
1	12.71	25.45	63.66	127.32
2	4.30	6.20	9.92	14.09
3	3.18	4.17	5.84	7.45
4	2.78	3.50	4.60	5.60
5	2.57	3.16	4.03	4.77
6	2.45	2.97	3.71	4.32
7	2.36	2.84	3.50	4.03
8	2.31	2.75	3.36	3.83
9	2.26	2.68	3.25	3.69
10	2.23	2.63	3.17	3.58
11	2.20	2.59	3.11	3.50
12	2.18	2.56	3.05	3.43
13	2.16	2.53	3.01	3.37
14	2.14	2.51	2.98	3.33
15	2.13	2.49	2.95	3.29
16	2.12	2.47	2.92	3.25
17	2.11	2.46	2.90	3.22
18	2.10	2.44	2.88	3.20
19	2.09	2.43	2.86	3.17
20	2.09	2.42	2.84	3.15
21	2.08	2.41	2.83	3.14
22	2.07	2.41	2.82	3.12
23	2.07	2.40	2.81	3.10
24	2.06	2.39	2.80	3.09
25	2.06	2.38	2.79	3.08
26	2.06	2.38	2.78	3.07
27	2.05	2.37	2.77	3.06
28	2.05	2.37	2.76	3.05
29	2.04	2.36	2.76	3.04
30	2.04	2.36	2.75	3.03
40	2.02	2.33	2.70	2.97
60	2.00	2.30	2.66	2.92
120	1.98	2.27	2.62	2.86
infinity	1.96	2.24	2.58	2.81

Values for Chi Square

df	p= 0.10	p=0.05	p= 0.025	p=0.01	p=0.001
1	2.706	3.841	5.024	6.635	10.828
2	4.605	5.991	7.378	9.210	13.816
3	6.251	7.815	9.348	11.345	16.266
4	7.779	9.488	11.143	13.277	18.467
5	9.236	11.070	12.833	15.086	20.515
6	10.645	12.592	14.449	16.812	22.458
7	12.017	14.067	16.013	18.475	24.322
8	13.362	15.507	17.535	20.090	26.125
9	14.684	16.919	19.023	21.666	27.877
10	15.987	18.307	20.483	23.209	29.588
11	17.275	19.675	21.920	24.725	31.264
12	18.549	21.026	23.337	26.217	32.910
13	19.812	22.362	24.736	27.688	34.528
14	21.064	23.685	26.119	29.141	36.123
15	22.307	24.996	27.488	30.578	37.697
16	23.542	26.296	28.845	32.000	39.252
17	24.769	27.587	30.191	33.409	40.790
18	25.989	28.869	31.526	34.805	42.312
19	27.204	30.144	32.852	36.191	43.820
20	28.412	31.410	34.170	37.566	45.315
21	29.615	32.671	35.479	38.932	46.797
22	30.813	33.924	36.781	40.289	48.268
23	32.007	35.172	38.076	41.638	49.728
24	33.196	36.415	39.364	42.980	51.179
25	34.382	37.652	40.646	44.314	52.620
30	40.256	43.773	46.979	50.892	59.703
35	46.059	49.802	53.203	57.342	66.619
40	51.805	55.758	59.342	63.691	73.402
45	57.505	61.656	65.410	69.957	80.077
50	63.167	67.505	71.420	76.154	86.661
55	68.796	73.311	77.380	82.292	93.168
60	74.397	79.082	83.298	88.379	99.607
65	79.973	84.821	89.177	94.422	105.988
70	85.527	90.531	95.023	100.425	112.317
75	91.061	96.217	100.839	106.393	118.599
100	118.498	124.342	129.561	135.807	149.449

Values for Pearson’s R

df	p < .05	p < .01
3	.997	.999
4	.950	.990
5	.878	.959
6	.811	.917
7	.754	.874
8	.707	.834
9	.666	.798
10	.632	.765
11	.602	.735
12	.576	.708
13	.553	.684
14	.532	.661
15	.514	.641
16	.497	.623
17	.482	.606
18	.468	.590
19	.456	.575
20	.444	.561
40	.312	.403
80	.224	.292
100	.195	.254
200	.138	.181
500	.088	.115
1000	.062	.081

