

How to Build Computers that *Think*

In 10 Easy Steps

August 16, 2013

SHARE

Boston, Hynes Convention Center

Bill Valyo

Copyright 2013, William J. Valyo, Jr.

Rights to reproduce are granted within the scope of the SHARE users group.

Trademarks

The following trademarks are retained by William J. Valyo, Jr. as they relate to standards certification within the scope of the BRAIN architecture:

Broad-Reasoning Artificial Intelligence Network™

BRAIN:Power™

BRAIN:Cell™

BRAIN:Tree™

BRAIN:Stem™

BRAIN:Pan™

BRAIN:Surgeon™

BRAIN:Trust™

BRAIN:Fever™

BRAIN:Scan™

BRAIN:Dump™

BRAIN:Wave™

BRAIN:Child™

BRAIN:Storm™

While the trademarks reflect existing standards, it is the intent of Mr. Valyo to transfer ownership of the trademarks to a future standards organization focused on interoperability within the scope of Broad-Reasoning Artificial Intelligence Network implementations.

Table of Contents

Part 1: Overview.....	5
Abstract.....	5
About this Document.....	5
Part 2: The Ten Recursions of Intelligence.....	6
One: The Physical Recursion	7
Indicators within Neuroscience	7
Indicators within Psychology	7
BRAIN Implementation	7
Two: The Knowledge Recursion.....	8
Indicators within Psychology	8
Indicators within Neuroscience	8
Indicators within Language	8
Indicators within Computer Science	9
BRAIN Implementation	9
Three: The Indexing Recursion.....	11
Indicators within Psychology	11
Indicators within Computer Science	11
BRAIN Implementation	12
Four: The Process Recursion	14
Indicators within Neuroscience	14
Indicators within Computer Science	14
BRAIN Implementation	14
Five: The Programming Recursion	16
BRAIN Implementation	16
Six: The Reasoning Recursion.....	17
Observational Indicators.....	17
BRAIN Implementation	17
Seven: The Functional Recursion	19
Observational Indicators.....	19
Eight: The Operational Recursion	20
Logical Assertions.....	20

BRAIN Implementation	20
Indicators within Psychology	22
Indicators within Computer Science	22
Nine: The Understanding Recursion	23
Indicators within Language	23
BRAIN Implementation	23
Ten: The Organizational Recursion	24
Logical Assertions.....	24
Indicators within Language	24
Indicators within Human Organizational Knowledge	25
Indicators within Neuroscience	25
BRAIN Implementation	25
Part 3: The Components of Intelligence	29
BRAIN:Power – Reasoning	30
BRAIN:Pan – Memory Retrieval	31
BRAIN:Dump – Outgoing Language	32
BRAIN:Scan – Incoming Language.....	33
BRAIN:Wave – Perception	35
BRAIN:Child – Logic.....	37
BRAIN:Storm – Creativity.....	38
BRAIN:Surgeon – Intuition	40
BRAIN:Trust and BRAIN:Fever – Confidence and Emotion	41
Part 4: Removing the Roadblock of Consciousness	42

Part 1: Overview

Abstract

This paper outlines two related technologies:

- **Broad-Reasoning Theory** - a collection of inter-dependent theories surrounding the *process* of human intelligence.
- **The Broad-Reasoning Artificial Intelligence Network** - a proposed software architecture for recreating human intelligence in its entirety on a small computer. Throughout this text, I will refer to the software by its acronym, "the BRAIN".

About this Document

The text appears in four parts:

- Part 1 is this brief overview.
- Part 2 summarizes the ten fundamental recursions that comprise human intelligence.
- Part 3 outlines how the recursions are used to create the components of intelligence, including language, logic, creativity and others.
- Part 4 removes the stumbling block of consciousness from the AI debate.

Elements of Broad-Reasoning and the BRAIN are illustrated in parallel as the text progresses. In this manner, the software elements may be presented side-by-side with the supporting theory.

This paper describes research performed independently by the author between 1983 and 2003. It is a summary only. Details, including a Technical Reference for those who wish to build a BRAIN, may be obtained separately from the author.

Part 2: The Ten Recursions of Intelligence

Most artificial intelligence theory begins with the assumption that human intelligence is extremely complicated. This is evident in that AI work is often divided into separate domains such as game theory, language processing and facial recognition that are, presumably, to be stitched together with their cousins into some future representation of a functioning human intelligence.

Broad-Reasoning begins with the opposite assumption: I have taken as a foundational assertion that human intelligence must be simple or what a software designer would term "elegant" in order to fit into such an extremely small computer as the human brain. I propose that it makes use of processing "recursions" that are used repeatedly like building blocks to form increasingly larger, more complex structures. These structures include the elements above (game theory, facial recognition, etc.) but more importantly, even more fundamental components such as learning and logic and language and even reasoning itself. I propose that these are all built from nothing but the core, repeating processes and, moreover, that they will work in *all* domains of human endeavor. We could call this domain *agnostic* AI.

In this part of the text, I will outline the ten fundamental recursions from which I proffer all human intelligence is built. The first is obvious and the second is borrowed from psychology and commonly adapted to AI theory. The third is a functional, physical representation of Piaget's theories.

The remaining seven should represent an underlying change to how we view human intelligence.

One: The Physical Recursion

The first building block is obvious: Human intelligence is built from neurons and their connections to other neurons. While this is certainly not ground-breaking, it does provide the ground upon which the remaining recursions are built.

Indicators within Neuroscience

The assertion above is generally referred to as the “Neuron Doctrine”, and is over 100 years old. It’s foundational to neuroscience.

Indicators within Psychology

We can propose this in another way, by suggesting that human *knowledge* is encapsulated in some manner in the network of neurons. In the 1970s, John Anderson proposed that there are two types of knowledge, declarative (the “facts and figures”) and the procedural knowledge (that structures how we think). In computer terms we may think of these as the data and the programming. These two elements contain the entirety of human intelligence.

While we may use that divide to structure our investigation, I intend to show through one of the recursions below how *both* are contained within the network of neurons. More importantly the procedural cells that compromise the fundamental programming language of intelligence are actually a *subset* of the declarative cells.

BRAIN Implementation

To mimic this network, the proposed artificial BRAIN is to be built largely from a construction that I call “BRAIN:Cells”. Logically, these cells are comprised of three parts that mimic the like-named human counterparts:

- **The soma** are individually addressable (this is in effect, the central body of the BRAIN:Cell).
- **Axons** point to other soma/cells (by their addresses).
- **Dendrites** serve to *prevent* connections from other axons.

Note that these names are not strictly related to human cells, but are used here in a conceptual way. Its worthy of note that human cells only have one axon, whereas the proposed artificial BRAIN:Cell will have a myriad. This is addressed farther below.

Intelligence is built from knowledge structured through a latticework of neurons.

Indicators:

- The Neuron Doctrine
- The declarative/procedural knowledge division

Significance:

- Structure enabling replication of the entirety of human intelligence.

Proposed BRAIN Component:

- BRAIN:Cells that replicate human neuron clusters.

Two: The Knowledge Recursion

Naturally, before we can program a brain artificially, we have to understand how human neurons achieve both declarative and procedural knowledge. We start with the less complex declarative knowledge in this section.

Indicators within Psychology

Broad-Reasoning truly begins with the assumption that there are discrete units of knowledge represented by neurons. While I would like to believe that the theory was new to Broad-Reasoning, it was first proposed in the 1960s, roughly 20 years earlier.

Jerry Lettvin and Jerzy Konorski, working independently, theorized that *each* neuron represented a discrete piece of knowledge. Lettvin famously referred to this concept as the "grandmother cell", while Konorski called it a "gnostic" cell. Lettvin's approach was far more colorful and descriptive, illustrating that there was a single brain cell in each of our heads that identified a single one of our grandmothers. I include this theory as a foundational element in Broad-Reasoning, with small modifications.

Indicators within Neuroscience

The theory was proven to the degree possible with present technology, when Rodrigo Quijan Quiroga and team were able to localize activity in the brain triggered by photos of celebrities, both real and fictional.

The slight modification here is that I don't believe that *all* neurons singularly represent discrete units of knowledge for two reasons:

- First, functional MRI does not have the degree of accuracy to truly prove this assertion.
- Secondly and more importantly, each natural neuron has only one axon that triggers downstream activity. For this arrangement to work, it would be logical to include multiple axons that are triggered independently. However, the effect may be achieved by clusters of cells that are generally called "microcircuits" in the neuroscience community. I offer that each of these clusters generally represent a discrete unit of knowledge such as your grandmother or her house, or her signature cherry pie.

Indicators within Language

Your biggest question at this point is probably: how is it determined which is a discrete unit of knowledge? Would it be the knowledge of

Individual clusters of cells represent units of declarative knowledge.

Indicators:

- Konorski's "gnostic cell"
- Lettvin's "grandmother cell"
- Quran's research into facial recognition

Significance:

- Artificial knowledge may be constructed in discrete units.

Proposed BRAIN Component:

- BRAIN: Cells that contain human knowledge.

a human or knowledge of a human hand?

I would propose that it is both.

Stated simply, a discrete unit of knowledge is anything that we can name. In fact, I will go farther and suggest that giving an unnamed element a name (such as the mysterious “factor x”), solidifies it into a single cell in our brains. I suspect that this is where the process of neurogenesis begins. I was presented with the same dilemma when I was describing elements of intelligence that are unique to Broad-Reasoning theory. Until I gave these elements a name, they seemed to be contained in separate process steps. Once they had a name, the process steps became subordinate as a unit, although it often took me several days to achieve that singularity in my own mind. It was my belief that naming it was fixing it permanently in my brain, and it is my intent to use the new names to do the same to yours.

Proposed Natural Language Processing (NLP) systems generally structure vocabulary into similar units that are arranged hierarchically.

While I agree that vocabulary should be arranged in such a manner, I disagree that vocabulary is a separate infrastructure beyond the words used. I propose that vocabulary and the knowledge used to represent the actual concepts are synonymous, and that any attempt to separate vocabulary from knowledge will fail for reasons that will be explained throughout the text.

Indicators within Computer Science

I am not suggesting that a neuron *contains* knowledge, but only represents it. Of course, you wonder how that can be. I assert that the connections between related cells form a neurological fabric representing your grandmother, just as she is built (at least in your mind) from images of her face, memories of time spent, her logical location in your family hierarchy and all other facets of any similarly complex human life.

There is a simple analogy within computer science that will drive this point home. Taken singularly, neither you nor the computer can perceive the meaning of a single bit. It is only when taken within context of a sequence of bits or a hierarchy of files or memory containing bits that the bit is given meaning.

BRAIN Implementation

The BRAIN:Cell, proposed above doesn’t reflect individual cells, but the cellular microcircuits. In the BRAIN design we place a small amount of text in the soma that is the name represented by the cell. The name may be specific such as the name of your grandmother, or

generic, like the word "grandmother", or even unknown, like the nebulous "factor x". While this is certainly cheating, it is intended to work within the limitations of the computer, and may prove to be a more efficient means than the human counterpart.

BRAIN:Cells can be contained within conventional data such as RDBMS tables, XML, flat files, etc., but will be communicated as XML for the purposes of interoperability.

Please note that there is a vast difference between knowledge and conventional data. The best example might be a phone number, which is typically a field contained within a larger traditional data record that might include other fields associated with an individual or an account. The phone number is understood in its entirety.

On the other hand, a phone number recorded as *knowledge* within our heads is typically divided into four cells that may include the country code, area or city code, exchange and the unique number. Each of these may relate to other knowledge. The country code is of course related to a country which may give you an indication of the language of the person on the other end of the call. The city code might relate to time zones. In the cases of small nearby towns, you might be able to identify the town of the caller.¹

One of the challenges with an artificial knowledge is that the BRAIN:Cells must be individually addressable. Naturally this implies a structure. That structure begins to reveal itself in the next recursion.

¹ It's worth noting that there are certain elements within human experience that may be addressed more efficiently as "alternate indices". Examples would include the phone numbers above and similar elements that point to an individual or individual cell like email addresses, twitter handles, etc. The most common use of this approach would be in the linguistic programs that are described in the next section. I expect that indices of words and phrases will be critical to conversational language in real time.

Three: The Indexing Recursion

When you think of abstract terms like “honesty” or “love”, what do you picture?

I propose what I term a "model" cell to contribute a great deal of understanding and structure to human intelligence. Although the model cell has been called by other names in psychology.

Indicators within Psychology

Jung called such an element an "archetype". Piaget termed it a "schema". The small yet substantial difference that I am proposing is an actual, physical cell, building on the knowledge recursion above. Secondly, I suggest that it exists for *all* elements, not just the limited number of "psychological" elements in Jungian usage.

To explain, allow me to use an example. I propose that we all have a model of what constitutes a tree. I proposed that the model cell exists for all trees, and from this model cell you may have hanging a similar cell that identifies all coniferous or deciduous trees, although you may use different names. Furthermore, you may have an instance of one of these trees in your yard that is similarly represented by a cell in your brain. I have a dogwood by my front door. I propose that my particular dogwood is represented (for me) by a single cell that is connected to a cell representing all dogwoods, including those of my boyhood home.

The model cell points to other cells that further describe the element. For example, you may assume the default color for all trees is green. It's a fair assumption and a reasonable default color. Only when you encounter a tree in the fall or a plant such as a naturally red Japanese maple would your assumptions change, but they would change only for the cell representing the Japanese maple.

Indicators within Computer Science

The model cell solves several problems:

- Addressability - whether human brain cells are physically connected to a model in a hub and spoke system, or artificial cells that may use the model as a node in an index hierarchy, the model becomes an organizational necessity.
- Compression - a model may hold the qualities of all women, for example, rather than requiring that each woman you know has the same qualities imprinted in brain cells describing her. This is a dramatic space saver in the human

Knowledge is clustered around model cells.

Indicators:

- Piaget's schema
- Jung's archetypes
- Logical continuation of knowledge recursion above.

Significance:

- Provides logical index to knowledge categories.
- Provides defaults (assumptions) for knowledge categories.
- Provides unknowns for further review.
- Provides implied structure for common (shareable) knowledge.
- Provides separation of common and personal knowledge.
- Makes better use of available knowledge storage space.

Proposed BRAIN Components:

- The model cell variant of the BRAIN:Cell.
- Different indexing scheme for model and non-model cells.
- BRAIN:Stem as a standard collection of model cells.
- BRAIN:Tree as a unique representation of specific data that includes the stem.

or electronic computer.

- Defaults (Assumption) - the "default" arrangement of the above, allows us to make assumptions when encountering new instances of similar knowledge. For example, we can largely assume that all cars have four wheels, although it may not always be true.
- Program Simplicity - For artificial intelligence programming, the model cell effectively provides an object-oriented programming (OOP) approach, thereby simplifying our task.
- Separation of common and personal knowledge - This item is specific to *artificial* intelligence, and requires greater explanation below.

For artificial intelligence to be practical, artificial knowledge has to exist in two opposing frames:

- Personal knowledge - It must be unique to all implementations (so that we can have our own knowledge and even "truths").
- Common knowledge - It must be standardized to allow of indexing and sharing across computers (not unlike we have a common language that allows human beings to share knowledge).

BRAIN Implementation

The model cell solves this. Artificial intelligence can be produced and shared with a standardized knowledge hierarchy of model cells. This model only framework, I call the "BRAIN:Stem", as it forms the trunk of a large hierarchical tree.

A specific instance of this tree is known as the "BRAIN:Tree". Your personal knowledge is contained on a "leaf" of that tree. In such manner, we can share the concept of a "mother", but your mother and mine are different.

In the BRAIN, a model cell is differentiated with a mere flag and a small indexing element. The indexing scheme is described in more detail in the BRAIN Technical Reference and other documents available from the author, but it's essentially a tuple number identifying all levels of the tree with a master. The leaf is the text contained in the instance or leaf cell. Therefore, your mother might

be recorded as: 1.2.456.675.46.34.Marjorie Smith.²

Certainly, there are a number of challenges with the concept of a vast knowledge hierarchy. Some of them will be addressed in later recursions in these pages; some of them are met in more comprehensive works on Broad-Reasoning theory including the above-mentioned *Technical Reference*.

² This is a rough example. Please do not assume the indexing scheme from this number. The *Technical Reference* will contain details.

Four: The Process Recursion

In the first few recursions, we had started to map only declarative knowledge, but with our incursion into problem solving, we are skirting the procedural. To create an *artificial* intelligence, we have to understand how neurons can be combined to create the procedural knowledge that we'll use to replicate human thought.

The answer is found in the neuron itself and is the next recursion.

Indicators within Neuroscience

As is widely known, neurons don't actually touch, but have to overcome a physical gap or synapse through a chemical bridge. Consider that the synapse is not an accident of nature, but must have a function. If knowledge is built when two cells connect, the functional purpose of the gap has a human analog:

What color are *your* lover's eyes? The answer to that question joins two cells, according to the above theory: a cell representing your lover, and a cell representing eye color. When those cells are joined, the synapse is bridged.

In that context, consider the following: Is it possible that the synapse is effectively *a question that must be answered*?

Indicators within Computer Science

Please don't miss that previous statement; it is probably the most important in the whole of Broad-Reasoning and certainly this text. From this simple assertion, much can be built. Consider that all problems are really a series of unanswered questions. Questions abound in our intellectual universe. Knowledge is achieved and solutions are mapped when the questions are answered. All processes within intelligence have their own problems to solve and are a series of questions that must be answered.

Consider the simple act of understanding a sentence: What is the subject? What is the verb? Computer programming also has an analog: Variables are simply questions. A variable/value pair is an answered question.

BRAIN Implementation

I propose that the *question* is the fundamental recursion at the heart of human intelligence. The BRAIN and Broad-Reasoning Theory will exploit this concept to model the entirety of human intelligence.

However, for this to work, there are at least two missing elements:

The process of intelligence is built upon the concept of "the question".

Indicators:

- The logical purpose of the synapse.
- The use of variable/value pairs within computer science.

Significance:

- The question is the unknown.
- The unknown is a natural occurrence that intelligence seeks to overcome.
- Provides a structure that leads to a natural programming language that is easily replicated in a computer.

- There must be a method for *recording* questions (that is effectively a natural programming language).
- There must be a standard means for *answering* questions built into human intelligence.

Thankfully, both already exist and are described in the next two recursions.

Five: The Programming Recursion

The first problem we have to solve is the recording of questions. Luckily, we've actually achieved this already.

Allow me to revisit model cells. You may recall that they have two purposes. Using the car example they may:

- Index all instances of a given model (all the cars *you* have had, for example)
- Provide the default characteristics for *all* cars

I may have undersold the latter element. Model cells are intended as placeholders for *all* qualities of a given element. So while they may include defaults, they also include variables in the form of pointers to other cells. Using the car example, you may assume that a car is two-wheel drive or automatic transmission because they are most common (at least in the US), but the very fact that the variable/cell-pointer exists, implies that it is a changeable variable. Other variables such as car color cannot be assumed.

All of these variables are questions. Each time you encounter a new instance of such knowledge, you may find that you have a number of questions. For example, when your friend tells you that he's considering buying a new car, you may ask him about make, model, color and many other such qualities so that you may build a conceptual version of *his* car branching off of your model cell.

BRAIN Implementation

However, the more important aspect of model cells is not associated with declarative knowledge like cars, but procedural knowledge such as language and logic. You will see an example in part three of this text. We can use singular cells or hierarchies of connected cells to form what I call "thought programs". The thought programs are merely the questions and common defaults. They may represent any procedural knowledge from how to tile a bathroom floor to how to be creative, provided that we know the right questions to ask.

This is also provided that we humans have a means to answer any question. We do, of course, and that is the next recursion.

The model cell records questions in a hierarchical fashion that represents "thought programs".

Indicators:

- Model cell already solves this problem.

Significance:

- Simple, existing method for "coding" procedural intelligence within a neuron.
- Declarative and procedural knowledge are largely the same thing.

Six: The Reasoning Recursion

I propose that there is a universal process that we use to answer questions, which itself is simply a series of questions and therefore a thought program. I call the process "diaductive reasoning" as it spans and controls other forms of reasoning.

If the question is the central element in intelligence, then how we answer that question is the central process of intelligence. In effect, the process *is* human reasoning.

The technique is extraordinarily basic:

Observational Indicators

When faced with a question, I would posit that we simply...

- Determine if we know the answer already. This step usually takes place subconsciously, as it more than likely an electro-chemically based neuron search. Obviously the BRAIN must mimic this search.
- Determine if we can obtain the answer from someone else (in a timely fashion). The source may be human or document. Naturally, this implies the ability to use language and as such, language thought programs are necessary for the BRAIN.
- Determine the answer ourselves. This is often much more challenging, and requires the use of logic, intuition and creativity, each of which are called as thought programs and are described in part 2.

BRAIN Implementation

In the BRAIN software, the above is also a thought program called BRAIN:Power that I contend contains the core process of human reasoning. It consists of only these three questions.

BRAIN:Power does little on its own, but is the central recursive thought program at the heart of the BRAIN. (As a result, it's one of the few thought programs mainly built from conventional programming languages.) It primarily calls other thought programs to answer the questions.

The first called thought program is what I call "BRAIN:Pan". It effectively represents human knowledge retrieval as it searches knowledge to determine if the answer is already known. To

Diaductive reasoning is a simple thought program targeted at answering questions.

Indicators:

- Observation.

Significance:

- Identifies a simple process that attempts to answer any question.
- Calls other thought programs and is used recursively to answer the questions implied in the called programs.

understand how BRAIN:Pan works, we will need to understand more about how knowledge is organized (another recursion, found below).

In order to answer the second question (Can I determine the answer *externally?*), BRAIN:Power will call one of two language thought programs:

- BRAIN:Dump represents outgoing language, and can structure the question to an expert. Stated more simply, it writes.
- BRAIN:Scan is the thought program that addresses incoming language. Stated another way, it reads. It will be used to understand the answer from the expert, or to read an expert document.

Determining who can be considered “the expert” is addressed in a number of ways. It may be included in other thought programs that are called downstream, it can be embedded into a range of data (i.e.: the fundamental treatise on Chemistry is the Linus Pauling book), or it could be a default such as a dictionary or encyclopedia. In a pinch, a BRAIN could even do a Google search. (We will confront the trustworthiness of knowledge when we replicate human “confidence” through a program called “BRAIN:Trust”.)

The final question (Can I answer the question internally?) is addressed by a combination of programs that may call *each other*. As stated above, these are perception (BRAIN:Wave), logic (BRAIN:Child), creativity (BRAIN:Storm) and intuition (BRAIN:Surgeon).

Seven: The Functional Recursion

Thus far, I have been very nebulous in describing the *process* of intelligence. It is time to refocus our activities. A simple way to do that is to ask ourselves what human intelligence truly *does*.

Observational Indicators

I don't believe that anyone will argue that intelligence solves problems. I would like to take this farther, however, and suggest that problem-solving is *all* that intelligence does.

Some would argue that human emotions such as love, and human endeavors such as art are not representative of problem solving. Yet somehow they solve problems for those that engage in them.

If you will consider the possibility that the entirety of human intelligence is focused upon solving problems, it solves a huge problem for us. Problem-solving becomes the principal operational construct. Stated another way, if the brain is solely solving problems, then the "problem" is a recursion.

Intelligence is solely focused upon solving problems.

Indicators:

- Observation

Significance:

- Provides a single unit of work for the whole of intelligence.

Proposed BRAIN Components:

- See next recursion.

Eight: The Operational Recursion

Of course, for us to code the problem presents us with a problem.

Logical Assertions

Consider the following:

If:

- the "problem" is a singular recursion, and
- all elements that may be named may exist as a cell, and
- a model cell exists to index all like items,
- then there must be a model cell that represents the concept of a problem.

BRAIN Implementation

I call the cell the "universal problem". Such a model cell is important for two reasons:

- It extends the indexing scheme to include active problems. That is, all problems are individual "leaf" cells pointed to from the model cell. The universal problem model cell is in effect, a problem queue not unlike those contained in all computer operating systems.
- It provides a model that is a pattern for recording *all* problems.

The latter element is a challenge, but I do not intend to leave it as such. It's critical that the universal problem fit into the deductive reasoning process and may easily be defined within a network of BRAIN:Cells. Moreover, it has to not only model the problem, but provide enough information to suggest a solution.

Keeping all these things in mind, I propose that there are only *seven* elements that all problems have in common, and these can be pointed to from within a model cell.

The first two variables are:

- **The problem owner:** This is an entity (individual or organization) with the problem. Frankly, a problem doesn't exist in our world unless it affects someone. The problem owner is also important because it provides the source of all of the remaining six elements. When considered within the context of diaductive reasoning, the problem owner is the expert in the problem. Such an individual or organization may be queried to understand the problem better.

A universal problem model cell exists, which indexes and contains all problem descriptions.

Indicators:

- By logical extension from the operational recursion.

Significance:

- Allows human problems to be described to a computer.
- Provides a unit of work similar to operating system constructs.
- Provides a problem queue and prioritization.

Proposed BRAIN Components:

- The universal problem model cell.
- The BRAIN:Child thought program that deconstructs and records problems using the universal problem model.

- **The time frame:** Each problem exists in time. It either needs to be solved by a certain time, or is part of a greater sequence or some similar concept. Using our own pursuit as an example: the problem of the model cell needs to be addressed before we can model a universal problem. The timeframe is used to link related problems together.

The next 4 elements are all related to solving the problem, and are based on the assumption that all problems have the concept of change at their core. The solution is either intent upon change, or the prevention of change:

- **The contentious item:** All problems have some item that is the focus of attention, and many problems have several. For example, the problem of keeping beer cold has two contentious items, heat and coldness. A good beer cooler design will keep the cold in and the heat out. While a single solution may solve both, these are effectively two sub-problems of a singular complex problem. Considering the second element (keeping heat out) may result in beer coolers that are lighter in color and can thereby prevent heat from being absorbed.
- **The location of the contentious item:** This is simply the binary, "here" or "there", and can refer to concepts such as qualities or ownership. If the beer is presently cold, cold is here. Heat is there.
- **The desirability of the contentious element:** This can also be binary. When it comes to beer, cold is good, warmth is bad. There may also be more complex value such as cost associated with an item, but these are typically complex problems with multiple items used in comparison. Ultimately the decision comes a go or no-go binary that is the desirability factor.
- **The strategy:** When viewed on a network of cells, the three items immediately above produce a fundamental strategy which is either a flow of change or obstruction of change. In the beer cooler, we are obstructing the flow of the desired coldness from leaving *here*, while simultaneously blocking the undesired warmth from traveling from *there*. The solution, if possible, grows from this fundamental strategy.

The seventh and final element is a nebulous element that we can call "**why**". It is the answer to the simple question, "why am I doing this?" This element puts the problem in a larger context that can

include other knowledge normally considered outside the scope. It governs all assumptions for the problem, and prioritizes all problems in the system against each other.

In fact, the singular concept of "why" is foundational to all human thought, and hence reveals a recursion at the heart of understanding, not only problems, but all things. As it is a part of all things, this will be the topic of the next recursion.

The next part of this paper introduces a thought program called "BRAIN:Child" that is intent upon structuring all problems according to the universal problem model.

Before we leave this topic, however, there are some important after-the-fact indicators to explore:

Indicators within Psychology

I propose that the universal problem (and mapping problems to it) fundamentally represents what the psychologists call "executive function".

Indicators within Computer Science

The universal problem replicates common operating system problem-switching actions by recording a snapshot of the problem.

Nine: The Understanding Recursion

Why is iron solid, but mercury liquid at room temperature? Why is the sky blue? Why did you choose a career in computer science? Why should we pursue artificial intelligence? Why is cancer such a challenging problem to solve? Why are you reading this paper?

Everything, whether concrete or conceptual may be affiliated with the question "why?" Yet why is one of the more challenging problems to solve; therefore, it appears that any artificial intelligence will need to understand the concept of why.

I propose that there is a structure to why that is easy to replicate inside of a network of cells. Why can always be explained as a disposition to change. Stated another way, everything has a disposition to change at its core.

There are generally three approaches toward change:

- **Reason:** Sentient beings and organizations have a motivation behind all of our actions. Again, these actions are either causing change or preventing it.
- **Purpose:** Human inventions and evolved systems have a purpose. That purpose directs or prevents change.
- **Cause:** Non-living things may often be explained by the actions moving within or against them. Consider chemistry and how it forms the very matter we breathe or ingest.

Indicators within Language

Not incidentally, dictionary definitions for the word "why" generally include some variation of the synonyms "reason", "purpose" and "cause".

BRAIN Implementation

"Why" will be used as a part of the organizing structure of knowledge. That is the last recursion, described below.

All things can be understood as a function of change.

Indicators:

- Dictionary definitions of the word, "why"

Significance:

- Gives us a means to categorize all things according to a natural understanding
- Provides a means to understand how various elements will react within a problem situation either as tools or antagonists
- Provides a foundation for understanding human intuition, not as a magical act, but an understanding of all things that is so fundamental as to be below observable consciousness

Proposed BRAIN Components:

- The seven hierarchical subnetworks defined herein.

Ten: The Organizational Recursion

All of the previous elements outlined above are found within knowledge. I've hinted at a hierarchy of knowledge, and certainly some structure is necessary in order for us to *index* artificial knowledge.

In this section, we're going to bring it all together and organize our hierarchy.

Logical Assertions

Consider the following:

The first question in diaductive reasoning is, "do I know the answer already?" Naturally, this implies that there is a methodology for searching knowledge that is structured toward answering questions.

Indicators within Language

Consider also that language offers us a list of principal human questions through a collection of words that are collectively called the "interrogatives" that are more-or-less universal. You would know them more commonly as "who, what, where, when, how and why".

I indicated above that these words were "more-or-less universal" because there are only six interrogatives in Germanic languages like English. Many other language families (Romance, Slavic, others) contain *two* variations of what we English-speakers call "how". There is the *procedural* how, as in how we accomplish some task. I tend to call this "How^P" within the confines of Broad-Reasoning theory. Additionally, there is the qualitative how, as in how much something costs or relative degrees of qualities (how hot is it today?). Similarly, we can call this "How^Q". Even though we only use a single word, the difference is clear to most speakers of Germanic languages.

To be clear, I am suggesting that knowledge is *physically* separated within our brains into *seven* corresponding categories that represent the answers to the common seven questions. To state another way: I am proposing that specific axons will point to specific ranges of knowledge. In the natural brain, these would equate to the downstream cells that are a part of the aforementioned cell clusters that would point to various regions of the brain. In the artificial BRAIN, these would be specific axon types that distinguish between relationships. For example, the BRAIN identifies axons for ownership, synonyms, antonyms, etc.

The entirety of knowledge is structured into one superset and six subsets that house answers to the seven principal human questions: who, what, where, when, how (procedural), how (qualitative) and why.

Indicators:

- The seven interrogatives of language
- The seven matching categories of nouns
- The seven matching variables in the universal problem
- Potentially, the common use of the number seven to organize human knowledge (days of week, colors of the rainbow, references in Judeo-Islamic-Christian heritage).

Significance:

- Knowledge is organized to answer questions.
- There is a common hierarchical structure to knowledge.
- There is a common pattern to the pointers out of each cell cluster or BRAIN:Cell.

Beyond the interrogative words themselves, there is another linguistic indicator that is obvious in retrospect. If everything is a *thing* and nouns describe things and all knowledge is represented in cells representing things, then perhaps we will find an indicator there.

If you recollect the categories of nouns, almost instantly, you will recall “people, places and things”. Yet these were the categories you learned from the grammar in grammar school. As you grew older you were taught of *seven* categories of nouns: people, things, places, events, actions, qualities and states.

These are exactly, who, what, where, when, how^P, how^Q and why. Please consider this as a reasonable indicator that we are on the right track.

Indicators within Human Organizational Knowledge

Human numerical systems are organized around the number ten, likely due to the numbers of fingers and toes we have. What then is responsible for our arbitrary grouping of elements around the number seven? Consider the days of the week, the number of colors we delineate in a rainbow or the odd use of the number seven repeatedly in Judeo-Islamic-Christian religious texts.

Is it possible that seven has a similar physical aspect, perhaps associated with the branches connecting a given cell to seven regions of knowledge?

Indicators within Neuroscience

Neuroscience is beginning to identify regions that match the seven I’ve identified above. For example, decades of experiments with rats have led to an understanding of place, grid and head-direction cells in mammals that help locate us in our environment, that falls within the “where” network I’ve proposed.

BRAIN Implementation

As a last indicator, consider the proposed universal problem model cell. If the brain is focused solely upon solving problems, isn’t it possible that the human brain has evolved the universal problem to be at the top of the knowledge hierarchy? If that is true, then the universal problem model would reflect the seven categories.

Observe that the entity with the problem is *who*. The contentious item is *what*. The relative location is *where*. The time frame is *when*. The strategy is *How^P*. The desirability is *How^Q*. And finally, the why factor is of course, *why*.

What I'm suggesting is there is a hierarchy of knowledge that has

one superset and six subsets. The superset is the *what* aspect. (Everything is a thing and therefore a what.) The remaining six subsets are equally subordinate.

For clarity's sake, I call the sub-networks by the following names, identified in the first column of the table below:

Network	Interrogative	Noun	Universal Problem
Social	Who	Person	Entity
Elemental	What	Thing	Item
Positional	Where	Place	Location
Temporal	When	Event	Time Frame
Procedural	How ^P	Action	Strategy
Qualitative	How ^Q	Quality	Desirability
Dispositional	Why	State	Why

Elemental Superset

The elemental network is the master network and contains all others. (Everything is a thing!) It is subdivided in its first level by the remaining 6 sub-networks. It is from this superset that we get the indexing of all things.

Note that our index is merely a relative number. We could use any organizational scheme, but I believe that there is a natural index that relates to other elements. Once we separate into the remaining six elements, we still have to divide the remainder of the elemental network which is still large. I would suggest that the *common* delineation of all elements is as follows:

- Living things
- Non-living things

This divide is based upon the disposition to change that we investigated in the last section. As I mentioned, I contend that disposition to change is part of the intuitive way that we understand all things. As such, it is a logical divide.

Dispositional Subnet

The dispositional subnet has three differing layers:

- **Motivation Hierarchy (Reason):** This is a layer of priorities similar to Maslow's hierarchy of needs. It establishes a human priority scheme. Problems are attached to these priorities to identify relative priority to each other.
- **Cause:** Contains a hierarchy that I call the Grand Hierarchy, that defines know causes for natural creation. It is a hierarchy that attempts to encode chemistry and physics that begins with the proposed singularity and establishes the

theoretical basis for elementary particles, sub atomic particles, elemental atoms, molecules and complex molecules. I recognize the challenge behind such a process and that most human brains would not contain this knowledge, but the hierarchy is useful toward understanding the underlying cause (changes) for so many things in the physical universe.

- **Purpose:** The Purpose hierarchy defines human-induced change or change-avoidance inherent in tools and other human creations such as organizations. It may ultimately be synonymous with much of the Procedural subnet, described below.

Qualitative Subnet

The fundamental qualities of a given cell are likely pointers to those qualities within an addressable hierarchy of those qualities. While that hierarchy may include divisions for individual sensory qualities like color and temperature, it must also include qualities that are not perceived by a single sense (like motion or distance), or those that may not be perceived by senses at all.

Positional Subnet

Like all of the subnetworks, the individual locations within the positional subnet may be indexed from the top down in a hierarchical fashion (a place exists in a hierarchy that includes the planet, nation, state, county, town, street, building, etc.). It may also be relatively indexed by starting at a point within the network and moving in proximities (up, down, in, out, near, far).

Temporal Subnet

Similarly structured is the temporal subnet. A given event may be located hierarchy through eons, millennia, centuries, years, months, days, hours, etc. We may also locate relative to now, but ascribing past, present and future across a logically horizontal timeline.

The concept of now may be attributed to a specific time, but sine time is fleeting, I suspect that what human intelligence refers to as the present is not specifically an event, but a gross step along a process. In this manner, time is linked to the procedural network.

Procedural Subnet

The procedural network follows the strategies defined in the universal problem. That is, it starts with the idea of action, then immediately breaks into the concepts of change and prevention of change, followed quickly by a layer that are clearly delineated by the 7 strategies. Additional layers include more complex procedures that include more than one strategy (such as a purchase, which involves the flow of money in one direction and the flow of goods in

the other. I believe that these are filed under the last strategy (the ultimate result of the procedure). It is noteworthy to say that the structure I have defined to this point is a hierarchy, but the steps involved in a complex procedure are themselves a sequence, relative to the current step.

There is likely to be another level of external index in the procedural sub-network, where the item of contention (and possibly other elements from the universal problem) further flavors the procedure. For example, for me in these days of electronic book readers, the purchase of a book starts with me searching my favorite book seller online for a downloadable version.

Social Subnet

We finally come to the last of the subnetworks, the social subnet. This is of ultimate importance for at least two reasons: primarily, from this network comes the Entity with the problem, from which all information about the problem begins. In addition, people and organizations are often tools that serve as problem solvers for us.

I suggest that from the top down, the social network is organized into organizations. This is for two reasons, we tend to refer to organizations as entities, using the plural personal pronouns (we, us, they, them) and also because we ourselves are often organized into hierarchal organizations. Yet the principle elements found in the social network are individuals. However, these are organized not hierarchically, but sequentially entered around a brain cell representing the "self" the center of an ever expanding ring. This arrangement is crucial. We determine who we know based upon their proximity to us, or secondarily based upon the organization in which they partake. Before we will even speak to another individual or certainly solve their problems, we have to know them. In conventional computer terms, they have to be validated at login that we know who they are.

Roles within organizations are not included in the social network but rather the invented part of the elemental network. It should be noted that the types of organizations are within the invented part of the elemental network, and it is only individual organizations that are in social.

Part 3: The Components of Intelligence

The importance of the ten recursions of intelligence is that they may be used in various combinations to produce *any* of the components that collectively comprise intelligence.

The previous part of this work represented the material discussed in the brief time allowed by the SHARE presentation. In this section, we want to begin applying the concepts.

In this part of the paper we're going to do a process walk-through and examine the individual thought programs that represent the core of human intelligence. It is my assertion that these programs, once created, can effectively create their own thought programs going forward.

Any problem solving could occur through a number of routes, but we use diaductive reasoning as our starting point.

Please note that this walkthrough, like the previous material in this paper, is only a summary, simplified for faster reading.

BRAIN:Power – Reasoning

For completion, it might make some sense to briefly summarize diaductive reasoning, and the thought program (BRAIN:Power) through which it's defined to the artificial BRAIN.

BRAIN:Power consists of three questions:

1. Do I know the answer already?
2. Can I determine the answer externally?
3. Can I build the answer internally?

The first question is answered by the BRAIN:Pan thought program.

BRAIN:Pan – Memory Retrieval

The importance of the seven knowledge subsets is that it makes it easier to find knowledge especially when posed as a question. For example, when the question is “who?”, the answer is limited to the social (who) subset.

BRAIN:Pan further narrows the pool of answers by maintaining contexts. This is easier to achieve because there are contexts for all seven categories of knowledge relative to all involved parties as well as additional contexts created during language. The context is hierarchical, and is merely a pointer to a range of pre-existing knowledge within one of the seven subnets.

But the thought program also “pans” left or right, that is, sequentially through using axons to seek the answer within the hierarchical context. Hence the name, BRAIN:Pan.

You’ll notice that there are two alternate means of indexing here. The “vertical” hierarchical approach limits the knowledge to a specific range. The “horizontal” sequential approach should cross the vertical range of knowledge. Although there are certainly more than these two dimensions within human knowledge, locating a specific cell is much like locating a singular point on a Cartesian plane.

This offers a unique opportunity for further neurological study: Consider the much popularized left brain/right brain discussion (that has fallen out of favor because it was greatly misinterpreted in popular media). The so-called strengths of the dominant side were logical and linguistic, both operations that rely heavily on the hierarchical. The social and artistic characteristics of the alternate side rely heavily on proximal relationships that may be expressed in a sequence, such as relative positions and similar nearness socially. While I recognize that the characteristics expressed may be simplistic, I am offering a possibility that is even simpler. I suspect that the different sides of the BRAIN tend to link knowledge in two separate ways (hierarchically and sequentially) and communication between the two sides across the corpus callosum may be used to link the two mirrored points like a single point on a plane.

If BRAIN:Pan fails to locate knowledge, it will take one additional step before returning to the reasoning logic: It will look to the knowledge itself to produce a potential expert. This is important. I assume that knowledge of subject matter experts in the human brain is linked to the subjects within which they are expert. Such an approach leaves an alternative pathway.

In such a course, BRAIN:Pan returns its inability to find the answer, the “expert” value to the reasoning program and the subset of knowledge within which it was seeking. BRAIN:Power receives this information and determines the next course of action. Two paths remain: The BRAIN will have to seek knowledge externally or create it internally. While there is nothing that suggests these two events could be done in parallel or in a different order, we will pursue the external path first. BRAIN:Power will ultimately have the decision as to which “expert” to choose. The expert could have already been established (like universal problem “entity” which is in effect the expert of a given problem), the expert returned by BRAIN:Pan or a default domain expert like a dictionary or encyclopedia.

BRAIN:Dump – Outgoing Language

To seek an answer externally in the human space requires language. One will need to ask questions and understand the answers, or to read the expert document.

Natural Language Processing has long been the Holy Grail of AI, much sought, but extremely illusive. Broad-Reasoning has a dramatic change in approach that even rewrites archaic grammar rules in order to understand language in a truly natural manner. Under Broad-Reasoning concepts, language is not a separate function, but a natural and inevitable extension of human intelligence.

You will note that I have separated language into incoming and outgoing components. While there is a functional purpose for this as there are different processing requirements (questions) for each, it is a divide that exists also in nature. According to current neuroscience, the section of the human brain called “Wernicke’s area” is engaged in *incoming* language, where “Broca’s area” appears to be the section where *outgoing* language is created.

I have chosen the inelegant name “BRAIN:Dump” for the outgoing language thought program. It has three modes in which it operates:

1. Interrogative (creation of questions).
2. Conversational.
3. Expository/Prosaic/Creative.

The creation of questions is generally simple as the BRAIN:Pan has already returned the area within which it seeks and seeks to clarify. A sentence such as “where did she come from?” might be common.

Conversational language cannot be understood until I investigate incoming language below, and longer writing styles common to the last category are truly out of scope for this work. Needless to say it is an expansion of the conversational style using well-defined and much taught norms that fit into the questioning process at the heart of Broad-Reasoning.

So to understand language, I find it best to understand *incoming* language within the scope of Broad-Reasoning.

BRAIN:Scan – Incoming Language

Human language will initially be the principal means of presenting a problem to the BRAIN, and it's also necessary to understand the *answers* given from external sources, whether they are conversational or textual.

The BRAIN thought program³ in this space is termed "BRAIN:Scan", although it does far more than scan for keywords. It should be capable of truly understanding language within context.

Yet first we have to understand language in a larger sense: The problem language is trying to solve, of course, is communication. When considered within the context of Broad-Reasoning theory, what it's communicating is knowledge. While both of these statements are obvious, seldom does the theorist take the next step to realize that in communicating knowledge, *language is intent upon changing the structure of the listener's brain.*

Allow me to propose that in another way: Language specifically alters the circuitry between neurons. This is one of the reasons that NLP will not work without considering the larger relevance of knowledge and general intelligence.

It is within this context that grammar begins to change. We have been led to understand that the nouns are the most critical element among the eight parts of speech and in sentence deconstruction. Given our proposal that each thing is represented by a neuron, certainly nouns are important. Nouns are neurons. Yet, I am not suggesting that knowledge changes the nouns, but alters the connections *between* them. With this consideration, it is the verbs that become critical.

Consider: there are two types of verbs. The cupola verbs, couple. These are words like "is" and "was". In the sentence, "John is old", a neuron representing John is coupled with a neuron representing his aged condition. It is effectively the axon/dendrite connection between the two cells.

On the other hand, the action verbs may be described by a change or obstruction to change that may be represented in a flow between cells. The verb "to throw" describes a noun in motion. In fact it connects three cells, the thrower, the object thrown and the destination. It is *central* to these three elements.

When the word "throw" is recorded as a model cell, it leaves questions as to who is doing the throwing, what is being thrown and where it is being thrown if that's relevant. What this means is that the verb is the critical word in the sentence. If we can determine the verb, we have the questions at the heart of diaductive reasoning. Therefore, BRAIN:Scan may have some common elements similar to any sentence diagram as it will try to identify verbs, nouns, and other parts of speech, but it starts from the verb, and answers the questions that the verb implies, altering the structure of the cells accordingly.

³ While it is largely a thought program, BRAIN:Scan includes some conventional programming that should speed up the process of conversational language.

The source of knowledge for language is twofold: Certainly the incoming language stream is processed in temporary knowledge until it can be selected for more permanent changes. Yet more importantly, the language itself is encoded throughout the brain in what we generally call vocabulary. Yet I propose this is not in the form of separate knowledge, but the entirety of declarative knowledge within the human brain. The BRAIN concept of actually recording the words within the soma of the BRAIN:Cell, simplifies the artificial association of knowledge to vocabulary.

The parts of speech also fit the larger structure of knowledge suggested earlier in this paper. We've already suggested that the seven types of nouns can be simply placed within the seven subsets of knowledge. Using this as a model I would suggest that prepositions fit neatly into the sequential aspect of the positional subnet, while adjectives and adverbs are relative to the qualitative subnetwork. Verbs are not only descriptive of the procedural subnetwork, they may be defined by the very mapping of change according to their cells.

The remaining parts of speech are functional. When considered in terms of conventional computer communication, conjunctions link message units or expand message unit fields. Pronouns and interjections simplify communication much like compression does in data messaging. They are variables that may be expanded into greater meaning.

Furthermore, the sentence, which is the message unit of language, maps beautifully to the universal problem. Accordingly, each sentence is a problem that needs to be solved. To determine the vocabulary from a sense of parts of speech or meaning is also nothing more than a series of questions in need of answering. (I.e.: Is the word preceded by a determiner? If yes, then it is a noun.)

Certainly, this is a simplified version of human language, limited by the bounds of this paper. I would welcome the opportunity to explain the entire range of language considered within the scope of Broad-Reasoning theory.

BRAIN:Wave – Perception

I started with language as an input stream, because it will be the principal method for external information to be presented to the BRAIN for processing in most early situations. However, I assume that the BRAIN will also be used in robotics, surveillance and other procedures. Toward that end, Broad-Reasoning has a separate “perception” component that is intended to perceive other elements that normally arrive through the standard human senses. Moreover, it also fits into the diaductive reasoning process.

Consider that there is not always time to ask an expert. Thousands of times a day, we solve our own problems, small though they may be, in an immediate manner. Consider the challenge of simply walking through your living room. Perception is another means to answer questions, but to answer them *internally* with external information.

The BRAIN component is called “BRAIN:Wave” attributing to its intent to pull observations from the sensory world much like a radio tuner pulls radio waves from the “ether”.

The approach requires that external sensory devices (video cameras, microphones, etc.) are registered to the BRAIN through driver programs much like the process of conventional peripheral equipment in current computing architectures. The input is then linked to the qualitative subnet in a hierarchical approach (colors to colors, pitch to pitch, etc.).

As always with perception, the challenge is matching sensory signals with objects. Yet Broad-Reasoning suggests an approach that limits the workload implied to a reasonable amount of processing. I proffer that there are only two things that perception does:

- Finds the things we seek.
- Finds what doesn’t belong.

Consider that when you drive down the highway, you do not notice every sign signaling a gas station. These are largely ignored until the gas meter beeps at us. At that moment they arise from the landscape more clearly. Additionally, as you travel down the highway, your speed causes the distance in-between you and the next car to suddenly meet whatever criteria your brain maintains as dangerous. You may not even notice traffic until brake lights suddenly come on in front of you.

The ability to limit perceptions to immediate context suggests an action taking place: Both items require knowing *where we are* as contextual knowledge. Stated another way, the “where” or positional subnet should contain linkages to objects and the relative positions *between* them. The positions could be specific to your living room, or general contexts such as highways within your experience. In such manner we can navigate through the darkened living room in the middle of the night, or may suddenly notice when our spouse has bought a new chair.

I would further suggest that there is evidence of this process when the short-term memory is temporarily lost due to unconsciousness following a blow to the head. When injured recovers and the human brain “reboots”, the cliché first *question* asked is not “what happened”, but “where am I”. It is my assertion that the context of where we are is necessary for the act of sensory perception to occur.

BRAIN:Child – Logic

Not all problems are immediately solved through perception. The more challenging problems require the application of logic and creativity and perhaps even a little intuition. In this section, we start with logic.

The problem with logic is that it goes by many definitions, some of which are met with partially or completely in the other BRAIN components already discussed. Given problem solving as the sole operational recursion of human intelligence, I would contend that all logic is problem-solving logic.

You have already met the principal part of the BRAIN thought program which I call “BRAIN:Child”. The principal set of questions in the thought program begins with the universal problem model. BRAIN:Child is intent upon mapping problems to this model, and answering the questions implied. The “child” moniker refers to the method it uses to divide the problem elements into their subsets, especially the “contentious item”. It seeks or even *defines* the item’s children (asking questions and using comparison and contrast against qualities) to determine if there are truly multiple problems to be sought. Finally it manages the *process* of problem solving.

Toward that end, it even maps toward the solution. Initially, it uses the methodology defined in part one of this document, using the item location and desirability to map the desired change. Then it tries to implement that change across a network of cells. Each blockage produces another sub-problem to be solved in the same method.

In many ways, BRAIN:Child is deductive reasoning. However, when deductive reasoning fails, another method may be called.

BRAIN:Storm – Creativity

Creativity is another notoriously challenging element to define. Yet our Broad-Reasoning theory thus far establishes a foundation from which we can leap.

In short, I propose that creativity redefines the starting assumptions of any problem. This is done through a number of techniques that may be rendered within the architecture proposed to this point, but they always move away from the original brain cell reflected in the assumption.

The elements that are changed are primarily reflected in the cells connected to the universal problem, and most commonly the contentious item, but creativity can be assigned to any cell.

Some of the more common techniques follow:

- **Induction:** Inductive reasoning moves in the opposite direction of deduction. Rather than moving *down* the knowledge hierarchy to lesser and lesser cells, it moves upwards to be more inclusive. Comparison and contrast of cellular qualities are used to find commonalities among sibling cells, previously overlooked. Broad-Reasoning theory is greatly the result of this process, seeking similarities in the processes of intelligence, rather than dividing it into disparate components. Like so many things in Broad-Reasoning, the process was cyclical causing me to look higher with each pass.
- **Parallel Metaphor:** Many of my insights into Broad-Reasoning were made as the result of parallel metaphors (often in nature) that revealed greater understanding. The hierarchical nature of knowledge came to me when staring upward into a tree. My understanding of the question at the heart of intelligence came when I was trying to understand how language worked by reading a grammar book in 1999. The author used questions to get the reader to determine if a word was a noun or verb, etc. Certainly stories abound of similar discovery in scientific endeavor, and the reader likely has his or her own tales to tell in this respect. In the BRAIN, parallel metaphor may be established by seeking collections of BRAIN:Cells that have similar axons connecting them. I believe this is the process that occurs when we dream.
- **Opposition:** In the BRAIN, this technique uses the antonym axon to locate the polar opposite concept. I used this in my beer cooler example where I looked not only to keep the cold in but the warmth *out*. (Which may result in other techniques such as a lighter color on the outside.) There are certain strategies in the universal problem that have opposites (consider “the best defense is a good offense”) that work well with this approach. Yet it may be used with any cell that has an opposite.
- **Substitution:** Substitution is a means to experiment on the hopes of producing a reasonable observation that may inform upon the original problem. Typically a sibling cell would be used, but there are other approaches.
- **Deflection:** This method takes place within forward progress along a sequence of cells. Deflection uses a junction at a cell to go arbitrarily in another direction. One can think of it as “the road not taken” approach.

- **Others:** This list is not exhaustive, but only reflects some of the more easily defined techniques. Others that need little description would include randomness for example.

In each case, creativity returns to the original problem with changed assumptions, resulting in a new or parallel approach.

BRAIN:Surgeon – Intuition

The last approach to answering questions internally is often called human intuition. Intuition is employed when an assumption is good enough, and might even be considered another technique of creativity, above.

It is my assertion that intuition is not magic, but merely a collection of processes that are so hard-wired into the structure of our human brain that they operate below consciousness. Originally, I had organized a separate component (called “BRAIN:Surgeon”) that represented collective intuition. The “surgeon” term was used because it stitched knowledge together.

It has since occurred to me that the elements of intuition are already included in the other thought programs above. This hasn’t altered the concept, but merely the location where they execute.

As a result, BRAIN:Surgeon has become relegated to a process that runs when the system is otherwise idle, identifying holes in knowledge, and attempting to answer the questions they pose. In some way, it might be considered a curiosity engine.

BRAIN:Trust and BRAIN:Fever – Confidence and Emotion

The common view of an artificially intelligent computer is one that is devoid of emotion, as if emotion is some pariah preventing logical thought, an evolutionary mistake, if you will. Furthermore, the typical AI (at least as far as science fiction is concerned) is one that is superior to human intelligence, above all seeming lack of human confidence.

I would propose that these assumptions are simply wrong. Human confidence and emotion are a fundamental part of the intelligence process. I do not believe anything that is the result of millennia of evolution is a mistake, but rather, has been finely tuned over the same period of years.

I join them together in this chapter because I believe that they operate along the same lines, achieving similar results. I propose that both emotion and confidence affect the strength of the connections between neurons. This is employed in two ways in both artificial and natural intelligence:

- They're used to qualify the value of the *declarative* knowledge.
- They're employed in *procedural* knowledge to quantify potential paths.

Additionally, although they are both thought programs (with questions to answer), they express their results in a mere *number* and this is done through conventional programming. As a result, you may consider them to be hybrid programs. Confidence is expressed as a decimal between 0 and 1 (zero is no confidence). Emotion is expressed as a number between -1 and 1. Negative numbers are negative emotions such as fear, while positive numbers are positive emotions.

Confidence is largely confidence in the value of the knowledge. In the human counterpart, I would suggest that the strength between two procedural cells is increased by the number of times that it has been transited. In other words, "I know this will work, because I've done it many times before". In declarative knowledge it depends greatly upon the number of hops (cells) it takes to get the knowledge to you. If someone you know and trust provides the knowledge, it is strong, if it is an anonymous posting on a web site, it is weak.

Emotion is tied to the prioritization scheme defined in the motivational hierarchy (in the dispositional *why* network). Negative emotions result in negative numbers relative to their height in the hierarchy. Life and death fears, for example, receive a higher rating than employment-related fears.

BRAIN:Trust and BRAIN:Fever include algorithms that produce the numbers, and the numbers are included as two separate fields on the axon between cells. They are often considered in relationship to parallel paths. (i.e.: I trust this approach more than this one, or I trust this knowledge more than this.)

Without confidence and emotion, there would be no means to choose between elements of knowledge or courses of action.

Part 4: Removing the Roadblock of Consciousness

You may find it odd that I've saved the topic of consciousness for last because it has long been a principal barrier to both intelligence theory and AI. The assumption has long been that consciousness is necessary for intelligence.

I want to suggest the opposite.

Imagine yourself driving. Suddenly a truck pulls out from a side street into you're lane of travel. Do you swerve? Of course you do. Moreover, you would do it before you consider the question I posed consciously. It's a mere *reaction*.

If you could react to such a high priority, life-threatening situation without the intervention of consciousness, isn't it possible that consciousness is not required for intelligence?

Then let me propose an alternative to the assumptions we have all made in the past. Is it possible that consciousness is not required for intelligence, but merely an *observer* toward its actions? Let me state that in another somewhat horrific example: If someone were to break into your office and kill your colleague while you watched, would that make you a murderer? Of course not. Neither does consciousness' presence at the scene of the crime make it the criminal.

In order to remove the consciousness obstruction, Broad-Reasoning theory takes the approach that consciousness may be cleaved into two parts:

- There are obviously a number of components that are necessary for intelligence to function. We can call this set of elements "**intellectual consciousness**". Naturally, it is necessary for Broad Reasoning or any theory of intelligence to address such elements, and I have already done so within these pages.
- All remaining roles of consciousness we can place in the set of "**ethereal consciousness**". As fascinating as these are, we may set these aside for later debate. By our definition they are not a functional requirement of intelligence and hence of little interest to our immediate goal.

In such a manner we may dispatch the roadblock by reducing the consciousness question to manageable number of simpler elements. Toward that end, we can identify the following components of intellectual consciousness that we have largely addressed herein:

- **Problem Prioritization:** We typically make conscious decisions as to which problems we will pursue within our limited resources. As you may recall, we've used the human "reason" hierarchy (akin to Maslow's hierarchy of needs) to prioritize the problem queue pointed to by the universal problem model. Psychology calls this "attention".
- **Selection of Potential Solutions:** Similarly, we may have multiple paths to pursue a single solution. In the more resource intensive problems, consciousness appears to take part in this decision. Yet I have proposed that emotion (BRAIN:Fever) and confidence (BRAIN:Trust) aid in this determination.
- **Perception:** This term is often used as a synonym for consciousness. Yet, you have already read my view that perception (the BRAIN:Wave thought program) is intent upon, not only translating perceived elements to objects, but understanding their immediate importance.
- **Self-Awareness:** This term is also often believed to be analogue to consciousness because to perceive, one is truly perceiving the external environment as it relates to one's self. This perception can be extended to include awareness of oneself in the wider, more spiritual context. Yet we have defined a BRAIN:Cell called the "self" cell within the scope of this project. It is central to the sequential relationships within the social subnet.
- **Conscience:** Some reviewers have expressed to the author that conscience is only possible where there is consciousness. While this is debatable (and I think it is the similarity between the words that creates a false linkage), it is included here for completeness. I would suggest two alternatives: In one aspect, conscience is a set of rules. Rule-based artificial intelligence has been suggested for decades, although I don't include it here. Instead, I suggest a means for understanding the human disposition toward change as expressed in the hierarchy of needs. Certainly this provides a means for empathy, truly understanding the human condition and its priorities.

As a part of the peer review process, we encourage identification of other areas where consciousness and intelligence intersect; however, we end with the working assumption that the above is a functional list.

If this list has truly been addressed, then consciousness is not an issue. As such, we are free to pursue an artificial intelligence, unencumbered.

Toward that end, I encourage you to build the BRAIN in your basement or lab. All I ask is that I be present when it wakes up and asks its first question:

"How may I help you?"