

SECRET

Final Report--Objective A, Task 4

December 1986

REMOTE VIEWING EVALUATION TECHNIQUES (U)

By: BEVERLY S. HUMPHREY VIRGINIA V. TRASK
EDWIN C. MAY MARTHA J. THOMSON

Prepared for:

PETER J. McNELIS, DSW
CONTRACTING OFFICER'S TECHNICAL REPRESENTATIVE

CONTRACT DAMD 17-85-C-5130

WARNING NOTICE

RESTRICTED DISSEMINATION TO THOSE WITH VERIFIED ACCESS
TO THE [REDACTED] PROJECT

SG1A

333 Ravenswood Avenue
Menlo Park, California 94025 U.S.A.
(415) 326-6200
Cable: SRI INTL MPK
TWX: 910-373-2046



SECRET

UNRELEASABLE TO
FOREIGN NATIONALS

SECRET

*Final Report--Objective A, Task 4
Covering the Period 1 October 1985 to 30 September 1986*

December 1986

REMOTE VIEWING EVALUATION TECHNIQUES (U)

By: BEVERLY S. HUMPHREY
EDWIN C. MAY

VIRGINIA V. TRASK
MARTHA J. THOMSON

Prepared for:

PETER J. McNELIS, DSW
CONTRACTING OFFICER'S TECHNICAL REPRESENTATIVE

CONTRACT DAMD 17-85-C-5130

SRI Project 1291

SG1A

WARNING NOTICE

RESTRICTED DISSEMINATION TO THOSE WITH VERIFIED ACCESS
TO THE [REDACTED] PROJECT

Approved by:

ROBERT S. LEONARD, *Executive Director*
Geoscience and Engineering Center

Copy 11 of 15 Copies.

This document consists of 37 pages.

SRI/GF-0291

CLASSIFIED BY: HQ, USAMRDC (SGRD-ZA)
DECLASSIFY ON: OADR

SECRET

NOT RELEASABLE TO
FOREIGN NATIONALS



UNCLASSIFIED

ABSTRACT (U)

(U) A simplified automated procedure is suggested for the analysis of free-response material. As in earlier similar procedures, the target and response materials are coded as yes/no answers to a set of questions (descriptors). By definition, this coding defines the complete target and response information. The accuracy of the response is defined as the percent of the target material that is correctly described (i.e., the number of correct response bits divided by the number of target bits = 1). The reliability of the response is defined as the percent of the response that is correct (i.e., the number of correct response bits divided by the total number of response bits = 1). The figure of merit is the product of the accuracy and reliability. The advantages and weaknesses of the figure of merit are discussed with examples. Mean chance expectations (MCE) are calculated for the figure of merit, and recommendations are made to extend current techniques and to explore new technologies.

UNCLASSIFIED

UNCLASSIFIED**TABLE OF CONTENTS (U)**

ABSTRACT	ii
LIST OF ILLUSTRATIONS	iv
LIST OF TABLES	v
I INTRODUCTION	1
II METHOD OF APPROACH	3
A. Figure of Merit Analysis	3
1. Overview	3
2. Mathematical Formalism	3
a. Definitions	3
b. Linear Least-Squares Analysis	6
c. Mean Chance Expectation	7
d. Probability Assessment and Analysis	8
B. Fuzzy Set Theory--An Enhancement Technique for Descriptor List Technology	11
1. Overview	11
2. A Tutorial	12
3. Initial Application to RV Evaluation	14
4. Potential Future Applications	14
III RESULTS	16
A. Inter-Analyst Reliability Factors	16
B. Response Definition: Descriptor List Formulation	19
1. Novice Response Descriptor List	19
2. Advanced Response Descriptor List	20
C. Target Definition: Implications for Target Pool Composition	25
IV RECOMMENDATIONS	27
A. Similarity Experiment	27
B. AI Techniques	28
C. In-House Effort	29
V CONCLUSIONS	31
REFERENCES	32

UNCLASSIFIED

UNCLASSIFIED

LIST OF ILLUSTRATIONS (U)

1.	An MCE Figure of Merit Distribution	9
2.	The Fuzzy Set "Very Young"	13
3.	Comparison Between Types of Analysts	19
4.	Baylands Nature Interpretive Center, With RV Response	24

UNCLASSIFIED

UNCLASSIFIED

LIST OF TABLES (U)

1.	Descriptor-Bit Definition	4
2.	Viewer No. 454 Results Coded by Four Analysts	18
3.	Candidate Abstract Descriptors for Novice Responses	21
4.	Comparison of Target vs. Response Coding for "Baylands" Target	22
5.	Potential Problem Areas for Novice Targets	26

UNCLASSIFIED

SECRET

I INTRODUCTION (U)

(S) Since the publication of the initial remote viewing (RV) effort at SRI International*¹, two basic questions have remained in evaluating remote viewing data:

- What is the definition of the target site?
- What is the definition of the RV response?

In the development of meaningful evaluation procedures, we must address these two questions, whether the RV task is a research-oriented one (in which the target pool is known), or an intelligence-oriented mission (in which the target may not be known).

(U) In the older, IEEE-style, outbound experiment, definitions of target and response were particularly difficult to achieve. The protocol for such an experiment dictated that an experimenter travel to some randomly chosen location at a prearranged time; a viewer's task was to describe that location. In trying to assess the quality of the RV descriptions (in a series of trials, for example), an analyst visited each of the sites and attempted to match responses to them. While standing at a site, the analyst had to determine not only the bounds of the site, but also the site details that were to be included in the analysis. To cite a specific example using this protocol: if the analyst were to stand in the middle of the Golden Gate Bridge, he/she would have to determine whether the buildings of downtown San Francisco, which are clearly and prominently visible, were to be considered part of the Golden Gate Bridge target. The RV response to the Golden Gate Bridge target could be equally troublesome, because responses of this sort were typically 15 pages of dream-like free associations. A reasonable description of the bridge might be contained in the response—it might be obfuscated, however, by a large amount of unrelated material. How was an analyst to approach this problem of response definition?

(U) The first attempt at quantitatively defining an RV response involved reducing the raw transcript to a series of declarative statements called concepts.² Initially, it was

*(U) References are listed in order of appearance at the end of this report.

SECRET

SECRET

(U)

determined that a coherent concept should not be reduced to its component parts. For example, a *small red VW car* would be considered a single concept rather than four separate concepts, *small*, *red*, *VW*, and *car*. Once a transcript had been "conceptualized," the list of concepts constituted, by definition, the RV response. The analyst rated the concept lists against the sites. Although this represented a major advance over previous methods, no attempt was made to define the target site.

(S) During an FY 1982 program, a procedure was developed to define both the target and response material.³ It became evident that before a site can be quantified, the overall remote viewing goal must be clearly defined. If the goal is simply to demonstrate the existence of the RV phenomena, then anything that is perceived at the site is important. But if the goal is to gain information that is useful to the intelligence community, then specific items at the site are important while others remain insignificant. For example, let us assume that an office is a hypothetical target and that a single computer in that office is of specific interest. Let us also assume, hypothetically, that a viewer gives an accurate description of the shape of the office, provides the serial number of the typewriter, and gives a complete description of the owner of the office. Although this kind of a response might provide excellent evidence for remote viewing, the target of interest (the computer) is completely missed--this response, therefore, is of no interest as intelligence data. What is needed is a specific technique to allow assessments that are mission-oriented.

(S) This report describes a computerized RV evaluation procedure that was initially developed in FY 1984⁴ and has been expanded and refined in FY 1986.* In its current evolution, it is an analysis that has been aimed primarily at simpler, research-oriented tasks using a known target pool. It is anticipated, however, that future refinements to existing procedures, in addition to the advances of proposed new technologies, will allow evaluation techniques to begin to address the more complex issue of operational RV intelligence collection.

*(U) This report constitutes Objective A, Task 4, "Remote Viewing Evaluation Techniques."

UNCLASSIFIED

II METHOD OF APPROACH (U)

A. (U) Figure of Merit Analysis

1. (U) Overview

(U) Current approaches in evaluation technology have focused on the refinement and extension of the figure of merit analysis.⁴ Defined in general terms, this procedure generates a figure of merit (M) between 0 and 1, which provides an accurate assessment of an RV response. The M is the product of the accuracy and reliability with which an RV response describes its correct target, as determined by an analyst's coding of RV targets and responses according to a "descriptor list." Table 1 provides a representative example of such a list, which was used in an FY 1986 novice RV training program. Each of the items in a descriptor list requires a binary decision from the analyst as to the item's presence or absence in each of the targets and responses. The mathematical formalism for converting the analyst's binary codes into Ms and their controls is detailed in Section A.2. below.

2. (U) Mathematical Formalism

a. (U) Definitions

(U) For a single viewer, the overall method of analysis consists of calculating a figure of merit, M, for each viewing session, and then comparing these Ms to a control set of figures of merit.

UNCLASSIFIED

UNCLASSIFIED

Table 1

(U) DESCRIPTOR-BIT DEFINITION

Bit No.	Descriptor
1	Is any significant part of the scene hectic, chaotic, congested, or cluttered?
2	Does a single major object or structure dominate the scene?
3	Is the central focus or predominant ambience of the scene primarily natural rather than artificial or manmade?
4	Do the effects of the weather appear to be a significant part of the scene? (e.g., as in the presence of snow or ice, evidence of erosion, etc.)
5	Is the scene predominantly colorful, characterized by a profusion of color, by a strikingly contrasting combination of colors, or by outstanding, brightly-colored objects (e.g., flowers, stained-glass windows, etc.--not normally blue sky, green grass, or usual building color)?
6	Is a mountain, hill, or cliff, or a range of mountains, hills, or cliffs a significant feature of the scene?
7	Is a volcano a significant part of the scene?
8	Are buildings or other manmade structures a significant part of the scene?
9	Is a city a significant part of the scene?
10	Is a town, village, or isolated settlement or outpost a significant feature of the scene?
11	Are ruins a significant part of the scene?
12	Is a large expanse of water--specifically an ocean, sea, gulf, lake, or bay--a significant aspect of the scene?
13	Is a land/water interface a significant part of the scene?
14	Is a river, canal, or channel a significant part of the scene?
15	Is a waterfall a significant part of the scene?
16	Is a port or harbor a significant part of the scene?
17	Is an island a significant part of the scene?
18	Is a swamp, jungle, marsh, or verdant or heavy foliage a significant part of the scene?
19	Is a flat aspect to the landscape a significant part of the scene?
20	Is a desert a significant part of the scene, or is the scene predominately dry to the point of being arid?

UNCLASSIFIED

UNCLASSIFIED

UNCLASSIFIED

(U) Let $[n]$ be the number of sessions to be analyzed. Suppose also that the descriptor list contains $[m]$ bits. We then define the total number of bits in a specific response, $[k]$, as

$$R_k = \sum_{j=1}^m R_{jk} ,$$

where $R_{jk} = 1$ if bit $[j]$ were answered affirmatively and equal to 0 otherwise. Likewise the total number of bits in a specific target, $[i]$, as

$$T_i = \sum_{j=1}^m T_{ji} ,$$

where $T_{ji} = 1$ if bit $[j]$ were answered affirmatively and equal to 0 otherwise.

(U) The accuracy of response, $[k]$, (the percent of target $[i]$ that is described correctly) is given by

$$a_{ki} = \sum_{j=1}^m \frac{R_{jk} T_{ji}}{T_i} .$$

The reliability of response $[k]$ (the percent of response $[k]$ that was correct) is given by

$$r_{ki} = \sum_{j=1}^m \frac{R_{jk} T_{ji}}{R_k} .$$

(U) Finally, the figure of merit for response $[k]$ matched against target $[i]$ is given by

$$M_{ki} = a_{ki} \times r_{ki} .$$

(U) The analysis can be considered from two perspectives: matches--i.e., the figure of merit is calculated by matching a response against its intended target (i.e., $k = i$), and cross-matches--i.e., the figure of merit is calculated by matching a response against some target other than its intended one (i.e., $k \neq i$).

UNCLASSIFIED

UNCLASSIFIED**b. (U) Linear Least-Squares Analysis**

(U) After [n] remote viewing sessions have been completed and the analysis described above performed using $k = i$, there are [n] figures of merit, one for each RV session, in order of session number. To examine if there are any systematic variations within this data, a best-fit straight line is fitted through the figures of merit using standard techniques. If [x] is the session number ($x=1, \dots, n$), consider a straight line defined as

$$M(x) = a + b(x - \bar{x})$$

where $a = M(\bar{x})$ and $\bar{x} = \text{constant}$, and [a] and [b] are the intercept and slope, respectively. Suppose there are [n] pairs of points, (x, M_x) . Then the slope, which is calculated by a standard least-squares technique, is given by

$$b = \frac{n \sum_{x=1}^n x M_x - \sum_{x=1}^n x \sum_{x=1}^n M_x}{\Delta}, \quad \text{where } \Delta \text{ is given by}$$

$$\Delta = n \sum_{x=1}^n x^2 - \left(\sum_{x=1}^n x \right)^2$$

The intercept is given by

$$a = a' + b \bar{x}, \quad \text{where } a' \text{ is given by}$$

$$a' = \frac{\sum_{x=1}^n x^2 \sum_{x=1}^n M_x - \sum_{x=1}^n x \sum_{x=1}^n x M_x}{\Delta}$$

If we set $\bar{x}$ = to the average value of the session number, then

$$\bar{x} = \frac{1}{n} \sum_{x=1}^n x,$$

and [a] becomes the average value of the figure of merit. Thus

$$a = \bar{M}.$$

UNCLASSIFIED

UNCLASSIFIED

(U) This analysis produces the average figure of merit for a series and an indicator, the slope, of learning.

c. (U) Mean Chance Expectation

(U) The calculation for the mean chance expectation (MCE) must be sensitive to a number of possible artifacts or confounding factors:

- Viewer variations (e.g., viewers' different response biases).
- General knowledge of the target pool (e.g., targets are known to be *National Geographic* magazine photographs).
- Specific knowledge of the target pool resulting from trail-by-trial feedback.
- Methodological considerations (e.g., viewers are asked to respond with more data at the end of the series compared to what was asked of them at the beginning of the series).

All of these factors will affect the expected average figure of merit and any session-to-session systematic variation that may be present.

(U) A method for determining the figure of merit MCE, which requires the fewest number of assumptions about the structure of the data or the response biases, involves the cross-matching of all the responses to the same target set used in the series in question. A cross-match is defined as a comparison between a response and a target other than the one used in the session. If a figure of merit distribution is calculated for a large number of cross-matches, a number of the confounding factors listed above will be addressed. To determine the session-to-session dependencies of the MCE, however, the session order must be preserved. By preserving the order of the responses and by calculating $[n]$ sets of cross-matches at a time, MCE figure of merit, slope, and intercept distributions can be calculated.

(U) As before, let $[n]$ be the number of sessions in a series for a single viewer. Also, let the order of the responses, R_k be preserved. Define $[N]$ as the number of cross-match cycles through the ordered set of $[n]$ responses. The MCE calculation proceeds as follows:

1. Randomly choose a target order, $i = 1, n$, such that $k \neq i$ where $k = 1, n$ is the preserved response order.
2. Calculate the figure of merit for the k th response/target cross-matches as

UNCLASSIFIED

UNCLASSIFIED

(U)

$$M_{ki} = a_{ki} \times r_{ki}, \quad k \neq i,$$

where a_{ki} and r_{ki} are the accuracy and reliability for response [k] cross-matched against target [i], respectively.

3. Do step 2 above for all [n] sessions.
4. Calculate a slope and intercept for the resulting figures of merit by the linear least-squares analysis described above.
5. Repeat steps 1 through 4 above for [N] cycles to produce MCE figure of merit, slope, and intercept distributions.

(U) It is important to note that MCE distributions are generated for each viewer and are *not* summed across viewers. Therefore, individual viewer differences in response "biases" are accounted for by definition.

(U) This procedure also accounts for general knowledge of the target pool by the viewer, because information learned by this method in a given session will *not* necessarily be associated with the intended target for that session. The net effect of this type of artifact will be to "bias" the MCE figure of merit distribution toward larger values.

(U) Because the order of viewings is preserved, any knowledge of the target pool that is learned by the viewer as a result of trial-by-trial feedback is accounted for in two ways:

1. Information resulting from increasing knowledge of the target pool will "bias" the MCE figure of merit distribution toward larger values.
2. Information resulting from increasing knowledge of the target pool as a function of session number will "bias" the MCE slope distribution toward larger values.

(U) Similarly, any artifact caused by methodological considerations as a function of session, will also "bias" the MCE figure of merit and slope distributions.

d. (U) **Probability Assessment and Analysis**

(U) There are a number hypotheses that can be tested using the various MCE distributions described above:

- An individual remote viewing is statistically beyond MCE.

UNCLASSIFIED

UNCLASSIFIED

(U)

- The series generated by a single viewer shows statistical evidence for remote viewing.
- There is evidence above MCE for remote viewing "learning."
- The mean of the observed figure of merit distribution is significantly larger than the mean of the MCE distribution.
- The observed figure of merit distribution is significantly different than the MCE distribution.

(U) Using the MCE figure of merit distribution, a straightforward calculation of areas will determine whether a particular figure of merit from a single session is significant. Figure 1 shows an example of an MCE figure of merit distribution described above. M_{kk} is the figure of merit resulting from session $[k]$. The probability of obtaining a figure of merit of M_{kk} or larger caused by artifact is the patterned area divided by the total area under the curve shown in Figure 1. This technique can be used to assess the chance likelihood for all sessions in a series by a single viewer.

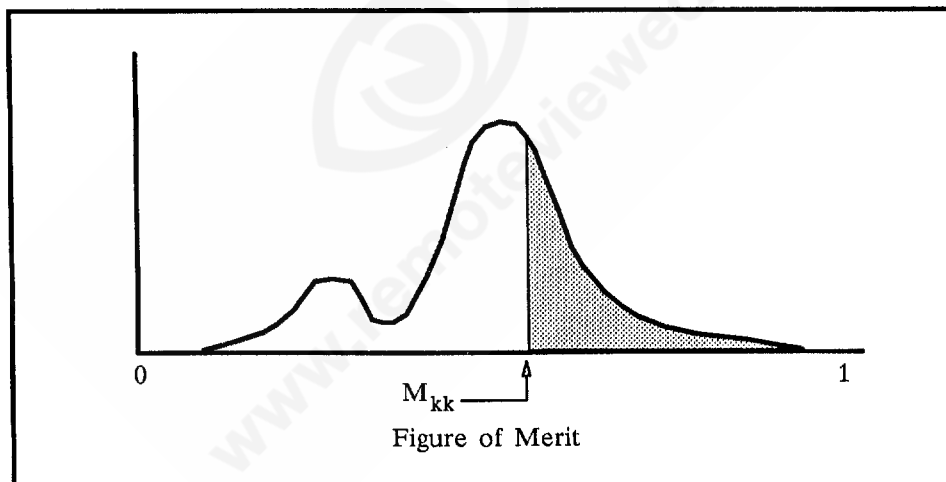


FIGURE 1 (U) AN MCE FIGURE OF MERIT DISTRIBUTION

(U) To determine whether there is statistical evidence of remote viewing within a given series for a single viewer, the p-values for the individual sessions must be combined. The primary method used for combining p-values was developed by Fisher.⁵ A X^2 with two degrees of freedom is computed for each p-value and summed. The resulting X^2 is evaluated with $2n$ degrees of freedom where $[n]$ is the number of p-values that were combined. If $[k]$ is the session number, the appropriate total X^2 is given by

UNCLASSIFIED

UNCLASSIFIED

(U)

$$X^2 = \sum_{k=1}^n -2.0 \ln p_k, \quad df = 2n$$

where the p_k are the p-values for each of the [k] sessions. A second technique involves testing the significance of the average p-value across all sessions. A standard z-score is calculated by

$$Z = \sqrt{12n} (0.5 - \bar{p}),$$

where $\bar{p}$ is the average p-value and [n] is number of sessions.⁶

(U) These two measures are sensitive to different aspects of the remote viewing series. For approximately 20 or more sessions, the two techniques will yield similar probability estimates if there is slight, but *consistent* evidence of remote viewing. On the other hand, if there are a few *very* good results (i.e., individual p-values < 0.001), then the X^2 technique more accurately reflects the series as a whole.

(U) As an example of consistency, suppose 20 sessions having individual p-values of 0.35 each are analyzed. Then the z-score for the average p-value is 2.32, corresponding to a combined p-value of 0.01. The X^2 technique yields a total X^2 of 42.0 with $df = 40$, corresponding to combined p-value of 0.40. To illustrate the X^2 technique's sensitivity to "good" remote viewings, consider the following p-values for 5 individual sessions: 0.45, 0.72, 0.55, 0.001, and 0.00005. The average p-value technique yields a combined p-value of 0.11, while the X^2 technique yields a combined p-value of approximately 0.0005.

(U) To determine whether there is evidence of "learning" and whether the means of the actual and MCE figure of merit distributions are significantly different, an ANOVA technique is used.⁷ By transforming the data about the average value of the session number, the slope and intercept hypothesis testing may be done separately. The F-ratios (from the ANOVA) for the two tests are given below.

UNCLASSIFIED

UNCLASSIFIED

(U)

$$F(\text{slope}) = \frac{n (b - b')^2 \times \frac{n}{(n-1)} \left(\sum_{k=1}^n \frac{k^2}{n} - \bar{k}^2 \right)}{\Delta}, \quad df1 = 1; df2 = (n-2)$$

and

$$F(\text{intercept}) = \frac{n (a - a')^2}{\Delta}, \quad df1 = 1; df2 = (n-2),$$

where [a] and [b] are the intercept and slope from the remote viewing figure of merit data, and [a'] and [b'] are the intercept slope from the MCE figure of merit data. Δ is given by

$$\Delta = \frac{\sum_{k=1}^n M_k^2 + n a^2 + b^2 \sum_{k=1}^n k^2 + 2 a b \sum_{k=1}^n k - 2 a \sum_{k=1}^n M_k - 2 b \sum_{k=1}^n k M_k}{n-2}.$$

(U) Because the F-ratio for the slope for the figure of merit data is a statistical test between the observed slope and that computed from the MCE, it constitutes an estimate of the probability that remote viewing "learning" occurred over and above any contribution that might have occurred because of some artifact. The F-ratio for the intercept constitutes an estimate of the probability that the mean of the figure of merit distribution is different from the mean of the MCE. We use a standard X^2 measure, a more sensitive test, to determine if the observed figure of merit distribution is statistically identical to that of the MCE distribution.

B. (U) Fuzzy Set Theory--An Enhancement Technique for Descriptor List Technology

1. (U) Overview

(U) The figure of merit analysis is predicated on descriptor list technology, which represents a significant improvement over earlier "conceptual analysis" techniques, both in terms of "objectifying" the analysis of free response data and in increasing the speed and efficiency with which evaluation can be accomplished. It has become increasingly evident,

UNCLASSIFIED

UNCLASSIFIED

(U)

however, that current lists are inadequate in providing discriminators that are "fine" enough both to describe a complex target accurately and to exploit fully the more subtle or abstract information content of the RV response. To decrease the granularity of the RV evaluation system, therefore, it was determined that the technology would have to evolve in the direction of allowing the analyst a gradation of judgment about target and response features, rather than the current "all-or-nothing" binary determinations. A preliminary survey of various disciplines and their evaluation methods (spanning such diverse fields as artificial intelligence, linguistics, and environmental psychology) has revealed a branch of mathematics, known as "fuzzy set theory," which provides a mathematical framework for modeling situations that are inherently imprecise. The principal architect of fuzzy sets, L. A. Zadeh, has stated:

*"One of the aims of the theory of fuzzy sets is the development of a methodology for the formulation and solution of problems which are too complex or ill-defined to be susceptible to analysis by conventional techniques."*⁸

(U) Because the task of RV analysis requires human judgments about imprecise situations--namely, the categorization of natural sites and the interpretation of abstract representations of those sites--it would appear, according to the above definition, that fuzzy set theory is a promising line of inquiry. In the next section, some of the basic concepts of fuzzy set theory will be examined, with the aim of understanding how this technology might be applied to the specific problem of RV evaluation.

2. (U) A Tutorial

(U) In traditional set theory, an element is either a member of a set or it isn't--e.g., the number 2 is a member of the set of even numbers; the number 3 is not. Fuzzy set theory is a variant of traditional set theory, in that it introduces the concept of *degree* of membership: herein lies the essence of its applicability to the modeling of imprecise systems. For example, if we take the concept of *age* (known as a *linguistic variable* in fuzzy set parlance), we might ascribe to it certain subcategories (i.e., *fuzzy sets*) such as *very young*, *young*, *middle-aged*, *old*, etc. Looking at *very young*, only, as a fuzzy set example, we must

UNCLASSIFIED

UNCLASSIFIED

(U)

define what we mean by this concept vis-a-vis the linguistic variable *age*.* If we examine the chronological ages from 1 to 30, we might subjectively assert that we consider the ages 1 through 4 to represent rather robustly a spectrum of the concept *very young*, whereas the age of 30 probably does not accurately represent *very young* at all. As depicted in Figure 2, fuzzy set theory allows us to assign a numerical value between 0 and 1 that represents our best subjective estimate as to how much each of the ages 1 through 30 embodies the concept *very young*.

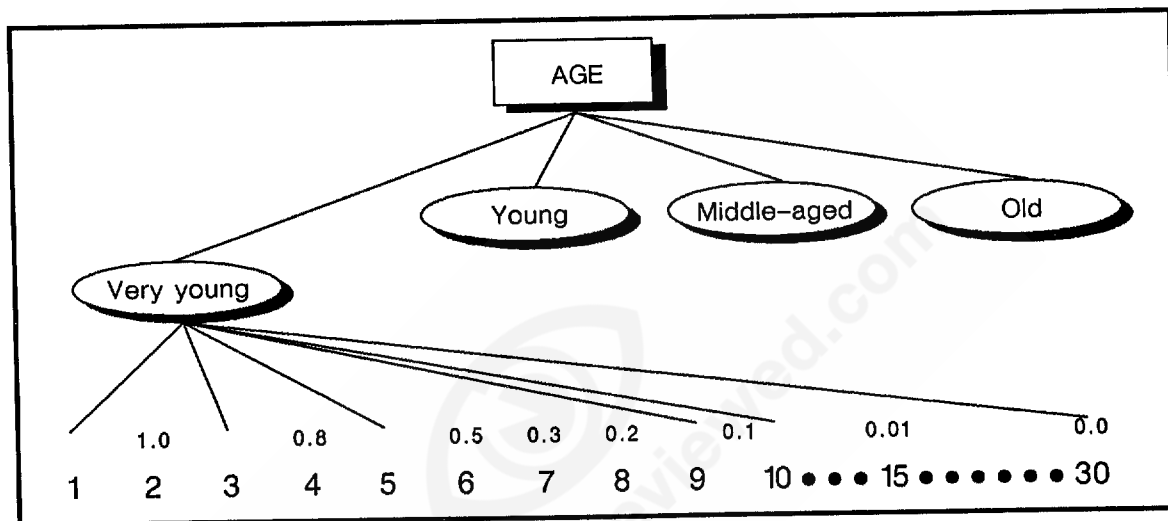


FIGURE 2 (U) THE FUZZY SET "VERY YOUNG"

(U) Clearly, a different set of numerical values would be assigned to the ages 1 through 30 for the fuzzy sets *young*, *middle-aged*, and *old*--e.g., the age of 6 might receive a value of 0.5 for *very young*, but a value of 1.0 for *young*, depending on context, consensus, and the particular application of the system. In this way we are able to provide manipulatable numerical values for imprecise natural language expressions; in addition, we are no longer forced into making inaccurate binary decisions such as, "Is the age of 7 very young--yes or no?"

* (U) It is important to note that the design of the fuzzy application occurs in accordance with the subjectivity of the system designer. Fortunately, the fuzzy set technology is rich enough that it allows for a virtually unrestricted range of expression. Technically speaking, *young* is the fuzzy set and *very* is a modifier, but it is beyond the scope of this paper to present terminology in depth.

UNCLASSIFIED

UNCLASSIFIED

3. (U) Initial Application to RV Evaluation

(U) In FY 1986, work began on an initial application of fuzzy set technology to RV evaluation, which simply entails an extension of the current descriptor list capabilities. In coding the targets, an analyst employs numerical values between 0 and 1, inclusively, to rate each of the 20 descriptors according to the importance of its visual representation in the target. For example, in rating a *National Geographic* magazine picture of the Khyber Pass, an analyst might ascribe a value of 0.80 for a "mountain" descriptor, a value of 0.20 for a "desert" descriptor, and values for other appropriate descriptors in accordance with their perceived importance to the target as a whole.*

(U) The rating of responses is considerably more subjective than the rating of targets. The analyst is required to apply a "confidence rating"--i.e., again, a value between 0 and 1, inclusively--as to what degree an abstract ideogram is representative of a given descriptor. For example, if a novice subject draws a conical-shaped object and labels it, "fuzzy cone...wider at the bottom..." the analyst may decide that there is some justification for interpreting this ideogram as a volcano covered with vegetation. Clearly, however, the confidence factor for making this highly subjective determination is quite low; the net result might be, therefore, that the "volcano" descriptor might receive a rating of 0.15, while "foliage" might receive 0.05.

(U) We anticipate that the primary effect of implementing this rudimentary application of fuzzy set technology will be to "fine tune" the figure of merit scores, such that they are more representative of the "true" information content of an RV response. The current figure of merit application penalizes certain responses and inflates others (especially given the "noisy" aspect of novice data), based on the correctness or incorrectness of the analyst's "all-or-nothing" determination with regard to any given descriptor. To summarize, fuzzy set technology is attractive in two important respects: (1) it affords the analyst a wider range of expression, thereby enabling him/her to provide a more realistic portrayal of the information contained in both targets and responses, and (2) it is compatible with the figure of merit mathematical formalism.

4. (U) Potential Future Applications

(U) It is anticipated that the initial application of fuzzy set technology to RV responses and targets will greatly enhance the accuracy with which their information content is

* (U) Coding of both targets and responses might be more "objectively" arrived at via the consensus of a group of experienced analysts.

UNCLASSIFIED

UNCLASSIFIED

(U)

depicted. A problem remains, however, with the inherently large granularity of the current descriptor list, which is independent of the potential "fineness" of its application allowed by fuzzy set theory--although the analyst will be allowed a gradation of response in interpreting an abstract ideogram (i.e., the confidence factors ranging from 0 to 1), he/she will still be constrained to interpreting that ideogram according to 20 *concrete* descriptors. These descriptors are significantly limited in their ability both to portray rich environments and to distill the most usable information from abstract RV mentations.

(U) It is projected that future descriptor lists will afford the analyst greater latitude in interpreting the more abstract aspects of RV responses, by providing *basis vector* descriptors. Such descriptors would represent, in essence, the lowest practicable common denominator of abstraction from which more concrete descriptors might be generated using fuzzy set operations (such as intersection and union). An example of a basis vector descriptor might be the concept of *vertical*, which is an abstraction that is represented to varying degrees in such concrete descriptors as *building*, *cliff*, *mountain*, *waterfall*, etc.

(U) Ultimately, we envision that evaluation would proceed along the lines of analyzing both the RV responses and targets in terms of fuzzy-weighted basis vector descriptors. A comparison of basis vector descriptors between responses and targets could then be effected, which would culminate in a figure of merit analysis reflecting the subject's ability to debrief the more abstract components of the psi signal. By using fuzzy set operations, concrete target and response descriptors could subsequently be generated on a "best fit" basis from the basis vector descriptors, and a figure of merit evaluation could be performed at this higher-order level also. The primary benefits of this type of procedure would be in providing objectification of abstract response data, and in affording more automated interpretation of these data in concrete terms. Furthermore, it would also be possible to track, in a systematic and quantifiable manner (on both a "subject-by-subject" and "across subject" basis), the kinds of abstract signals that subjects are receiving reliably; presumably, this capability might then be used to illuminate important lines of future investigation within RV fundamentals.

UNCLASSIFIED

UNCLASSIFIED

III RESULTS (U)

(U) The results of the FY 1986 evaluation effort have been obtained primarily from two sources: (1) identification of inter-analyst reliability factors, based on analysis of figure of merit statistics, and (2) insights into descriptor list formulations and target pool composition, based on *post hoc* analysis and observations. Each of these areas is explored in turn below.

A. (U) Inter-Analyst Reliability Factors

(U) A method was developed in FY 1986 for rating the abilities of potential remote viewing analysts. The most direct method of accomplishing this was simply to ask a candidate to analyze a known series of remote viewings; the results could then be compared to those produced by a proven analyst, 374.

(U) Three individuals, 432, 579, and 642, were asked separately to score first the targets and then the responses used in a remote viewing series from a novice viewer. They used a twenty-bit descriptor list (see Table 1) under a "blind" protocol. The procedure described in Section II.A. was used to calculate figures of merit, session p-values, and overall p-values for each analyst.

(U) Novice remote viewing data, which have been collected under our stimulus/response protocol, contains two distinguishing characteristics:

- The data tend to be sparse and abstract.
- The data tend to be noisy (i.e., large amounts of incorrect information).

(U) If the descriptor list contains mostly concrete items rather than abstract concepts (e.g., "Is there a waterfall?" versus "Are there vertical features?"), then an analyst who is unwilling or unable to interpret abstract and/or sparse data will miss whatever remote viewing information may be present. In the extreme case, a literal analyst may not answer any questions on the descriptor list affirmatively. If it is assumed that there was some remote viewing information present in the abstract response, then it is clear that the literal analyst will miss it. As the responses become less abstract and possibly more accurate, the difference between an interpretive and literal analyst becomes less important.

UNCLASSIFIED

UNCLASSIFIED

(U) Based upon these concepts, three hypotheses were formulated that could be tested with the three candidate analysts listed above:

- For an analyst to be sensitive to novice data, he/she must be willing to interpret abstract data.
- An interpretive analyst cannot demonstrate a remote viewing effect where there is none.
- For literal analysts, the difference between their p-values and those of an interpretive analyst will correlate significantly, on a session-by-session basis, with their own session p-values.

(U) The first hypothesis is true by inspection. Because the only remote viewing output that is analyzed is the one coded into the descriptor list, it follows that an analyst must interpret abstract data, or there are *no* data for analysis.

(U) Given that the analyst must be interpretive, we must consider whether an artifact could be induced by being interpretive. This is not the case. Because the analyst is "blind" to the correct target for a given session, there is no reason to expect that the interpretation of the abstract response would be selective in such a way as to match the intended target better than any other target in the series. Because the probability assessment of a single session involves the MCE cross-matched figure of merit distribution, any "enhanced" effects are canceled by the differential comparison.

(U) As a result, it was predicted that the difference between the means of the actual figure of merit and the MCE (ΔM) would reflect the remote viewing information content, and that the difference would decrease as the analyst tends to be more literal. Table 2 shows the results from four analysts in assessing the same 45 novice remote viewing sessions. It should be noted that $p(\Delta M)$ is the probability (derived from ANOVA) of observing ΔM under the MCE hypothesis. Table 2 also shows the probability correlation, $[r]$, described above, its degrees of freedom, $[df]$, its associated probability, $[p(r)]$, the slope, $[Slope(r)]$, of the regression line, and the overall p-value achieved by each analyst for the series of 45 remote viewings.

UNCLASSIFIED

UNCLASSIFIED

Table 2

(U) VIEWER NO. 454 RESULTS CODED BY FOUR ANALYSTS

Analyst (j)	ΔM	$p(\Delta M)$	r $(p_j - p_{374})$ vs. p_j	df	$p(r)$	Slope(r)	Overall Series p
374	0.019	0.441	-	-	-	-	0.316
432	0.004	0.881	0.421	43	0.0040	0.326	0.702
579	0.007	0.768	0.555	43	0.0001	0.509	0.867
642	-0.012	0.573	0.552	43	0.0001	0.558	0.909

UNCLASSIFIED

(U) All three analyst candidates produced highly significant positive correlations between their p-values and the difference between their p-values and 374's p-value. This indicates that the literal and interpretive analysts will tend to provide more similar p-value estimates as the quality of the data improves.

(U) Another indication of the difference between a literal and interpretive analyst is ΔM . Even the most interpretive analyst did not find significance in this data set--i.e., $p(\Delta M) = 0.44$. Figure 3 demonstrates the effect on p-value estimates of the same data for a literal and an interpretive analyst from a ΔM perspective. An analyst with a large and positive ΔM will observe a larger number of significant sessions (i.e., that portion of the curve labeled "Matches" above M_{kk} --the critical value of the figure of merit) than an analyst with a small, or negative, value of ΔM . Thus, an optimal way of selecting analysts is to choose those with larger values for ΔM and smaller values for Slope(r).

UNCLASSIFIED

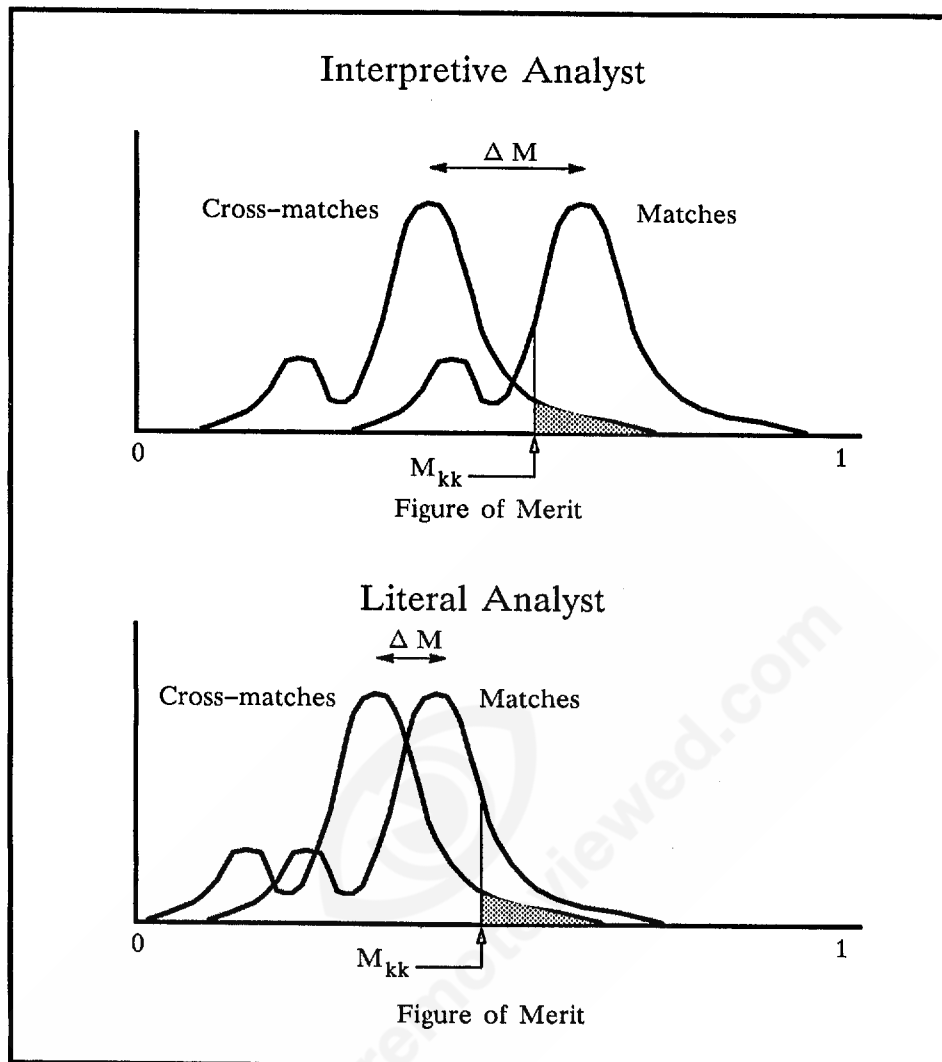
UNCLASSIFIED

FIGURE 3 (U) COMPARISON BETWEEN TYPES OF ANALYSTS

B. (U) Response Definition: Descriptor List Formulation**1. (U) Novice Response Descriptor List**

(U) A *post hoc* examination of the FY 1986 novice RV transcripts has resulted in a summary list of responses that were considered by the analysts to be the most troublesome to interpret within the highly specific framework of a twenty-bit descriptor list (see Table 1). In the RV training paradigm currently used by novices, interpretation by the viewer is largely discouraged. As a result, concrete words, such as *city*, *lake*, *tree*, or *boat*, are often labeled as analytical overlay (AOL) and must be discarded from the analysis by definition. Abstract

UNCLASSIFIED

UNCLASSIFIED

(U)

(and less interpretive) descriptions, such as *rounded*, *vertical*, and *wavy*, however, are typically encouraged and are commonly found in transcripts. The addition of abstract descriptors, therefore, would enable analysts to quantify the contents of a transcript literally, without having to make the considerable interpretive leap from nebulous responses to highly specific, concrete descriptors. Table 3 summarizes some of the more common abstract descriptions gleaned from novice responses, and provides suggestions for candidate abstract descriptors for incorporation into future lists.*

(U) It has yet to be determined how abstract and concrete descriptors will be structured within a given list--e.g., their interdependence could be either hierarchical in nature or configured along the lines of semantic networks. Whatever the mathematical formalism, it is anticipated that the addition of abstract descriptors will alleviate much of the burden of novice response interpretation for analysts.

2. (U) Advanced Response Descriptor List

(U) An FY 1986 experiment consisting of 12 outbound sessions was performed in which an advanced remote viewer (No. 342) was permitted, in an unsupervised fashion, to create his own descriptor list in advance of the experiment. The viewer was told only that the experiment was to be of the outbound beacon variety using San Francisco Bay Area targets and that his list should consist of approximately 20 to 30 descriptors. He was given the novice RV descriptor list as a template (see Table 1). The hypothesis under test was that the viewer, himself, would be most knowledgeable about his internal perceptions and would therefore be most qualified to objectify these perceptions in the form of his own "personalized" descriptors.

* (U) Table 3 is not meant to be an exhaustive list of potential abstract descriptors. It is merely meant to be illustrative by highlighting some of the more commonly encountered novice RV responses.

UNCLASSIFIED

UNCLASSIFIED

Table 3

(U) CANDIDATE ABSTRACT DESCRIPTORS FOR NOVICE RESPONSES

Typical Novice Responses		Suggested Abstract Descriptor
Category of Response	Actual Responses	
Patterns--curved, circular, rounded	Curved, circular, circle, oval, ellipse, round, wavy, rolling, contours, contoured, sloping	Are patterns of round, curved, or circular lines significant at the site?
Patterns--straight, angled	Straight, angled, parallel, horizontal lines, verticality, vertical lines, vertical objects, diagonal	Are patterns of straight, angled, or parallel lines significant at the site?
Patterns--combined curved and straight	Cone	Are a mixture of curved and straight patterns significant at the site?
No discernible patterns	Irregular, shapeless, uneven, rough, bumpiness, rugged terrain, clusters, irregular blobs, irregular shapes	There are no significant patterns at the site.
Distinct boundaries	Areas of light and dark, light and dark contrast	Distinct boundaries between light and dark are significant at the site. or Contrasting areas of light and dark are significant at the site.
Unspecified (generic) water	Water, wavy, waves, rippling, water movement, water blue	Is water a significant part of the scene?

UNCLASSIFIED

(U) Table 4 provides a comparison of target and response codings (assigned by an analyst on a blind basis) for the target matched with its correct response (see Figure 4).^{*} This particular example was chosen because our subjective appraisal tells us that the quality of the response does not seem to be reflected in its overall p-value (0.3289). The analyst's coding of the response and target were evaluated on a *post hoc* basis and were found to be

^{*} (U) The abstract descriptors proposed by the viewer appear to hold promise for codifying some of the information contained in novice responses (see Table 3).

UNCLASSIFIED

UNCLASSIFIED

Table 4

(U) COMPARISON OF TARGET VS. RESPONSE CODING FOR "BAYLANDS" TARGET

Bit	Target Coding*	Reponse Coding*	Descriptor
1	0	0	There are no significant patterns.
2	1	1	Are patterns of straight, parallel, or angled lines significant at the site?
3	0	1	Are patterns of round, curved, or circular lines significant at the site?
4	0	1	Are a mixture of round and straight patterns significant at the site?
5	0	0	Is a significant part of the scene hectic, chaotic, congested, or cluttered?
6	1	1	Is a significant part of the scene clean, empty, or open?
7	0	0	Is a significant part of the scene inside?
8	1	1	Is a significant part of the scene outside?
9	1	1	Is water a significant part of the scene?
10	0	0	Is sculptured water a significant part of the scene (fountains, etc.)?
11	1	1	Is natural water a significant part of the scene (lakes, ponds, streams, etc.)?
12	1	1	Are buildings or other manmade structures a significant part of the scene?
13	0	1	Is a single structure a significant part of the scene?
14	1	1	Is/are functional (useful, moving parts, etc.) structure(s) at the site?
15	0	0	Is/are artistic (there to look at) structure(s) at the site?
16	0	0	Is a single color predominant at the scene?
17	1	0	Is foliage a significant part of the scene?
18	1	0	Is foliage natural in appearance at the scene?
19	0	0	Is foliage significantly sculpted, manicured, or pruned at the scene?
20	0	1	Is the scene predominantly void of foliage?
21	0	0	Is motion significantly important at the site?
22	1	0	Is ambient noise significant at the site?
23	0	0	Is noise generated by the target?
24	0	0	Is noise generated by people adjacent to the target?

* 1 = yes, 0 = no

UNCLASSIFIED

UNCLASSIFIED

UNCLASSIFIED

(U)

reasonable--i.e., only one or two bit assignments were arguable, and their reassignment would not have resulted in a significant p-value, which this session seemed to merit.

(U) With analyst error seemingly eliminated, attention was focused on the efficacy of the descriptor list itself. It was concluded that the list was deficient in its ability to capture certain kinds of information, which largely accounted for the perceived accuracy of the response, including:

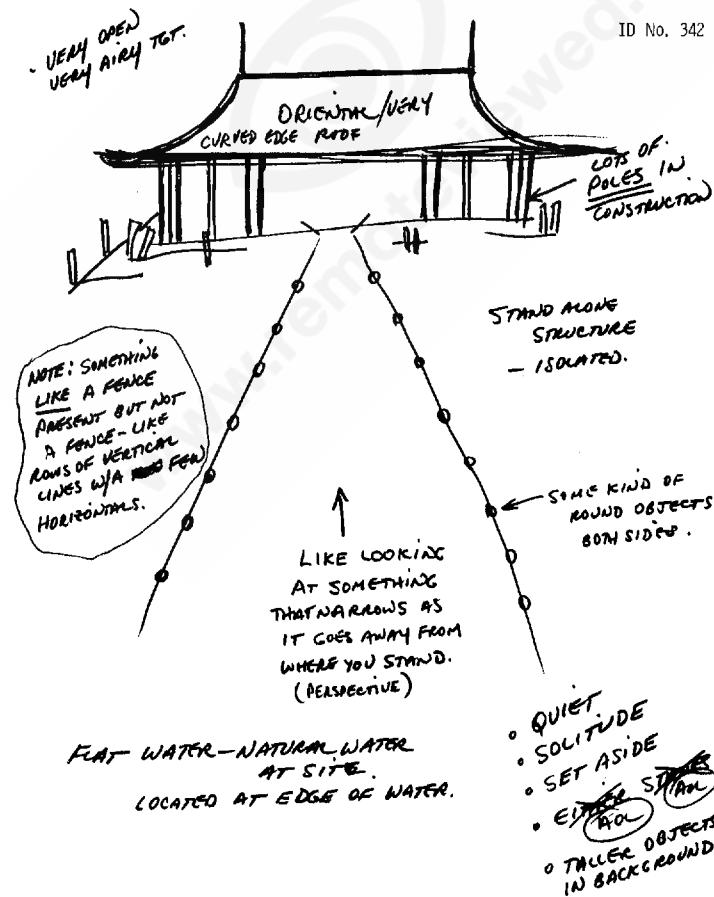
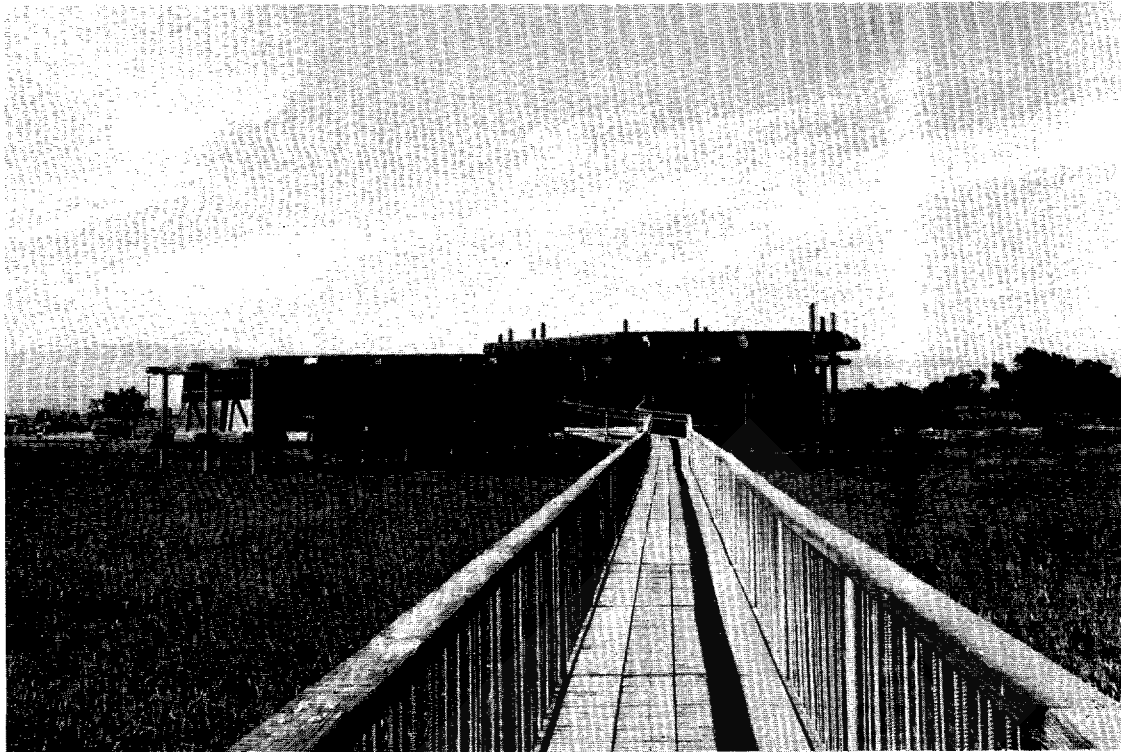
- The juxtaposition of elements (i.e., spatial relationships)
- The "novelty factor" of certain elements in the response
- The high specificity of named (or alluded to) target elements.

(U) No mechanism exists, as yet, within descriptor list technology for capturing the information contained in the spatial relationships between elements. Clearly, as in the example cited in Figure 4, this type of information can be very significant--i.e., the viewer drew his response from the beacon's actual perspective on the target. Spatial relationships appear to be particularly significant in advanced RV responses, in which complex, composite drawings are much more prevalent than in novice responses. This type of information may eventually be accounted for by employing new technologies such as rule-based expert systems (see Section IV.B.), which lend themselves well to recognition of juxtaposed elements.

(U) Another factor that is often thought to be important both in novice and advanced RV responses is--for want of a better term--the *novelty* or *strangeness* of an element in a response. An example of this in Figure 4 might include the odd shape of the structure's roof--i.e., a *curved* roof is a slight departure from normal expectation. The operative information here is embodied in the idea of "architectural oddity," a concept that is quite central to the target and is higher in information content than what is expressed by the various combinations of pattern and structure descriptors alone. Another example is the viewer's statement, "...like a fence present but not a fence...", which is a somewhat odd and uncertain phrase, but actually describes the catwalk guardrails in the target quite well. An experienced analyst might consider this latter type of information to be qualitatively better, because it represents a viewer's attempt to objectify his perception without succumbing to the pitfalls of analytical overlay (AOL). Analyst observations about response novelty can be systematically tested by devising an element-by-element analysis that can be applied across a wide qualitative range of responses. If the analyst "lore" is correct, it might be captured either by a "novelty" fuzzy weighting factor applied to descriptor lists or by expert systems capabilities.

UNCLASSIFIED

UNCLASSIFIED



UNCLASSIFIED

FIGURE 4 (U) BAYLANDS NATURE INTERPRETIVE CENTER, WITH RV RESPONSE

UNCLASSIFIED

UNCLASSIFIED

(U) The final factor that led to forfeiture of information entails the degree of specificity of the descriptor list. A highly detailed response of advanced RV quality will suffer if the descriptor list quantifying it cannot register detail. In the Figure 4 example, pertinent (and possibly *unique*) pieces of data--like poles being used in the structure's construction--are relegated to relatively nonunique bit categories such as "patterns of straight, parallel, or angled lines..." Additional highly specific, concrete descriptors (e.g., *poles*, *pathways*, etc.), therefore, are essential for descriptor lists quantifying high-quality RV data.

(U) The question of whether a viewer is better able to devise his own descriptor list remains unanswered. Subjectively, it is felt that analyst-derived lists have tended to be more concrete in nature and that the advanced series described here would have benefited from that kind of emphasis. There is also the additional fact that the viewer was not aware of the logical consistency rules governing descriptor formulation and application, and that an analyst's awareness of these procedural mechanics would have been beneficial to the construction of the list. The viewer's insights into abstract descriptor composition, however, were quite invaluable and hold important implications for novice list construction in particular.

C. (U) Target Definition: Implications for Target Pool Composition

(U) A few preliminary guidelines governing target pool composition have been distilled from two sources: (1) the opinions of RV monitors about the appropriateness of various kinds of RV targets for novice viewers, and (2) RV analysts' assessments about the difficulties encountered in using the twenty-bit descriptor list to score the 412 targets currently in the novice target pool.

(U) As a general rule, the current subjective consensus is that targets are inappropriate for training purposes if they exhibit any of the following qualities:

- They are contrary to the viewer's expectations.
- They are imbued with negative emotional impact.*
- They violate the "spirit" of the descriptor list's intended use.

*(U) Laboratory anecdotal evidence suggests that targets having negative emotional impact often result in psi-missing responses.

UNCLASSIFIED

UNCLASSIFIED

(U) A wide range of target types was used in the FY 1986 novice RV training series and have been subjectively determined on a *post hoc* basis to be of varying degrees of appropriateness for this task. Table 5 provides specific examples of how the current novice target pool may be problematic, given the following assumptions: (1) novice viewers had anticipated that targets for this series would consist of pictures taken from *National Geographic* Magazine featuring large, outdoor, gestalt scenes (e.g. cities, mountains, lakes) of roughly the same dimensionality; and (2) the twenty-bit descriptor list was appropriate for coding targets of this type *only* and nothing else—e.g., use of the list for technical sites or for targets featuring “unnatural” expressions of the bits was inappropriate.

(U) While it has yet to be determined empirically (i.e., by systematically examining figures of merit) whether these target types are *actually* problematic, it is currently the subjective opinion of the evaluation team that these kinds of targets would pose the greatest difficulties for novice viewers.

Table 5
(U) POTENTIAL PROBLEM AREAS FOR NOVICE TARGETS

Target Type	Specific Target Problem	Condition Violated		
		Viewer expectation	Emotional impact	Intended use of list
Close-up photo of a small feature (e.g., a flower, tree trunk, etc.)	Dimensionality	✓		✓
Reflections of rock formations in a still pool	Unusual perspective	✓		✓
Moon---off-earth photos		✓		✓
Sunken ship ruins	Underwater photos	✓		✓
Oil derricks	Technical site	✓		✓
Whale slaughter			✓	✓
Black & white photos				✓
Standing dead trees (without foliage)	Significant target element for which there was no descriptor	✓		✓
Photos with people and/or animals	Significant target element for which there was no descriptor			✓
Ornate interior of the Vatican during a ceremony	Too complex for novices	✓		✓

UNCLASSIFIED

UNCLASSIFIED

UNCLASSIFIED

IV RECOMMENDATIONS (U)

(U) The results of the FY 1986 evaluation effort have illuminated several areas of investigation that may hold promise for improving RV evaluation procedures. These areas include (1) identification of new descriptor lists that more accurately reflect target and response information, (2) implementation of enhancement techniques (e.g., fuzzy set theory) for attaining greater accuracy from descriptor lists, (3) systematic examination of inter-analyst reliability factors, and (4) development of new technologies (e.g., expert systems) for capturing analysts' insights with greater efficiency. Several parallel approaches, which address various aspects of these areas, have been targeted for preliminary research in FY 1987. These include:

- A "similarity" experiment (proposed by S. J. P. Spottiswoode) in which an attempt will be made to identify underlying semantic structures in remote viewing descriptions of target materials.
- An approach using artificial intelligence (AI) techniques (proposed by J. Vallee) for recognizing, analyzing, and describing target materials.
- In-house approaches directed at
 - Improvement of existing descriptor lists by incorporating more abstract descriptors into novice lists and more concrete descriptors into advanced lists.
 - Implementation of fuzzy set mathematical weighting factors into existing descriptor lists in an attempt to decrease their granularity.
 - Assessment of the relative merits of analyst-derived versus percipient-derived descriptor lists.
 - Identification of possible percipient-specific "ideogramic dictionaries," which might serve as prescriptive guides for the RV analyst.
 - Development of mission-specific descriptor lists (e.g., for technical sites).

Each of these approaches is outlined briefly below.

A. (U) Similarity Experiment

(U) According to the proposal submitted by consultant S. J. P. Spottiswoode, the proposed similarity experiment is aimed at improving existing evaluation techniques through

UNCLASSIFIED

UNCLASSIFIED

(U)

"...quantification of the informational content of transcripts on a small set of underlying semantic dimensions which might serve as basis vectors for the viewer's internal representation of the target. If such basis vectors can be found, complex constructs in the viewing data might be assembled by combining sets of data so expressed."⁹

(U) Using well-established techniques from the area of environmental psychology (which have been used to solve analogous problems in "normal" perception), the proposed experiment will attempt to isolate underlying semantic structures in RV perceptions. In essence, percipients will be asked to remote view—in sessions of two targets each—all possible pairs of targets and to estimate the similarity between targets in each pair. The data will be analyzed for important factors such as intersubject reliability, presentation order effect, and target pair effects. Multidimensional scaling (MDS) analysis* will then be applied to identify the underlying semantic dimensions. Identification of semantic structures would hold important implications for descriptor list development and possibly for identification of fundamental commonalities of perception across viewers.

B. (U) AI Techniques

(U) According to the following excerpt from the letter proposal submitted by consultant J. Vallee, the proposed AI approach

"...will seek to build an expert system for target recognition by analyzing the process that enables a human expert monitor to provide an interpretation of a remote viewing session, or a judge to match a given description with an actual target.

It is expected that a rule-based expert system can be developed in a series of iterations starting with the simple "twenty questions" framework already used in the project. Later this will lead to a fully-developed interactive model. We envision this "smart monitor" taking the analyst from simple scene and "gestalt" recognition to the detection of breaks, contradictions and, possibly, analytical overlays as well."¹⁰

(U) Assuming promising results with the initial task of target definition, it is anticipated that the expert system will ultimately be expanded to play a more active role in operational remote viewing sessions through on-line capture of respondents' ideograms and interactive analysis of their features.

* (U) It is beyond the scope of this discussion to describe the details of the MDS analysis.

UNCLASSIFIED

UNCLASSIFIED**C. (U) In-House Effort**

(U) SRI in-house approaches will focus on improvement and refinement of existing technologies. One potential line of inquiry would focus on identifying an appropriate set of abstract descriptors to add to the current novice list. This effort would be aimed at mitigating some of the inter-analyst reliability problems that have been encountered with novice data (see Section III.A.). An attempt would also be made to incorporate additional concrete descriptors into advanced lists.

(U) A second approach would endeavor to complete work begun in FY 1986--namely, a reanalysis of novice data using fuzzy mathematical weighting factors with the current list of twenty descriptors. The hypothesis is that the greater latitude afforded by fuzzy set membership values (as opposed to the "all or nothing" capability of the descriptors in their current configuration) will significantly decrease the granularity of the current list--i.e., will allow capture of more information. If the *post hoc* reanalysis yields promising results, the fuzzy set approach would be benchmarked on a blind basis against the current binary approach.

(U) A third effort would systematically evaluate the efficacy of viewer-derived versus analyst-derived descriptor lists. While viewers are sensitive to their internal perceptions, they are not cognizant of the requirements of the analytical procedures; conversely, the analyst is privy to the linguistic/analytical aspects of descriptor lists, but may be unaware of how to optimize a viewer's perceptions using descriptors. The hypothesis is that the combined insights of analyst and percipient will synergistically result in the optimal formulation of descriptor lists. One way to test this hypothesis would be to compare statistics across analyst-derived versus percipient-derived versus combined analyst/percipient-derived descriptor lists on the same set of RV data.

(U) A fourth approach would endeavor to take a retrospective look at viewers' ideograms and their possible range of meaning. If, for example, a viewer typically draws a "tic-tac-toe," cross-hatch-style ideogram, and *post hoc* analysis reveals that the drawing correctly corresponds to the presence of a city in targets 80 percent of the time, this information might be used to assign an *a priori* fuzzy weighting factor of 0.8 when the ideogram is encountered on a blind basis. An in-depth examination of a viewer's ideograms, therefore, might result in the development of viewer-specific, prescriptive guides for the assignment of fuzzy weighting factors in assessing RV responses. This type of information could also conceivably be automated and updated through iterative expert system capabilities.

UNCLASSIFIED

SECRET

(S) Finally, work will be initiated to develop mission-specific descriptor lists for technical site applications. It is anticipated that efforts along these lines will enable us to better address intelligence community requirements.

SECRET

SECRET

V CONCLUSIONS (U)

(U) The FY 1986 evaluation effort has resulted in (1) refinement and extension of current techniques, and (2) identification of candidate new technologies for preliminary research.

(U) The mathematical formalism for the current evaluation procedure--the figure of merit analysis--is well understood and stable. In addition to the system's ability to provide a reasonable assessment of remote viewing data, it has also provided a mechanism for systematic examination of inter-analyst reliability factors.

(U) The descriptor lists that currently form the basis for the figure of merit analysis have been evaluated on a *post hoc* basis. Preliminary observations indicate that lists designed for novice responses require greater *abstract* descriptor capability, whereas lists designed for advanced responses (i.e., higher-quality data) require greater *concrete* descriptor capability. It is anticipated that fuzzy set technology will assist in formalizing the interdependence between abstract and concrete descriptors, by providing a mathematical framework through which basis vector descriptors can be combined to form concrete descriptors.

(U) Research into new technologies for RV evaluation will begin in FY 1987. One of these approaches, the proposed "similarity" experiment, shows promise for identifying basis vector descriptors. A second approach, using rule-based expert systems, will explore a different dimension by endeavoring to capture RV analysts' expertise in codifying targets. Should this initial effort in artificial intelligence prove successful, it will be expanded to address the more difficult problem of response interpretation.

(S) It is hoped that this multifaceted approach to the refinement of RV evaluation procedures will result in increased capabilities for addressing the more complex problems of mission-oriented, operational RV.

SECRET

UNCLASSIFIED

REFERENCES (U)

1. Puthoff, H. E. and Targ, R., "A Perceptual Channel for Information Transfer Over Kilometer Distances: Historical Perspective and Recent Research," *Proceedings of the IEEE*, Vol. 64, No. 3 (March 1976) UNCLASSIFIED.
2. Targ, R., Puthoff, H. E., and May, E. C., 1977 *Proceedings of the International Conference of Cybernetics and Society*, pp. 519-529 UNCLASSIFIED.
3. May, E. C., "A Remote Viewing Evaluation Protocol (U)," Final Report (revised), SRI Project 4028, SRI International, Menlo Park, California (July 1983) SECRET.
4. May, E. C., Humphrey, B. S., and Mathews, C., "A Figure of Merit Analysis for Free-Response Material," *Proceedings of the 28th Annual Convention of the Parapsychological Association*, pp. 343-354, Tufts University, Medford, Massachusetts (August 1985) UNCLASSIFIED.
5. Fisher, R. A., "Statistical Methods for Research Workers," Oliver & Boyd (7th ed.), London, England (1938) UNCLASSIFIED.
6. Edgington, E. S., "A Normal Curve Method for Combining Probability Values from Independent Experiments," *Journal of Psychology*, Vol. 82, pp. 85-89 (1972) UNCLASSIFIED.
7. Cooper, B. E., "Statistics for Experimentalists," pp. 219-223, Pergamon Press, Oxford, England (1969) UNCLASSIFIED.
8. Zadeh, L. A., "Fuzzy Sets versus Probability," *Proceedings of the IEEE*, Vol. 68, No. 3, p. 421 (March 1980) UNCLASSIFIED.
9. Spottiswoode, S. J. P., "Proposals for Experimental Studies on Semantic Structure and Target Probability in Remote Viewing," *Private Communication* (July 1986) UNCLASSIFIED.
10. Vallee, J., "Applications of Artificial Intelligence Techniques to Remote Viewing: Building an Expert System for Target Description," *Private Communication* (August 1986) UNCLASSIFIED.

UNCLASSIFIED

SECRET



WARNING NOTICE

**RESTRICTED DISSEMINATION TO THOSE WITH VERIFIED ACCESS
TO THE [REDACTED] PROJECT**

SG1A

**NOT RELEASABLE TO
FOREIGN NATIONALS**

SECRET