
CURRENT RESEARCH IN SOCIAL PSYCHOLOGY

<http://www.uiowa.edu/~grpproc/crisp/crisp.html>

Volume 11, Number 12

Submitted: June 7, 2006

First Revision: July 7, 2006

Accepted: July 8, 2006

Published: July 10, 2006

NINE PSYCHOLOGISTS: MAPPING THE COLLECTIVE MIND WITH GOOGLE

Jack B. Arnold

Saint Mary's College of California, Retired

ABSTRACT

All pairs of names generated by the individual names of nine historically important psychologists were submitted as queries to the Google search engine. The resulting page counts were used to generate similarity/dissimilarity indices that were submitted to both cluster analysis and multidimensional scaling. Both of the analyses separated the names into three distinct clusters that were easily associated with three historically important schools of psychology. The purpose of the study was to examine the idea that the world-wide-web contains latent structures of the sort made familiar awhile ago by Charles Osgood. Earlier related data, gathered by the author in the last three or four years, is summarized and presented as further evidence that Osgood meaning may be latent in the world-wide-web. Questions regarding appropriate indices of similarity/dissimilarity and problems of the reliability and validity of these procedures and their results are discussed. Evidence is presented for all of these qualities in the results of this study. Finally, it is demonstrated that at least one of Osgood's connotative Semantic Differential factors is hidden in the structure of the world-wide-web.

INTRODUCTION

With the amount of discussion that has been generated lately about Google and the, so called, semantic web or Semweb (for example, Ford, 2002a, b), I am surprised that some psychologist has not noticed and remarked publicly that the coarse structure of the world-wide-web may be hiding semantic structures in which psychologists have, in the past, shown great interest. Here I am using the term "semantic structure" in an old fashioned sense that, I think, would warm Charles Osgood's heart (Osgood et al., pp. 25-124 and elsewhere). Events like the following sometimes occur during a Google search: (a) The search term, {+Freud +Jung}, finds 6760 documents. (b) The search term, {+Freud +Rembrandt}, co-occurs in only 645 documents. [1] Having noticed several of these occurrences, I am further surprised that our hypothetical psychologist wasn't moved to infer that those numbers might index the similarity of meaning of the co-occurring concepts (again using Osgood's sense of meaning). At least a couple of computer scientists (Cilibrasi and Vitynai, 2005) have made this inference and I suggest that much of their paper is likely to be both accessible and interesting to psychologists.

Osgood is remembered for, at last, three significant contributions to the behavioral/social sciences. He developed a mediational theory of meaning based on conditioning (Osgood, Suci & Tannenbaum, 1957, pp. 5-9). He developed a construct, the "semantic space," where the meaning of a concept is represented as a vector in a space spanned by an unknown (but discoverable) number of dimensions of meaning (Osgood, Suci & Tannenbaum, 1957, p. 25). And he developed a method, the Semantic Differential, for discovering the location, or meaning, of the concepts within the semantic space (Osgood, Suci & Tannenbaum, pp. 18-30, and elsewhere). The mediational theory is rarely mentioned now. However, the spatial model and the measurement method are alive in several areas of behavioral/social science research. I did a Google search for the term "+Semantic Differential" and found 149,000 references. The Google Scholar service (Google website, n.d) classified 920 of these as scholarly works published between 2004 and 2006.

To measure meaning, Osgood asked subjects to rate words (or more complex concepts) on Likert type scales defined by polar opposite adjectives, like good/bad or weak/strong. Low ratings were to indicate that the word was better characterized by the adjective defining the low end of the scale. High ratings were to indicate that the word was better characterized by the adjective defining the high end of the scale. Osgood called the activity of rating concepts on scales "differentiating" the concept's meaning (Osgood, Suci & Tannenbaum, 1957, p. 26). The number of possible SD scales is, clearly, very large. However, generally only a subset of the possibilities is used in a particular study. Each SD scale was conceived as a dimension in the semantic space. However, as most of the scales are correlated to some degree with other scales, some method is required to assess the number of actual independent qualities represented. Osgood and his group used factor analysis to determine the number and of independent dimensions. The number of salient factors was usually discovered to be three, a finding closely replicated over many separate studies involving numerous classes of concepts and numerous cultures and languages (Osgood, Suci & Tannenbaum, 1957, pp. 169-188). The most salient factor was usually one highly correlated with the good/bad scale when that one was included in the battery. The other two most salient factors were one closely correlated with active/passive and one correlated with weak/strong. Osgood called these factors connotative dimensions of meaning as distinguished from denotative dimensions of meaning. It is clear that the SD method is not very sensitive to the dictionary-like aspect of meaning which the Osgood group calls "denotative meaning" (Osgood, Suci & Tannenbaum, 1975, pp. 321-325, and elsewhere).

Although I have great regard for the utility of the Semantic Differential method, it is Osgood's very general idea of a semantic space that causes me to remark that his heart would be warmed by the idea of a semantic web, and especially a semantic web whose secrets might be open to psychometric methods. However, the best method to use with Google may not always be the Semantic Differential. Even if a Google parallel to SD scaling were devised, that method would likely not be the most efficient way to measure denotative meaning. Denotative meaning may better be measured using techniques like cluster analysis and multidimensional scaling applied to direct measures of the similarity of concepts, measures that are not derived from ratings on intermediary SD scales. Examples are judgments of the degree of similarity/dissimilarity of members of pairs of concepts, the frequency with which members of a pair are confused, or the number of Google pages in which a pair of concepts co-occur. Osgood's SD method is like any other approach that uses factor analysis. It can only discover the factors present in the original variables, in this case the SD scale battery. A direct similarity/dissimilarity for a pair of concepts will presumably include all of the dimensions relevant to discriminating the members of the pair. In fact, researchers have had some success using multidimensional scaling to solve SD problems. Osgood reports a doctoral study by Rowan (Osgood, Suci & Tannenbaum, 1957, pp. 143-146). Rowan apparently found some correspondence between the SD factors and his multidimensional scaling results and additional multidimensional scaling results that were difficult to interpret. Other workers (Flavell, 1961; Arnold, 1963, 1971) have obtained results that suggest that direct similarity/dissimilarity measures contain information that may be explained only partially by the connotative (SD) factors.

Consider a likely small study proposal, perhaps part of an M.A. thesis. The researcher is interested in one possible effect of taking a course in the history of psychology. She presents a number of course graduates with all of the possible pairs from a list of names for nine important psychologists. Say, Freud, Adler, Jung, Wertheimer, Köhler, Lewin, Watson, Skinner and Thorndike. Her instructions to her participants are to rate each pair on a scale ranging from zero to one-hundred to indicate how similar the members of the pairs appear to be. For her analysis of the data, she might average the resulting set of similarity matrices and use one or more of the readily available methods for cluster analysis or multidimensional scaling to search for a potential latent structure in the minds of her participants. An obvious expectation, assuming that the history course had any effect, would be that the results would display three clusters: three gestalt psychologists, three psychoanalysts, and three behaviorists.

METHOD

Data

Now consider another small study, one that I actually did, that might have been done as a preliminary to the one just described. I used the Google search engine to report on the frequency of occurrence in its index of web documents for each of the nine names mentioned above and for the frequency of occurrence for each of the possible 72 orders of pairs of names. Order sometimes makes a difference with Google's results. I did not use just the last names as search terms, since that would no doubt have inflated the document counts with noise caused by non-relevant persons with the same last names as the persons of interest. The search terms I did use were +"Sigmund Freud", + "Carl Jung", +"Alfred Adler", +"John B. Watson", +"B. F. Skinner", +"E. L. Thorndike", +"Wolfgang Kohler", +"Kurt Lewin", and +"Max Wertheimer". The quotation marks indicate to Google that a phrase is to be taken as a single term. I appended the plus sign (+) to indicate to Google that I did not want to count pages that failed to include the term.

The choices for inclusion as members in the three categories were not hard to make. My primary criterion was that members be well known, at least by social /behavioral scientists, as members of their respective schools. The psychoanalytic group was the easiest to form. Freud was, of course the most obvious choice. Adler and Jung were nearly as obvious, evidenced, for example, by Watson and Evans' (1991, p. 561) chapter title, "Adler, Jung and the third generation of dynamic psychologists." This was preceded by a chapter devoted to Freud and the early development of psychoanalytic thought. Several of Freud's more intimate and loyal collaborators are mentioned in this earlier chapter, but none are as well known outside psychoanalytic circles as Adler and Jung. The discussion of dynamic psychology in Boring's classic *History of Experimental Psychology* (1950, pp. 706-704) shows the same general distribution of emphasis, as does another classic, *Brett's History of Psychology* (Peters, pp. 715-730). Selecting behaviorists was also an easy task. J. B. Watson is the acknowledged philosophical founder of the group (Boring, 1950, p. 641, 643), and B. F. Skinner may be the best known of any of the Behaviorists likely because of his very radical philosophical stance, his willingness to offer practical advice on anything related to behavior, and his association with behavioral therapy (Watson & Evans, 1991, pp. 488-490). My selection for the third behaviorist was E. L. Thorndike. His major contributions occurred before Behaviorism was officially established, but he is discussed by most historians as an especially important precursor to Behaviorism (Watson & Evans, 1950, 471-472; Peters, pp. 694-699). The choices to represent the Gestalt school were, also, nearly automatic. Wertheimer and Köhler were two of the three founding members of the Gestalt movement, and Kurt Lewin, although not usually considered a founder of the movement, was an especially important social psychological theorist, noted for being profoundly influenced by the Gestalt movement (Watkins & Evans, pp. 501-518; Boeree, 2000).

Table 1 shows the matrix of occurrence/co-occurrence frequencies. Data for the 72 permutations of the nine names taken two at a time are displayed in off-diagonal cells. The numbers in the main diagonal are the page counts for the individual names. Notice that the matrix is not perfectly symmetrical that order, row name first and column name last, does sometimes make a difference. Queries taken at different times don't always yield exactly the same results either. These numbers were accumulated, one at a time, over the course of two days, 2/28/06 to 2/29/06. Changes in the numbers do occur over time, but I have not noticed any material changes over a period of a day or so. The numbers shown are the obtained frequencies divided by 1000.

Table 1. Google Page Counts/1000 for Pairs of Psychologist Names

	1.	2.	3.	4.	5.	6.	7.	8.	9.
1 Freud	1710.000	119.000	33.700	0.930	39.600	0.235	0.193	0.659	1.180
2 Jung	119.000	699.000	29.300	0.813	14.700	0.085	0.117	0.541	0.258
3 Adler	33.700	27.000	161.000	0.571	0.832	0.250	0.173	0.431	0.347
4 Watson	0.945	0.816	0.571	36.400	13.200	0.243	0.213	0.188	0.431
5 Skinner	39.500	14.600	0.831	13.200	342.000	0.407	0.307	0.605	0.537
6 Thorndike	0.234	0.083	0.250	0.245	0.407	13.600	0.050	0.082	0.117
7 Köhler	0.192	0.130	0.173	0.213	0.307	0.049	11.000	0.168	0.405
8 Lewin	0.653	0.537	0.431	0.188	0.605	0.080	0.168	93.300	0.540
9 Wertheimer	1.180	0.257	0.347	0.431	0.537	0.117	0.405	0.540	22.300

Data Analysis

I examined the data with both hierarchical clustering and multidimensional scaling procedures. However, before these methods could properly be employed, some pre-processing was necessary. The individual concepts vary markedly in the numbers of page counts they elicit, so the raw counts need to be normalized, otherwise the apparent similarities among the concepts are likely largely to be a function of individual concept page counts.

The steps in the pre-processing were determined by experience during three years of trial-and-error work with a number of data sets like Table 1 coupled with a small amount of very loose theory. First, the matrix was converted from asymmetric to symmetric by averaging the upper and lower triangles. The differences were not especially frequent nor were they usually large. It did not seem a stretch to assume that corresponding upper and lower numbers were estimating the same parameter. Second, the matrix data were used to calculate an index of similarity for each pair of concepts. There were a number of possibilities. I found the "cosine of pointwise mutual information," $cpmi$, (discussed by Terra & Clark 2002; Makkonen, Ahonen-Myka & Salmenkivi 2002; and others) to be convenient and to work well with the Google data. I present, below, the Makkonen, *et al.* formula, with slightly edited notation: $cpmi(A,B) = n_{A\&B} / \sqrt{n_A * n_B}$, where $n_{A\&B}$ is the number of co-occurrences of concept A and concept B and $n_A * n_B$ is the product of the number of individual occurrences of concepts A and B respectively. Relative to the matrix in Table 1, $n_{A\&B}$ corresponds to off-diagonal elements while n_A and n_B would each be main diagonal elements. As far as I have been able to determine, $cpmi$ is more often used to index the similarity of documents based on the frequency of the co-occurrence of concepts rather than the similarity of concepts based on their co-occurrence in documents, as I have done here. This index is, clearly, a kind of correlation. The maximum value is 1.0 in the case of perfect overlap, and the minimum value is 0.0 in the case of no overlap. Note, also, that the divisor on the left satisfies the need to normalize the effects of differences between n_A and n_B . [2] I did consider, and experiment with, more common measures, like the phi-coefficient, but to calculate phi one needs a frequency for a base population, and it is not at all clear what that number should be in the present context. (See, however, the discussion of this by Cilibrasi & Vitanyi, 2005.)

I used the non-metric, hierarchical clustering procedure described by Marascuilo and Levin (1983, pp. 254-258), to cluster analyze the $cpmi$ matrix. The $cpmi$ values tend to be small, even though based on large overlap frequencies. Non-metric clustering is based on the rank order rather than absolute values of $cpmi$ and tends to show clusters based on relative similarity rather than absolute similarity.

I programmed Marascuilo and Levin's instructions into a Microsoft Excel macro with the Visual Basic for Applications programming language (see, for example, Walkenbach, 1997).

For multidimensional scaling I used the *cmdscale* package from the R statistical environment (The R Development Team, 2005, pp. 868-869). This R package features the classical, metric, multidimensional scaling method developed by Torgerson (1958, pp. 247-297), rather than the newer methods introduced by Shepard (1962a, b) and Kruskal (1964a, b). The older method, when used carefully, is more likely to produce more tightly constrained results than the newer methods—a better sense of the dimensionality of the concept space and the importance of the several

dimensions. This is especially so for smaller sets of to-be-scaled objects. The non-metric methods often require the restraints imposed by a large number of dissimilarities to produce satisfying results. The classical method does require the assumption of Euclidian distances---an assumption that is supported by the results of this study.

My intuition is that, as a first approximation, a random population of squared Euclidian distances is distributed as Chi-square. If the cpmi function of the Google frequencies can be assumed to approximate Chi-square probabilities, then the square root of the associated Chi-square can be interpreted as a distance. [3] I used the cpmi coefficient as the probability for a one degree of freedom chi-square, resulting in the inter-concept distance formula, $\text{Distance}(A,B) = \text{square - root}(\text{chi-square}(\text{cpmi}(A,B)))$.

RESULTS AND DISCUSSION

Non-metric Cluster Analysis

The most salient feature of the hierarchical cluster analysis was the partitioning of the nine names into three intuitively satisfying clusters. Freud, Jung and Adler form a distinct cluster. Watson, Skinner and Thorndike form another. And Köhler, Wertheimer and Lewin form a third. The concept names are clearly sorted into Psychoanalysts, Behaviorists and Gestalters respectively. Also, Psychoanalysts appear to be substantially different from the other two groups than they are from each other. This could be an academic vs. professional distinction that is correlated with psychological school, but, since this was not predicted, and since I don't have any very convincing argument for that finding, I am happy to simply note the result and move to another point.

I commented earlier that Thorndike and Lewin were, in some sense, special cases. Thorndike accomplished his work before Behaviorism was officially founded, and Lewin was not one of the founders of the Gestalt movement. I was, therefore, surprised to find them fitting so cleanly into their clusters. So, out of curiosity, I repeated the cluster analysis using a different similarity measure, the Jaccard coefficient (discussed in Makkonen, 2002 and elsewhere) and got roughly the same results just discussed. The informative differences were that Thorndike was very loosely connected to the Gestalt cluster while Lewin was very loosely connected to the Behaviorist cluster. Not terribly surprising. Similarities based on Google apparently connect Lewin and Thorndike only marginally to the Gestalt and Behaviorist schools, respectively.

Classical Multidimensional Scaling

A scree diagram was plotted, depicting the eigen values associated with each of the nine dimensions extracted from the distances calculated from the data in Table 1. One important aspect of those results is that all eigen values were positive. This is an important finding, since a problem haunting use of the classical method has been the necessity to estimate an additive constant to reduce the size and number of negative latent roots (Torgerson, 1958, pp. 268-277). This finding is consistent with the assumption that the distance function used here produces Euclidian distances. The scree diagram also suggests that no more than two of the dimensions extracted using classical multidimensional scaling are likely to be of interest. Table 2 shows the coordinates for the nine psychologists on each of the two dimensions.

Table 2. Coordinates for the Nine Psychologists on the First Two Principal Component Dimensions Extracted by Classical Multidimensional Scaling

	Dim 1	Dim 2
Freud	-1.463	0.095
Jung	-1.597	0.220
Adler	-0.933	0.720
Watson	0.471	-1.207
Skinner	-0.379	-1.234
Thorndike	0.937	-0.636
Köhler	1.189	0.336
Lewin	0.818	1.240
Wertheimer	0.957	0.466

When plotted against their coordinates on these first two principal component dimensions, the nine concept names separate neatly into the same three clusters demonstrated in the previously described cluster analysis. The centers of the clusters formed a rough triangle with the Psychoanalytic group more distant from the Gestalters and the Behaviorists than these last were from each other.

Some More General Considerations

Validity

The findings described so far will be of general interest to psychologists only to the degree that they illuminate psychological issues. The obvious first issue is whether or not these kinds of findings are valid and reliable pointers to human mental structures. So far, I have been satisfied that the findings associated with several sets of concepts have comported with my intuitions of what they should be. Over a period of about four years, I have used the methods described to analyze several concept sets: seventeen academic disciplines, eighteen famous names, the twelve most recent U.S. presidents and fifteen religious vocations. Early in the present exploration, I used a different similarity index on a set of colleges and universities. All of these analyses produced intuitively satisfying results.

I would not have been convinced of the psychological validity of these results if they had not been intuitively satisfying, but my intuition is not a substitute for an objective cross validation of the procedures. Such a cross validation might be provided with a favorable comparison of results from Google data with results from human judgments. I have not, so far, undertaken a special study of human judgments to compare with the Google results. I have done a crude comparison of results from a multidimensional scaling study published by Henley (1962, cited by Snodgrass, 1985, pp. 83-86) with a comparable Google study. Henley asked participants to rate pairs among thirty mammals for similarity and used multidimensional scaling to discern the dimensional structure of her data. I applied the same method to the same animals using Google co-occurrences to estimate semantic distance. A visual inspection of the plots of the first two dimensions from the two data sets showed considerable, but far from perfect, similarity. I did not have Dr. Henley's table of concept coordinates, at hand, so I was not able to calculate a numeric index of similarity. I also reanalyzed similarity ratings among a set of eighteen adjectives obtained in a study of mine

(Arnold, 1971) and compared the multidimensional scaling results of that analysis with results for the same adjectives with similarities estimated from Google counts. I used the `cancor` utility from R (The R Development Team, 2005, pp. 860-861) to calculate canonical correlations among the first five dimensions from the human data and the first five dimensions extracted from the Google data. The five correlations were .85, .72, .63, .24 and .04, suggesting a fair amount of correspondence on at least two dimensions. Visual comparison of the two sets of results was not especially impressive, but the adjectives were a more or less random collection not picked with any sharply defined structure in mind. However, the canonical correlations suggest that a rotation of the two sets of principal component axes toward maximum similarity would likely improve the comparison.

Reliability

One might expect a high degree of reliability or repeatability over time with Google counts. As the index grows larger, it is to be expected that the page counts will keep pace, but it is reasonable to expect that relative or proportionate counts should show some constancy. I repeated observations for two sets of data, separated in time by a number of months, in order to examine this question. Google co-occurrences for fifteen religious vocations were first gathered in three or four days ending August 23, 2003 and again around April 4, 2006. The total number of independent observations was 225 for each set. The raw page counts were transformed to natural logarithms to compensate for the considerable range and extreme positive skewness of the page counts. The Pearson r between the two sets of data was .94. Google co-occurrences for twelve recent U. S. presidents were first gathered around May 1, 2005 and again around April 4, 2006. The Pearson r for these data was also .94. Surprisingly, the average page counts were slightly smaller for the newer data than for the older data. This was the case for both religious vocations and presidents.

Methodological Details Using Google

Since this paper was written as much to propose an area of research and a method as to assert any particular psychological substance, I am moved to comment on some methodological details. First, most of the concept sets that I have discussed here have been as simple as I could find. They have been simple both in the sense that I was careful to use sets that were homogeneous with regard to content and level of abstraction. I have experimented with heterogeneous sets (e.g., hierarchies) with mixed and generally opaque results.[4] I have also tried to restrict the concept names to be as short as feasible in keeping with the intent of the analysis. In situations where one word names had a chance of producing clean and meaningful results, I used one word names. The problem with many multiple word designations is that there is often more than one possible choice for a given concept. This is clearly true, for example, with names of persons. Consider President Roosevelt. Should one use Franklin Roosevelt, Franklin D. Roosevelt, Franklin Delano Roosevelt, F. D. R. or President Roosevelt? Each of these (appropriately enclosed in parentheses) will generally result in a different page count. Clearly, richer results than mine are likely to be obtained with concepts that are more complex than mine were. However, great care and imagination are likely to be required in their selection if useful results are to be expected.

The mechanical procedures for gathering the data appear also to have some affect on the numbers obtained. In general, I have found that I get more easily interpretable results when I get the page counts by querying Google one concept or one pair of concepts at a time. This can take a long time

with a large set of concepts (e.g., one-hundred queries for ten concepts). However, setting up a spreadsheet with the labels for the pairs pre-prepared can speed up the process considerably. It is possible to automate the query process a number of ways, with spreadsheet macros or special programs or scripts, but my experience with these is that they generate many inconsistent returns and more than an occasional error message instead of a number. In any case, the Google organization is apparently not favorably disposed toward automatic data extraction that does not follow the rules for using the Google API (Calliahain & Dornfest, 2003, p. 110, p. 306). Google will usually extend permission for 1000 queries in a twenty-four hour period to the user of a properly sanctioned program, and it is apparently possible to get this limit raised sometimes (Google website, n.d.), but I don't really recommend using a fully automated process.

Another consideration, in this context, is the choice of what to enter in the main diagonal of a co-occurrence matrix. A simple program for an automated query will likely present Google with all pairs in a given set, including doubles like {cow, cow} and {pig, pig}. My experience has been that the doubles do not generally yield the same sized page counts as do the corresponding singles, like simply cow or pig. The doubles counts are usually smaller than the singles. This makes a certain amount of sense. Further, the logic of the similarity indices discussed here suggests that the page counts for pages that include a particular concept at least once is the one properly inserted in the main diagonal of our co-occurrence matrices.

Finally, one needs to consider the many refinements or restrictions that Google provides that have the potential to refine a search. I will not list them here. The reader will find these easily at the Google website (n.d.). I have used only one of these with any frequency. I have restricted my searches to English language documents. One result, no doubt, is the presence of a considerable amount of "noise" left in my data, some of which might have been eliminated with the imposition of further restrictions. One can also get interesting variations on results by selectively excluding and including certain contextual concepts. For example, a search term that includes +Freud, +Jung and -Psychology (minus Psychology) will produce a different page count than one that includes +Freud, +Jung and +Psychology. Either of these might return results that, depending on the intent of the researcher, are less "noisy" than simply including Freud and Jung. I have not examined this proposition in any systematic fashion. Note that entering these proposed context-modifying terms can add considerable time and considerable potential typing error to the querying process.

THE SEMANTIC DIFFERENTIAL

Since this effort was largely inspired by Osgood's early work, and since I have dropped his name several times, it seems fitting to demonstrate that his specific insight regarding the three or four dimensions of connotative meaning (for example, Osgood, 1957, pp 31-75, and elsewhere) extend to the semantic structure of the world wide web. To accomplish this I borrowed twenty-one of the one hundred and twelve emotion names studied by Morgan and Heise (1988). The criteria for my selections were that each name have an extreme average rating on the Morgan and Heise E (evaluation) scale of greater than 2.5 or less than -2.5, and that they not be hyphenated terms. The scale values were derived from ratings on Semantic Differential type scales. Morgan and Heise present separate averages for men and women, I combined these into weighted means for men and women before making my selections. The twenty-one names are displayed in Table 3. The terms

+good and +bad were added as anchor references for my analysis. They do not appear in the Morgan and Heise list.

I presented Google with all of the pairs of the twenty-three terms displayed in Table 3, and subjected the resulting page counts to the same multidimensional scaling procedure described above. Table 3 shows the coordinates for the emotion names on the first five principle axis dimensions.

Table 3. Coordinates for 23 Emotions on the First Five Principal Component Dimensions Extracted by Classical Multidimensional Scaling

	Dim 1	Dim 2	Dim 3	Dim 4	Dim 5
+bad	-0.493	0.417	-0.008	-0.310	0.1385
+good	-0.753	0.203	0.503	-0.613	0.0979
+proud	-1.490	0.357	-0.307	-0.106	-0.7153
+ecstatic	-0.560	-2.562	0.261	-0.898	1.5465
+happy	-0.800	-0.689	-0.365	-0.703	0.5476
+overjoyed	1.396	-2.914	-2.135	-0.388	0.8225
+passionate	-1.819	0.201	0.844	-0.038	-1.5180
+thrilled	-1.528	-0.657	-0.989	-2.214	-1.9051
+joyful	-0.422	-2.197	0.762	2.409	-0.6031
+pleased	-1.628	0.001	-0.773	-0.052	-0.8029
+cheered	0.256	-0.128	-1.661	0.918	0.9121
+outraged	-0.069	2.865	-2.063	0.066	0.9468
+horrified	1.472	1.083	-1.208	-0.618	0.0484
+mortified	3.127	-0.332	-0.815	1.332	-1.7932
+ashamed	0.241	0.855	-0.195	1.105	-0.3142
+terrified	1.441	0.586	0.340	0.054	-0.1252
+empty	-1.181	0.759	0.958	0.392	0.8674
+hurt	-0.403	1.017	0.026	-0.066	0.2530
+lonely	-0.698	-0.295	1.984	0.100	0.0782
+miserable	0.361	0.318	0.445	0.784	-0.4225
+depressed	-0.630	0.242	1.274	0.160	0.6513
+crushed	0.072	0.597	0.906	0.479	1.4687
+petrified	4.108	0.275	2.218	-1.791	-0.1796

The multiple correlation for predicting Morgan and Heise E values from these scaled coordinates is .878. It is apparent that almost all of this is accounted for by Dimension 1 and Dimension 2.

Apparently, at least, one Semantic Differential factors, evaluation, does contribute to the structure of the web. It is not clear, from my analysis, how much activity and potency are present. The Morgan and Heise A and P ratings for the twenty-one concepts I used have high multiple correlations when predicted from the data in Table 3, but, for these twenty-one concepts, E, P and A all show considerable intercorrelation. For that reason, independent tests for E, A and P are not really feasible with my data.

REFERENCES

Arnold, J. B. (1963). A test of some unifying assumptions for meaningfulness and meaning. Unpublished doctoral dissertation, University of California, Berkeley.

Arnold, J. B. (1971). A multidimensional scaling study of semantic distance [Monograph]. *Journal of Experimental Psychology*, 90, 349-372.

Boeree, C. G. (2000). *Gestalt Psychology*. Retrieved 15 June 2006 from <http://www.ship.edu/~cgboeree/gestalt.html>.

Boring, E. G. (1950). *A history of experimental psychology* (3d ed.). New York: Appleton-Century-Crofts.

Calishain, T. & Dornfest, R. (2003). *Google Hacks*. Sebastapol, CA, O'Reilly.

Cilibrasi, R. & Vitanyi P. (2005). Automatic meaning discovery using Google. *arXiv:cs.CL/0412098 v2 15 Mar 2005*. Retrieved March 15, 2005 from http://arxiv.org/PS_cache/cs/pdf/0412/0412098.pdf.

Flavell, J. H. (1961). Meaning and meaning similarity II: The semantic differential and co-occurrence as predictors of judged similarity of meaning. *Journal of General Psychology*, 64, 321-335.

Ford, P. (2002a). August 2009: How Google beat Amazon and Ebay to the Semantic Web. *Ftrain*. Retrieved March 29, 2006, from http://www.ftrain.com/google_takes_all.html.

Ford, P. (2002b). A bit of commentary on Google and the Semantic Web. *Ftrain*. Retrieved March 29, 2006 from http://www.ftrain.com/google_semweb_commentary.html.

Google website (n. d.). Google web APIs (beta), frequently asked questions. Retrieved April 5, 2006 from http://www.google.com/apis/api_faq.html#gen7.

Henley, N. M. (1962). A psychological study of the semantics of animal terms. *Journal of Experimental Psychology*, 93, 366-378.

Kruskal, J. B. (1964a). Multidimensional scaling by optimizing goodness of fit to a monotonic hypothesis. *Psychometrika*, 29, 1-27.

Kruskal, J. B. (1964b). Nonmetric multidimensional scaling: A numeric method. *Psychometrika*, 29, 115-129.

Osgood, C. E., Suci, G. J. & Tannenbaum, P. H. (1957). *The measurement of meaning*. Urbana: University of Illinois Press.

Peters, R. S. (Ed.). (1962). *Brett's history of psychology*. Cambridge: The Massachusetts Institute of Technology.

Makkonen, J., Ahonen-Myka, H. & Salmenkivi, M. (2002). Applying semantic classes in event detection and tracking. Retrieved March 3, 2006 from <http://www.cs.helsinki.fi/u/jamakkon/papers/icon02.pdf>.

Marascuilo, L. A. & Levin, J. R. (1983). *Multivariate statistics in the behavioral sciences*. Monterey, Ca.: Brooks/Cole.

Morgan, R. L. & Heise, D. (1988). Structure of emotions. *Social Psychology Quarterly*, 51, 19-31. Retrieved May 4, 2006 from <http://www.indiana.edu/~socpsy/papers/MorganHeise.pdf>

The R Development Team (2005). *R: A Language and Environment for Statistical Computing Reference Index*, pdf. Retrieved April 6, 2006 from <http://www.r-project.org/>.

Shepard, R. (1962a). The analysis of proximities: Multidimensional scaling with an unknown distance function. *Psychometrika*, 27, 125-139.

Shepard, R. (1962b). The analysis of proximities: Multidimensional scaling with an unknown distance function. *Psychometrika*, 27, 219-246.

Snodgrass, J. G., Levy-Berger, G. & Haydon, M. (1985). *Human Experimental Psychology*. New York, Oxford University Press.

Terra, E. & Clarke, C. (2003). Frequency estimates for Statistical Word Similarity Measures. *Proceedings of HLT-NAACL 2003 Main Papers*, pp. 165-172 Edmonton, May-June 2003. Retrieved March 21, 2006 from <http://www.cs.mu.oz.au/acl/N/N03/N03-1032.pdf>.

Torgerson, W. S. (1958). *Theory and methods of scaling*. New York: John Wiley & Sons.

Walkenbach, J. W. (1997). *Excel 97 programming for windows for dummies*. Foster City, CA: IDG Books.

Watkins, R. I., & Evans, R. B. (1991). *The great psychologists: A history of psychological thought*. New York: HarperCollins.

ENDNOTES

[1] I use curly brackets here to avoid using parentheses or quotation marks which have special uses in making Google queries.

[2] It is, actually, possible for the cpmi to assume a value greater than 1.00, since the kinds of data observed here are not necessarily perfectly consistent. The way Google works, the estimated co-occurrence of concepts A and B could be larger than the sum of the estimated individual

occurrences of A and B. However, the likelihood of this ever occurring appears to be very small. I have not encountered such an instance.

[3] Note that the same result can be obtained by assuming that D has a folded normal distribution.

[4] The methods devised by Cilibrasi, R. & Vitanyi P. (2005) seem better able to handle hierarchical relationships than the ones that I discuss.

AUTHOR'S NOTES

I am indebted to Douglas M. Arnold for undertaking some very creative Java programming that made my work immeasurably easier than it would otherwise have been. He also applied to the Google organization for me to get some necessary authorizations.

I have collected further data to investigate the relation between the nine psychologists analyzed in this paper and their respective schools of thought. Students interested in pursuing thesis projects using such data may contact me at jb_arnold@comcast.net for access to the data.

AUTHOR BIOGRAPHY

Jack B. Arnold retired two years ago as Professor of Psychology at Saint Mary's College of California, where he taught Experimental Psychology and Research Methods for forty years. He was chair of the psychology department for three of those years. Earlier publications appear in *Human Relations*, *The Journal of Personality and Social Psychology* and *The Journal of Experimental Psychology*. They deal largely with style of clinical description and the measurement of meaning. Email: jb_arnold@comcast.net.